



El Model PIO (Principis, Indicadors i Observables) *Salut*:

Un enfocament d'autoavaluació per al compliment dels estàndards ètics i normatius aplicables als sistemes d'intel·ligència artificial en l'àmbit de la salut.



El Model PIO (Principis, Indicadors i Observables) Salut: Un enfocament d'autoavaluació per al compliment dels estàndards ètics i normatius aplicables als sistemes d'intel·ligència artificial en l'àmbit de la salut.

CIP 681.3.01(IA):61 OBS

Lillo Campoy, Alexandra, autor

El Model PIO (Principis, Indicadors i Observables) Salut : un enfocament d'autoavaluació per al compliment dels estàndards ètics i normatius aplicables als sistemes d'intel·ligència artificial en l'àmbit de la salut / Alexandra Lillo Campoy, Albert Sabater Coll, Susanna Aussó Trias, Juan González Fraile, Arlet Brufau Centelles i Hèctor Garcia Morales. – Girona : Oficina Edicions UdG : Observatori d'Ètica en Intel·ligència Artificial de Catalunya, abril 2025.
Descripció del recurs: 7 maig 2025
ISBN 978-84-8458-702-6

I. Sabater i Coll, Albert, autor II. Aussó Trias, Susanna, autor III. González Fraile, Juan, autor IV. Observatori d'Ètica en Intel·ligència Artificial de Catalunya V. Brufau Centelles, Arlet, autor VI. Garcia Morales, Hèctor, autor
1. Intel·ligència artificial – Aspectes ètics i morals 2. Ciències de la salut – Innovacions tecnològiques 3. Llibres electrònics

CIP 681.3.01(IA):61 OBS



Aquesta obra es troba sota una llicència Creative Commons (Reconeixement – Compartir Igual). Aquesta llicència permet l'ús, modificació i distribució d'aquesta obra amb qualsevol finalitat, sempre i quan es reconegui als autors i, en el cas de que es realitzin modificacions o desenvolupaments a partir de la mateixa, l'obra derivada es compartirà amb la mateixa llicència.

Observatori d'Ètica en Intel·ligència Artificial de Catalunya (OEIAC)

Carrer de la Universitat, 10 – Edifici Econòmiques

Universitat de Girona 17003 Girona

Girona, abril 2025

ISBN: 978-84-8458-702-6 (Editor: Oficina Edicions UdG)

Contingut

1.	Introducció: què és el Model PIO salut?	3
2.	Per què un model PIO salut?	6
2.1.	Les particularitats de l'ús de sistemes d'IA a l'àmbit de la salut	6
2.2.	Els productes sanitaris com a sistemes d'alt risc	8
2.3.	L'ús de sistemes d'IA generativa per part del personal sanitari	10
3.	Objectius	14
3.1.	La utilitat d'afrontar dilemes ètics	14
3.2.	Els objectius del Model PIO Salut	15
4.	La persona pacient en el centre	17
4.1.	Els riscos derivats de la implementació de la IA	17
4.2.	Els drets de les persones usuàries dels serveis de salut	20
5.	A qui va dirigit?	23
5.1.	Els subjectes del PIO Salut	23
5.2.	Subjectes afectats pel Reglament europeu sobre IA	24
5.3.	Les 10 preguntes inicials o bàsiques	26
6.	Instruccions	27
7.	Qüestionari	30
	Principi 1: Transparència i explicabilitat	30
	Principi 2: Justícia i equitat	48
	Principi 3: Seguretat i no maleficència	55
	Principi 4: Responsabilitat i retiment de comptes	63
	Principi 5: Privacitat	81
	Principi 6: Autonomia	89
	Principi 7: Sostenibilitat	97
8.	Guia per contractes en matèria d'IA en l'àmbit de la salut	100
	ANNEX I: Resum de les preguntes del qüestionari	125
	ANNEX II: Catàleg legal i de recomanacions	137

1. INTRODUCCIÓ: QUÈ ÉS EL MODEL PIO SALUT?

Aquests darrers anys hem experimentat una acceleració en l'ús de dades massives i sistemes d'Intel·ligència Artificial (IA) en molts sectors d'activitat i, especialment, en l'àmbit de la salut. Si bé aquesta situació s'espera que sigui encara més transformadora i significativa per al benestar de les persones¹, també cal subratllar que qualsevol avenç tecnològic requereix d'uns usos responsables i, per tant, d'unes recomanacions clares i precises, comunament acceptades i directament aplicables per assegurar que els sistemes d'IA siguin segurs i de confiança.

Les guies existents, en especial les directrius com les desenvolupades pel grup d'experts de la Comissió Europea (HLEG)² han suposat una passa fonamental en aquesta direcció. Aquestes defineixen aspectes necessaris per a una IA de confiança, incloent-hi perspectives legals, principis ètics i aspectes sobre la robustesa dels sistemes d'IA. No obstant, moltes vegades trobem com les guies ètiques són d'alt nivell o no són prou concretes i, per tant, no sempre serveixen per ser utilitzades directament com a eines d'avaluació o de millora dels sistemes d'IA o de les dades que aquests utilitzen. Per aquest motiu, des de diversos àmbits es promou la necessitat de desenvolupar protocols, guies i eines, amb l'objectiu de què siguin aplicables per a desenvolupadors, desplegadors i usuaris³, fent que les directrius es transformin en un avenç tangible per una IA de confiança⁴.

Prenent com a referència iniciatives similars, dins del marc de l'Estratègia d'Intel·ligència Artificial de Catalunya sorgeix el Model PIO (Principis, Indicadors i Observables), desenvolupat per l'Observatori d'Ètica en Intel·ligència Artificial de Catalunya (OEIAC)⁵. El Model PIO es construeix al voltant d'una pregunta clau i efectiva que qualsevol persona que desenvolupa, gestiona o dirigeix un projecte on hi ha dades i sistemes d'IA pot entendre fàcilment: **Ho hem fet?** Al seu voltant, el Model PIO proposa una metodologia d'autoavaluació que inclou una sèrie de preguntes senzilles perquè creadors, gestors i usuaris es qüestionin l'ús ètic de dades i sistemes d'IA⁶.

El Model PIO es basa en set principis ètics fonamentals: (1) Transparència i explicabilitat; (2) justícia i equitat; (3) seguretat i no-maleficència; (4) responsabilitat i retiment de

¹ Catalonia.AI (2020). Estratègia d'Intel·ligència Artificial de Catalunya, Generalitat de Catalunya, Disponible a: <https://politiquesdigitals.gencat.cat/ca/economia/catalonia-ai/> (consultada el 17 d'abril de 2024).

² Comissió Europea (2019). Ethics guidelines for trustworthy AI. Disponible a: <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai> (consultada el 17 d'abril de 2024).

³ Zicari, R. [Roberto] V., Brusseau, J. [James], et al. (2021). On Assessing Trustworthy AI in Healthcare. Machine Learning as a Supportive Tool to Recognize Cardiac Arrest in Emergency Calls. *Frontiers in Human Dynamics*, 3. Disponible a: <https://doi.org/10.3389/fhumd.2021.673104> (consultada el 21 d'abril de 2024).

⁴ Karimian, G. [Golnar], Petelos, E. [Elena], i Evers, S. [Silvia] M.A.A. (2022). The ethical issues of the application of artificial intelligence in healthcare: a systematic scoping review. *AI Ethics* 2, 539–551 (2022). Disponible a: <https://doi.org/10.1007/s43681-021-00131-7> (consultada el 21 d'abril de 2024).

⁵ OEIAC (2022). El Model PIO (Principis, Indicadors i Observables): Una proposta d'autoavaluació organitzativa sobre l'ús ètic de dades i sistemes d'intel·ligència artificial. Revisió 2024: Model PIO Compliance i Dret. Disponible a: <https://oeiac.cat/> (consultada el 17 d'abril de 2024).

⁶ *Idem*.

comptes; (5) privacitat; (6) autonomia; i (7) sostenibilitat. La selecció d'aquests principis segueix les directrius del grup d'experts d'alt nivell en IA de la Comissió Europea⁷. L'objectiu del Model PIO és sensibilitzar els diferents agents involucrats en el desenvolupament i l'ús de dades i sistemes d'IA sobre la importància de considerar aquests principis ètics per minimitzar riscos i maximitzar oportunitats en el seu ús. A més, el Model PIO té com a objectiu identificar accions adequades per tal de valoritzar, recomanar i avançar en l'ús ètic-legal de dades i sistemes d'IA⁸.

Sota el paraigua del Model PIO, l'OEIAC està desenvolupant diferents eines d'autoavaluació que, seguint mateix format i metodologia, posen el focus en diferents aspectes relacionats amb l'ús de la IA per tal d'adaptar el model segons les necessitats. En aquest sentit el 2024 es va fer públic el Model PIO sobre Obligacions i Drets, al qual acompanya un informe homònim publicat el gener de 2025⁹. Aquest està dissenyat per poder-se aplicar a qualsevol tipus de sistema d'IA, i serveix com a referent per a confeccionar la resta de variacions, que es tractaran de models sectorials, això és, centrats en un àmbit concret.

Així és com sorgeix el Model PIO Salut, presentat en aquest informe. I és que l'àmbit de la salut, lluny de ser una excepció, està experimentant una de les acceleracions més significatives en l'ús de dades i sistemes d'IA. La prova és que les aplicacions d'IA s'utilitzen àmpliament com a eines de diagnòstic clínic, per portar a terme tractaments o per fer el seguiment de pacients en temps real. Malgrat que la integració de la IA promet una precisió millorada, intervencions més ràpides i una atenció més personalitzada, d'aquesta nova realitat també sorgeixen riscos particulars que requereixen d'un abordament especialitzat, com s'explicarà en seccions posteriors.

Així doncs, el Model PIO Salut promou una revisió a través d'una interrogació detallada i continuada d'aspectes diversos, de possibles errors i de resultats inesperats o de risc potencial, abans i durant el desplegament d'un sistema d'IA. En aquest sentit, el model d'autoavaluació i les eines associades, anàlogament a altres avaluacions com les auditories¹⁰, proposen una metodologia exhaustiva per treballar en el desenvolupament d'una IA de confiança. Com descrivim més endavant, el procés d'autoavaluació requereix diferents usuaris i perspectives, per aprofitar habilitats tècniques, així com el coneixement del context, per poder així anticipar-se millor als efectes potencials un cop el sistema es desplega i les vulnerabilitats poden ser exposades de manera més evident.

⁷ Comissió Europea (2019) Ethics guidelines for trustworthy AI. *Op. cit.* 2.

⁸ OEIAC (2022). El Model PIO (Principis, Indicadors i Observables). *Op. cit.* 5.

⁹ OEIAC (2025). El Model PIO (Principis, Indicadors i Observables) sobre Obligacions i Drets: un enfocament d'autoavaluació per al compliment de la regulació en dades i sistemes d'intel·ligència artificial en el marc de la Unió Europea. Disponible a: <https://dugi-doc.udg.edu/handle/10256/26104> (Consultada el 16 de gener de 2025).

¹⁰ Liu, X. [Xiaoxuan], Glocker, B. [Ben], McCradden, M. [Melissa] M., Ghassemi, M. [Marzyeh], Denniston, A. [Alastair] K., i Oakden-Rayner, L. [Lauren] (2022). The medical algorithmic audit. *The Lancet. Digital health*, 4(5), e384–e397. Disponible a: [https://doi.org/10.1016/S2589-7500\(22\)00003-6](https://doi.org/10.1016/S2589-7500(22)00003-6) (consultada el 27 d'abril de 2024).

En aquest punt cal recordar que un dels principals factors que fan que molts sistemes d'IA desenvolupats mai arribin a tenir un impacte efectiu és el conegut com a “performance gap”: la diferència de rendiment entre el sistema d'IA quan es prova al laboratori i quan s'avalua ja en un entorn de desplegament real. Aquesta diferència és deguda a causes diverses segons cada cas, però té generalment un impacte notable en sectors com el de la salut. Necessitem eines i estratègies per identificar i prevenir aquesta diferència de rendiment, i el Model PIO Salut és una de les metodologies reconegudes per identificar el problema i trobar vies per mitigar-lo.

Els models en què estan basats els sistemes d'IA cobreixen moltes tipologies, i no només les xarxes neuronals i d'aprenentatge profund (Deep Learning), amb els quals darrerament es tendeix a associar qualsevol aplicació d'IA. El Model PIO és aplicable a qualsevol sistema d'IA, independentment de si es tracta d'un sistema d'aprenentatge automàtic simple, basat en regles fixes; si és un model lineal, o basat en arbres de decisió; o si, en contrast, és un sistema més complex com els de les grans xarxes neuronals profundes, incloent-hi els grans models de llenguatge. En aquest sentit, el Model PIO té l'avantatge de ser flexible i d'abast general.

A més de les consideracions ètiques i aspectes sobre la robustesa tècnica d'un sistema d'IA, el Model PIO sobre Obligacions i Drets va prendre com a principal fonament el contingut del Reglament europeu d'Intel·ligència Artificial, que entrà en vigor el 2 d'agost de 2024, més conegut com a RIA o AI Act. A part d'aquesta norma, es van prendre com a referència altres textos legislatius rellevants, com per exemple el Reglament general de protecció de dades (RGPD) o la Directiva 2016/2102 de 26 d'octubre de 2016 sobre l'accessibilitat dels llocs web i aplicacions mòbils dels organismes del sector públic, entre molts altres.

En el cas del Model PIO Salut, a les regles ja introduïdes al Model PIO sobre Obligacions i Drets, s'ha afegit regulacions específiques sobre els dispositius mèdics, en concret, el Reglament 2017/745, sobre els productes sanitaris, i el Reglament 2017/746 sobre productes sanitaris per a diagnòstic in vitro. Així mateix, s'han pres en consideració normes relatives als drets de les persones usuàries dels serveis de salut, com Llei bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informacions i documentació clínica; regles de funcionament del sistema sanitari com la Llei General de Salut Pública i la Llei General de Sanitat; i altres lleis i reglaments que afectaven d'una o una altra manera la prestació dels serveis de salut. D'aquesta forma, el Model PIO Salut, més enllà de d'avaluar si els sistemes s'adhereixen a les directrius ètiques determinades per autoritats en la matèria, permet fer una primera revisió sobre la conformitat del sistema respecte del marc regulador actual.

El present document està estructurat amb una primera secció sobre la necessitat d'un Model PIO específic en l'àmbit de la Salut, seguida per la presentació dels objectius del Model PIO Salut. A continuació, fem èmfasi en l'existència de diversos dilemes ètics i de com els afrontem. Tot seguit, exposem la importància de posar en el centre de tota decisió

i acció relacionada amb la intel·ligència artificial a la persona pacient, abordant els possibles impactes sobre el seu benestar i els seus drets (especialment aquells reconeguts al Conveni d'Oviedo). En la següent secció expliquem quins són els subjectes als quals està dirigit el Model PIO Salut. Més endavant s'enumeren les anomenades “preguntes 0”, això és, un recull de 10 interrogacions inicials i prèvies a l'autoavaluació. Seguidament, després d'aquestes explicacions es mostren les preguntes que comprenen el qüestionari del Model PIO Salut. Per finalitzar, es presenta una guia per a contractes relacionats amb sistemes d'IA en l'àmbit de la salut, basada en les consideracions ètic-legals que conformen el Model PIO Salut.

2. PER QUÈ UN MODEL PIO SALUT?

2.1. Les particularitats de l'ús de sistemes d'IA a l'àmbit de la salut

Per una banda, degut al gran impacte i les implicacions de l'ús de dades i els sistemes d'IA, les consideracions ètiques troben en l'àmbit de la salut un context especialment sensible. La vessant ètica ja és de vital importància a la pràctica clínica, on un gran nombre de decisions preses dins de l'àmbit clínic ha de passar una avaluació ètica i seguir protocols aprovats per comitès específics.

Però malgrat l'extensió de la recerca i la inversió en sistemes d'IA, la seva aplicació al món de la salut és encara, en contraposició, molt limitada. Per exemple, el nombre de sistemes de diagnòstic amb IA que han passat la validació necessària mitjançant assajos clínics i han estat exitosos és encara reduït¹¹. El món sanitari és conegut per ràtios altes de resistència en l'adopció d'innovació digital. Els sistemes de suport a la decisió no en són una excepció i molts d'ells mai superen la fase preclínica o de pilot¹².

Al món mèdic, a més, hi ha un canvi de paradigma des d'un enfocament hipertrofiat en el rendiment i la precisió dels nous sistemes d'IA, cap a cada vegada més un enfocament en l'anàlisi dels possibles errors i conseqüències associades¹³, convertint-se aquesta anàlisi, fins i tot, en un requisit dels protocols mèdics per reportar en estudis clínics¹⁴.

¹¹ WHO (2021). Ethics and governance of artificial intelligence for Health. Disponible a: <https://www.who.int/publications/i/item/9789240029200>. (consultada el 28 d'abril de 2024).

¹² Procter, R. [Rob], Tolmie, P. [Peter], i Rouncefield, M. [Mark] (2023). Holding AI to Account: Challenges for the Delivery of Trustworthy AI in Healthcare. *ACM Trans. Comput.-Hum. Interact.* 30, 2, Article 31, 34. Disponible a: <https://doi.org/10.1145/3577009> (consultada el 7 de maig de 2024).

¹³ Liu, X. [Xiaoxuan], Glocker, B. [Ben], McCradden, M. [Melissa] M., Ghassemi, M. [Marzyeh], Denniston, A. [Alastair] K., i Oakden-Rayner, L. [Lauren] (2022). The medical algorithmic audit. *Op. cit.* 10. En el mateix sentit: Vickers, A. [Andrew] J., i Elkin, E. [Elena] B. (2006). Decision curve analysis: a novel method for evaluating prediction models. *Medical decision making : an international journal of the Society for Medical Decision Making*, 26(6), 565–574. Disponible a: <https://doi.org/10.1177/0272989X06295361> (consultada el 5 de maig de 2024).

¹⁴ *Idem*.

En contrast amb altres sectors, a l'àmbit sanitari hi ha una visió unànimement acceptada: un dels objectius innegociables de qualsevol sistema és el benefici tant en la salut de les persones pacients, com de la població. Però, molt sovint, no és la mateixa persona pacient la usuària directa o la persona operadora d'un sistema d'IA, sinó que ho és el personal sanitari. Aquesta separació entre la persona usuària directe o la persona operadora i la principal persona afectada per l'ús de dades i sistemes d'IA, és un punt distintiu de l'àmbit de la salut que es propaga a cadascun dels reptes i requisits per una IA de confiança.

En particular, el principi de responsabilitat i retiment de comptes pren una importància capital, convertint-se en un aspecte prioritari per a les persones usuàries directes i operaris de la solució (personal mèdic) a l'hora de decidir si adopten l'ús i els resultats d'un sistema d'IA. I seguint la mateixa línia també el Principi de Transparència i explicabilitat. Això ha donat lloc a la recerca i proliferació de metodologies per l'anàlisi de decisions centrades al món clínic, de vegades massa reduccionistes i enfocades en aspectes concrets¹⁵. Malgrat que sempre s'estableix el benefici en la salut de la persona pacient com a prioritat innegociable, per avaluar de forma realista l'adopció d'un sistema d'IA s'ha de tenir en compte el benefici i cost no només per aquest, sinó també per al metge i responsable de la decisió, a part de per al centre mèdic al qual pertany, i al sistema sanitari en general.

Això només és possible amb una anàlisi en el context de salut del principi de Responsabilitat i retiment de comptes, i dels mecanismes existents associats, entorn de l'ús de dades i sistemes d'IA. Els diferents principis ètics per a una IA de confiança estan, en tot cas, fortament interconnectats i es troben d'entre els principis consensuats també pel grup d'experts de l'Organització mundial per la Salut sobre l'ús d'IA fiable en salut, on s'inclou de forma destacada: *"Assegurar la transparència, explicabilitat i intel·ligibilitat, i fomentar la responsabilitat i el retiment de comptes"*¹⁶.

La discussió de les interconnexions entre els principis ètics i, sobretot, l'anàlisi dels dilemes ètics que es plantegen en la creació i l'ús de dades i sistemes d'IA són necessàries per assolir solucions de confiança. No es tracta només d'afrontar una discussió teòrica sobre els dilemes i el seu interès acadèmic. És mitjançant la presa de decisions conscients enfront d'aquests dilemes en l'àmbit de la salut que com a creadors, desenvolupadors, usuaris, i receptors de l'impacte dels sistemes, podem avançar en l'ús ètic de dades i sistemes d'intel·ligència artificial. Aquesta idea vertebrava l'informe i motiva diverses preguntes del Model PIO Salut que presentem.

Finalment, la necessitat d'una visió tecnològica on la persona pacient tingui un rol central ha originat iniciatives prometedores al voltant de la idea que els sistemes d'IA siguin

¹⁵ Vickers, A. [Andrew] J., i Elkin, E. [Elena] B. (2006). Decision curve analysis: a novel method for evaluating prediction models. *Op. cit.* 13.

¹⁶ WHO (2021). Ethics and governance of artificial intelligence for Health. *Op. cit.* 11.

contestables per llur part¹⁷. En el camí cap a una IA de confiança, aquest subjecte hauria de poder contestar aspectes sobre l'ús de llurs dades personals, biaixos potencials del sistema d'IA, el rendiment del sistema d'IA o la divisió de contribució a una decisió entre el metge i el sistema d'IA¹⁸. En aquest procés destaca la necessitat d'involucrar els set principis, així com la utilitat de les metodologies basades en el qüestionament dels sistemes. No és sorprenent que la contestabilitat es consideri, de fet, un dels vuit principis ètics per treballs com el del govern d'Austràlia en el seu marc d'ètica per a la IA¹⁹.

El desenvolupament i l'ús de sistemes d'IA en l'àmbit de la salut presenta, per tant, particularitats que fan necessària l'existència d'un Model PIO específic pel sector sanitari, els seus usuaris i tots els agents involucrats.

2.2. Els productes sanitaris com a sistemes d'alt risc

El treball i discussions sobre les consideracions ètiques pel desenvolupament d'una IA de confiança d'aquests darrers anys, en particular en directrius com les del grup d'experts d'alt nivell de la Comissió Europea (HLEG), han representat un punt d'inflexió. També, tota una referència en el treball socio-tecnològic cap a l'assoliment de metodologies i procediments que promoguin el desenvolupament i l'ús de sistemes d'IA de confiança. Aquest tipus de directrius són, però, genèriques, i el mateix HLEG admet i motiva la necessitat d'adaptar-les i adequar-les a contextos d'ús concrets: *"diferents situacions donen lloc a diferents reptes"*²⁰.

Aquesta necessitat d'adaptació, concretament per l'àmbit de la salut²¹, ha motivat l'aparició de treballs específics²², encara incipients, com indica la contínua crida per a aprofundir en la seva recerca²³. Al món sanitari en particular està proliferant encara més ràpidament el desenvolupament de nous sistemes de presa de decisions automatitzades

¹⁷ Ploug, T. [Thomas]; Holm, S. [Søren] (2020). The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI. *Artificial Intelligence in Medicine*, Volume 107, 2020, 101901, ISSN 0933-3657. Disponible a: <https://doi.org/10.1016/j.artmed.2020.101901>. (consultada el 17 de maig de 2024).

¹⁸ *Idem*.

¹⁹ Govern of Australia (2019). The Artificial Intelligence (AI) Ethics Framework. Disponible a: <https://www.industry.gov.au/publications/australias-artificial-intelligence-ethics-framework>. (consultada el 17 de maig de 2024).

²⁰ Comissió Europea (2019). Ethics guidelines for trustworthy AI. *Op. cit.* 2.

²¹ WHO (2023). WHO outlines considerations for regulation of artificial intelligence for Health. Disponible a: <https://www.who.int/news/item/19-10-2023-who-outlines-considerations-for-regulation-of-artificial-intelligence-for-health> (consultada el 28 d'abril de 2024).

²² Bærøe, K. [Kristine], Miyata-Sturm, A. [Ainar], & Henden, E. [Edmund] (2020). How to achieve trustworthy artificial intelligence for health. *Bulletin of the World Health Organization*, 98(4), 257–262. Disponible a: <https://doi.org/10.2471/BLT.19.237289> (consultada el 28 d'abril de 2024).

²³ Selecció de crides (2023). Calls for Papers on trustworthy AI for health. Disponible a: <https://www.frontiersin.org/research-topics/51826/trustworthy-ai-for-healthcare>
https://www.mdpi.com/journal/sensors/special_issues/Healthcare_Human_Machine
<https://www.embs.org/ibhi/special-issues/trustworthy-and-collaborative-ai-for-personalised-healthcare-through-edge-of-things/> (consultada el 2 de maig de 2024).

basats en l'ús de dades i sistemes d'IA, sent l'àmbit de la salut un dels sectors amb més previsió de creixement durant els pròxims anys²⁴.

En aquest context, la necessitat d'assegurar amb el compliment dels criteris ètics han conduït a la regularització dels sistemes d'IA, donant lloc a iniciatives legislatives, com per exemple és la recentment adoptada Convenció marc del Consell d'Europa sobre IA i drets humans, democràcia i l'estat de dret. També destacà la Proposta de Directiva del Parlament Europeu i del Consell sobre l'adaptació de les normes de responsabilitat civil extracontractual a la intel·ligència artificial (AI Liability Directive), tot i que aquesta va ser retirada per la Comissió el passat 11 de febrer de 2025.

En aquest sentit, i en relació amb l'àmbit de la salut, destaca el Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial. El RIA estableix un sistema de categorització dels sistemes d'IA en funció del risc que suposen per a la seguretat, els drets humans, la democràcia i l'estat de dret i, evidentment, la salut.

Per aquest motiu, el RIA ubica els dispositius mèdics que incorporen intel·ligència artificial en la categoria de risc més alta a excepció d'aquells usos de la intel·ligència artificial que directament han estat prohibits. En conseqüència, aquests instruments hauran de complir amb les nombroses restriccions que els imposa el reglament sobre intel·ligència artificial de la UE, a més de qualsevol altra norma que els sigui aplicable.

Així doncs, amb el RIA augmenta la urgència per adaptar els principis ètics i també de les metodologies d'auditoria o d'autoavaluació, com el Model PIO Salut²⁵.

Nogensmenys, no tots els sistemes d'IA utilitzats en l'àmbit de la salut seran considerats als sistemes d'alt risc, sinó que l'article 6.1 del RIA limita aquesta categorització als sistemes afectats pel Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris i el Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro. Dit d'una altra manera, únicament als sistemes que poden ser considerats, bé un producte sanitari o un sanitari de diagnòstic in vitro, bé un component de seguretat d'aquests, d'acord amb les normes esmentades, els seran aplicables els requisits dels sistemes d'alt risc.

En concret, el Reglament 2017/745 defineix com a producte sanitari tot instrument, dispositiu, equip, programa informàtic, implant, reactiu, material o un altre article

²⁴ AI Act, Comissió Europea (2021). Proposal for a Regulation of the European Parliament and of the Council laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act) and amending certain Union legislative acts. Disponible a: https://eur-lex.europa.eu/resource.html?uri=cellar:e0649735-a372-11eb-9585-01aa75ed71a1.0001.02/DOC_1&format=PDF (consultada el 15 de març de 2024).

²⁵ Farah L. [Line], Murris J. [Juliette] M., et al. (2023). Assessment of Performance, Interpretability, and Explainability in Artificial Intelligence-Based Health Technologies: What Healthcare Stakeholders Need to Know, Mayo Clinic Proceedings: *Digital Health*, Volume 1, Issue 2, 2023, Pages 120-138, ISSN 2949-7612. Disponible a: <https://doi.org/10.1016/j.mcpro.2023.02.004> (consultada el 2 de maig de 2024). En el mateix sentit: MedTech Europe (2019). Trustworthy Artificial Intelligence (AI) in healthcare. Disponible a: https://www.medtecheurope.org/wp-content/uploads/2019/11/MTE_Nov19_Trustworthy-Artificial-Intelligence-in-healthcare-1.pdf (consultada el 2 de maig de 2024).

destinat pel fabricant a ser utilitzat en persones, per separat o en combinació, amb algun dels següents fins mèdics específics:

- Diagnòstic, prevenció, seguiment, predicció, pronòstic, tractament o Alleujament d'una malaltia;
- Diagnòstic, seguiment, tractament, alleujament o compensació d'una lesió o d'una discapacitat;
- Recerca, substitució o modificació de l'anatomia o d'un procés o estat fisiològic o patològic;
- L'obtenció d'informació mitjançant l'examen in vitro de mostres procedents del cos humà, incloent-hi donacions d'òrgans, sang i teixits;
- Productes de control o suport a la concepció.

Partint d'aquesta taxonomia, el Model PIO Salut ha estat dissenyat específicament com una eina d'autoavaluació dels sistemes d'IA que entrin en la definició de producte sanitari. Tanmateix, això no implica que altres sistemes utilitzats de manera més o menys directa en l'àmbit de la salut no comportin cap mena de risc a la seguretat humana. Per exemple, un sistema d'IA que atorga cites mèdiques a les persones que les sol·licitin segons la seva urgència no pot ser considerat un producte sanitari. Ara bé, aquesta eina pot ocasionar greus conseqüències si aquest filtratge es realitza de manera incorrecta, de manera que les persones poden no ser ateses d'acord amb la necessitat real de la seva condició, la qual cosa comporta un potencial perjudici per a la seva salut.

En situacions com l'esmentada, resulta recomanable sotmetre el sistema en qüestió a l'autoavaluació del Model PIO sobre Obligacions i Drets, el qual també es fonamenta en principis ètics i legislació aplicable als sistemes d'IA, i resulta aplicable a un ranc més ampli d'usos que el Model PIO Salut. Un dels trets característics del Model PIO Salut és que, al contrari que el Model PIO sobre Obligacions i Drets, no inclou una fase de Preavaluació prèvia per determinar el grau de risc del sistema abans de completar al qüestionari. Això és deu precisament al fet que aquest està enfocat a l'autoavaluació de productes sanitaris i, en conseqüència, es parteix de la idea que tots els sistemes avaluats recauran en la categoria d'alt risc.

2.3. L'ús de sistemes d'IA generativa per part del personal sanitari

Sense contradir tot el que es declara a l'apartat anterior, sorgeix una segona qüestió cabdal respecte a la introducció d'aplicacions d'IA en l'àmbit de la salut: l'ús de sistemes d'IA generativa. Som testimonis de com un creixent nombre de persones utilitza eines d'IA generativa, especialment grans models de llenguatge (LLM per les seves sigles en

anglès)²⁶, com per exemple Chat-GPT, per fer consultes relacionades amb el seu estat de salut i obtenir directrius de com actuar al respecte. Aquesta pràctica inclou persones que cerquen obtenir alguna mena de diagnòstic prèviament a consultar a una persona professional de la salut, o directament per evitar acudir als serveis de salut.

Però també existeixen situacions en les quals són les mateixes persones professionals de la salut les que empenen aquests sistemes per donar-se suport en la realització de les seves tasques, des de recopilar informació sobre una patologia en concret a introduir al sistema dades sensibles d'alguna pacient per després demanar-li que redacti un document a partir d'aquestes.

Els comportaments descrits comporten greus riscos per a la salut de les persones, puix que aquests sistemes, si bé poden transmetre una sensació de confiança i versemblança en els seus resultats, poden patir d'al·lucinacions, això és, sortides contextualment inversemblants, inconsistents amb el món real i incoherent amb la informació d'entrada²⁷. En aquests casos, la informació mèdica incorrecta o enganyosa generada per l'aplicació d'IA, pot afectar negativament la presa de decisions clíniques i els resultats dels pacients, la qual cosa condueix a una incorrecta atenció de la salut de les persones afectades, podent comportar importants conseqüències per a la seva salut²⁸.

El risc persisteix fins i tot respecte als models dissenyats específicament per al sector de la salut i entrenats amb vastos conjunts de dades que inclouen literatura mèdica i directrius ètiques²⁹. En relació amb aquests, les preocupacions es centren els mètodes d'entrenament, els presents biaixos de les dades d'entrenament i la capacitat dels models per adaptar-se a la constant evolució de la canviant disciplina de l'ètica mèdica³⁰. En aquest sentit, l'estudi "*Cross sectional pilot study on clinical review generation using large language models*", conclou que, en comparació amb les revisions clíniques escrites per humans, les revisions generades pels sistemes d'IA presenten importants deficiències en

²⁶ Els grans models de llenguatge són definits com "sistemes d'intel·ligència artificial dissenyats per aprendre gramàtica, sintaxi i semàntica d'un o més idiomes per generar un llenguatge coherent i rellevant per al context [...] com un tipus de sistema d'IA generatiu, els LLM creen nou contingut en resposta als comandaments de l'usuari en funció de les seves dades d'entrenament. Estan entrenats en grans quantitats de fonts de text (de milers de milions a milers de milions de paraules) d'una varietat de fonts". Lareo X. [Xabier] (2023). Large Language Models (LLM), a Supervisor europeu de protecció de dades (Ed.), *TechSonar Report 2023-2024*, 6-9. Disponible a: https://www.edps.europa.eu/data-protection/our-work/publications/reports/2023-12-04-techsonar-report-2023-2024_en (consultada el 31 de març de 2025).

²⁷ Ahmad, M. [Muhammad] A. [Aurangzeb], Yaramis, I. [Ilker], i Roy, T. [Taposh] D. [Dutta] (2023). Creating trustworthy llms: Dealing with hallucinations in healthcare ai. *arXiv Preprint arXiv:2311.01463*. Disponible a: <https://arxiv.org/abs/2311.01463> (consultada el 31 de març de 2025).

²⁸ Kim, Y. [Yubin], Jeong, H. [Hyewon], et al (2025) Medical Hallucination in Foundation Models and Their Impact on Healthcare, *MedRxiv*, 2025-02. Disponible a: <https://www.medrxiv.org/content/10.1101/2025.02.28.25323115v1> (consultada el 28 de març de 2025).

²⁹ Ning, Y. [Yilin], Teixayavong, S. [Salinelat], et al. (2024). Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist. *The Lancet Digital Health*, 6(11), e848-e856. Disponible a: <https://www.sciencedirect.com/science/article/pii/S2589750024001432> (consultada el 28 de febrer de 2025).

³⁰ Zohny, H. [Hazem], Mann, S. [Sebastian] P. [Porsdam], Earp, B. [Brian] D., i McMillan, J. [John] (2024). Generative AI and medical ethics: the state of play. *Journal of medical ethics*, 50(2), 75-76. Disponible a: <https://jme.bmj.com/content/medethics/50/2/75.full.pdf> (consultada el 14 de febrer de 2025).

quant a la fonamentació del text produït, la presència de biaixos, el grau d'innovació i de capacitat lògica i la qualitat general del resultat³¹.

Així doncs, sorgeixen dilemes ètics relacionats amb l'atribució de la responsabilitat per les decisions incorrectes preses arran del resultat generat per una IA i els danys i perjudicis que això pot causar a les persones pacients. Tot i la manca de regulació específica que contempli aquests supòsits, seria possible fins i tot considerar que la persona professional sanitària responsable de prendre la decisió o d'operar el sistema ha actuat de manera negligent, amb les conseqüències legals que se'n puguin generar d'aquest fet³².

En conseqüència, si una institució o persona professional de la salut decidís utilitzar un sistema d'IA generativa en el desenvolupament de les seves funcions, hauria de tenir present els riscos existents, i ser conscient en tot moment del fet que l'aplicació d'IA s'ha d'emprar com a una eina suplementària³³, i en cap cas com substituir el criteri professional de la persona usuària clínica, que haurà de validar els seus resultats.

No es pot deixar de subratllar que el fet d'introduir les dades d'una persona pacient a un sistema d'IA generativa per tal d'obtenir un resultat concret sobre el seu estat de salut, implica l'abocament de dades protegides que poden ser utilitzades per a l'entrenament i ajustament continuat del mateix sistema³⁴. Quan no es compta amb el consentiment exprés i informat de la persona afectada, aquesta pràctica suposaria una infracció de la normativa sobre protecció de dades (articles 6 i 7 el RGPD, sobre la licitud del tractament de les dades i les condicions per al consentiment), el que implica una vulneració del dret a la protecció de dades personals de la persona pacient (article 8 de la Carta de drets Fonamentals de la UE) així com el seu dret a la privacitat (article 10 de la convenció d'Oviedo i 18 de la Constitució Espanyola). És per aquest motiu que la implementació de sistemes d'IA generativa en l'àmbit de la salut pot suposar una greu amenaça per a la privacitat de les persones sobre les quals versen aquestes dades.

Aquest risc augmenta davant el cada cop més habitual accés en línia als historials mèdics. Si bé els historials d'accés en línia presenten l'avantatge d'oferir a la ciutadania una manera de consultar fàcilment informació relativa al seu estat de salut, el fet de situar aquesta informació a la xarxa porta aparellat riscos relacionats amb la privacitat de dades personals, fins i tot quan es prenen mesures de seguretat. En aquest sentit, Charlotte Blease adverteix l'accés en línia als historials clínics, combinat amb els reptes i les

³¹ Luo, Z. [Zining], Qiao, Y. [Yang], et al. (2025). Cross sectional pilot study on clinical review generation using large language models. *pj Digital Medicil*, 8, 170. Disponible a: <https://www.nature.com/articles/s41746-025-01535-z#citeas> (consultada el 25 de març de 2025).

³² Zohny, H. [Hazem], Mann, S. [Sebastian] P. [Porsdam], Earp, B. [Brian] D., i McMillan, J. [John] (2024). Generative AI and medical ethics: the state of play. *Op. cit.* 30.

³³ Mondal, H. [Himel] (2025). Ethical engagement with artificial intelligence in medical education. *Advances in Physiology Education*, 49, 163–165. Disponible a: <https://journals.physiology.org/doi/full/10.1152/advan.00188.2024>. (consultada el 25 de març de 2025).

³⁴ *Idem*.

pressions que enfronten els sistemes de salut, l'ús quotidià de l'internet i els interessos econòmics entorn de les dades personals “creen la tempesta perfecta per a l'exposició a la privacitat. És imperatiu que els pacients i els metges siguin més conscients d'aquests riscos per a la privacitat. Els sistemes sanitaris també haurien, com a qüestió d'exigència, esbossar polítiques que defensin la privacitat en l'ús de chatbots de LLM per ajudar els metges amb la documentació”³⁵.

Les consultes als sistemes d'IA generativa, s'estan abocant dades altament sensibles al sistema d'intel·ligència artificial. D'aquesta manera, s'hi atorga accés a l'empresa proveïdora o desplegada, donant lloc a situacions d'abús en les quals es pot fer un ús de les dades inadequat o no consentit per les persones afectades, que potser ni tan sols n'han estat informades³⁶.

Finalment, l'ús de la IA generativa té un impacte sobre l'autonomia i les capacitats de les persones que utilitzen el sistema. Un estudi publicat per Microsoft³⁷, mostra que la confiança en les aplicacions d'IA per portar a terme determinades tasques està relacionada amb un menor esforç de pensament crític. Aquest fet, a llarg termini, pot conduir a una dependència excessiva de l'eina i una disminució de la capacitat crítica de la persona usuària, juntament amb altres habilitats relacionades, particularment, la recopilació d'informació, la resolució de problemes i l'execució de tasques.

En conseqüència, l'ús de la IA generativa en l'àmbit de la salut pot arribar a malmetre, no només les habilitats del personal sanitari per portar a terme les seves funcions de la manera més adequada possible, sinó també la capacitat de la ciutadania per a prendre decisions relacionades amb la seva salut i el seu benestar.

Per tot el que s'ha explicat, hem considerat pertinent introduir en la justificació d'algunes preguntes articles del RIA que contenen obligacions referents als sistemes d'IA generativa, tot i que aquests siguin considerats, en principi, sistemes de risc limitat, així com unes poques altres que tracten concretament aquesta qüestió. D'aquesta manera, el Model PIO Salut també dona cobertura a situacions d'ús de la IA generativa com les descrites.

³⁵Blease, C. [Charlotte] (2024). Open AI meets open notes: surveillance capitalism, patient privacy and online record access. *Journal of Medical Ethics*, 50(2), 84-89. Disponible a: <https://jme.bmj.com/content/50/2/84> (Consultada el 3 de març de 2025).

³⁶ Al-kfairy, M. [Mousa], Mustafa, D. [Dheya], Kshetri, N. [Nir], Insiew, M. [Mazen], i Alfandi, O. [Omar] (2024). Ethical Challenges and Solutions of Generative AI: An Interdisciplinary Perspective. *Informatics*, 11(3), 58. Disponible a: <https://doi.org/10.3390/informatics11030058> (Consultada el 5 de febrer de 2025).

³⁷ Lee, H. P. [Hao-Ping] H. [Hank], Sarkar, A. [Advait], Tankelevitch, L. [Lev], Drosos, I. [Ian], Rintel, S. [Sean], Banks, R. [Richard], i Wilson, N. [Nicholas] (2025). The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects From a Survey of Knowledge Workers. *CHI Conference on Human Factors in Computing Systems (CHI '25)*. Disponible a: https://hankhplee.com/papers/genai_critical_thinking.pdf (Consultada del 26 de febrer de 2025).

3. OBJECTIUS

3.1. La utilitat d'afrontar dilemes ètics

Moltes de les particularitats que hem presentat en l'àmbit de la salut sobre els diferents principis estan fortament arrelades en l'existència de dilemes ètics. Aquests són presents en qualsevol de les fases del procés de creació, desenvolupament i ús dels sistemes d'IA.

Una part ineludible en el procés cap a una IA de confiança és de fet la identificació d'aquests dilemes. Un cop identificats, existeix una necessitat clara de no evadir la seva existència, sinó tot el contrari, afrontar la seva discussió i prendre les decisions que es derivin de forma conscient, involucrant en aquest procés a tots els actors rellevants al voltant del sistema d'IA³⁸. Trobar un balanç i prendre decisions adequades incrementarà la confiança en el sistema d'IA. El qüestionari que proposa el Model PIO Salut cerca centrar l'atenció en alguns d'aquests dilemes d'especial rellevància a l'àmbit sanitari.

En particular, parlem de dilemes que involucren la responsabilitat i el retiment de comptes, per una banda, amb el principi de seguretat i no maleficència. En el context de salut, l'excés en l'atribució de responsabilitat cap a l'usuari clínic d'un sistema d'IA, personal que desafortunadament no es troba sempre involucrat en el procés de creació del sistema, pot suposar un desincentiu en el seu ús efectiu final³⁹. Alhora, això dona com a resultat la maleficència del procés clínic, a més d'un possible conflicte amb el principi de justícia i equitat. L'excés d'atribució de responsabilitat en l'usuari clínic també suposaria la manca d'atribució de responsabilitats als creadors i desenvolupadors del sistema, disminueix l'incentiu perquè aquests enforteixin la seguretat en l'ús del sistema d'IA. Aquest fet és especialment rellevant si durant el desplegament del sistema les persones creadores es troben allunyades del seu control, donant lloc al mencionat *problema de control*⁴⁰.

Al contrari, l'excés d'èmfasi en l'atribució de responsabilitat als creadors i desenvolupadors dels sistemes d'IA, i la manca d'aquesta atribució a l'usuari clínic, pot al seu torn esdevenir en el temut *biaix d'automatització*, on la persona usuària clínic assumeix la recomanació del sistema d'IA sense l'atenció i anàlisi necessàries. Això resulta un dilema amb el principi d'autonomia⁴¹.

En el cas de la transparència i explicabilitat, darrerament s'està normalitzant l'ús de sistemes d'IA complexos, en els que la interpretabilitat i explicabilitat són compromeses. La motivació que a vegades justifica aquesta elecció és el possible rendiment més elevat

³⁸ Vetter, D. [Dennis], Amann, J. [Julia], Bruneault, F. [Frédéric] et al. (2023). Lessons Learned from Assessing Trustworthy AI. *Practice. DISO 2*, 35. Disponible a: <https://doi.org/10.1007/s44206-023-00063-1> (consultada el 2 de maig de 2024).

³⁹ WHO (2021). Ethics and governance of artificial intelligence for Health. *Op. cit.* 11.

⁴⁰ *Idem.*

⁴¹ *Idem.*

pel que fa a les mètriques de precisió en comparació amb sistemes més simples i explicables. Per tant, el dilema amb el principi de Seguretat i no maleficència és inevitable, i del seu estudi i decisions que se'n derivin depèn que el sistema d'IA es comporti de manera més o menys fiable. El Model PIO en l'àmbit de la salut fa èmfasi en aquest fet.

3.2. Els objectius del Model PIO Salut

Seguint les directrius del Model PIO, el Model PIO Salut té com a objectiu principal sensibilitzar sobre l'ús ètic dels sistemes d'IA a tots els agents involucrats dins del cicle de vida dels sistemes aplicats en salut - desenvolupadors, gestors, personal clínic i pacients - en cadascuna de les fases del seu cicle de vida: des del disseny, la modelització i la validació, al desplegament i el monitoratge un cop el sistema d'IA és operatiu. El Model PIO també té com a objectiu identificar accions adequades i inadequades que es donen durant tot aquest cicle a través del sistema d'autoavaluació que proposa i que permet valoritzar, recomanar i avançar en l'ús ètic de dades i sistemes d'IA⁴².

Part dels objectius del model d'autoavaluació PIO Salut són semblants als que persegueixen les auditories algorítmiques mèdiques⁴³. Mitjançant una investigació proactiva i dedicada d'errors algorítmics i de casos de mal funcionament existents, s'aconsegueix un monitoratge més efectiu d'un sistema d'IA. Aquest monitoratge permet entendre millor les limitacions del sistema i les conseqüències dels seus errors. De la mateixa manera, el Model PIO també serveix per adquirir informació més detallada del funcionament dels mateixos algoritmes i les possibles limitacions en diferents aspectes, amb l'objectiu d'incrementar la seguretat i rendiment dels algoritmes en la pràctica clínica⁴⁴.

Fomentar un acurat monitoratge, per tant, permet identificar errors i vulnerabilitats de forma ràpida, i permet també l'actuació immediata per minimitzar els efectes negatius que es generen. Això ha d'incloure mecanismes per reportar i documentar la informació rellevant, i també la planificació d'accions i protocols a dur a terme segons els resultats del sistema d'IA en aquest monitoratge⁴⁵.

El Model PIO Salut té com a objectiu qüestionar sobre aquest monitoratge i identificar els errors i vulnerabilitats dels sistemes d'IA en els aspectes clau relacionats amb els set

⁴² OEIAC (2022). El Model PIO (Principis, Indicadors i Observables). *Op. cit.* 5.

⁴³ ⁴³ Liu, X. [Xiaoxuan], Glocker, B. [Ben], McCradden, M. [Melissa] M., Ghassemi, M. [Marzyeh], Denniston, A. [Alastair] K., i Oakden-Rayner, L. [Lauren] (2022). The medical algorithmic audit. *Op. cit.* 10.

⁴⁴ *Idem*.

⁴⁵ Coalition for Health AI (2022). Blueprint for Trustworthy AI Implementation Guidance and Assurance for Healthcare. Disponible a: <https://www.coalitionforhealthai.org/papers/Blueprint%20for%20Trustworthy%20AI.pdf> (consultada el 2 de maig de 2024).

principis ètics per tal de promoure una IA de confiança al món sanitari⁴⁶. En el seu procés, el Model PIO proposa:

- Fer preguntes i proposar eines enfocades a autoavaluar l'ús i l'impacte previstos, així com fer un mapeig del sistema d'IA, identificant els diferents components, requisits, persones i processos associats.
- Promoure l'avaluació sobre els riscos de l'ús de dades i el sistema d'IA, en particular al voltant del flux de dades a cadascuna de les etapes de desenvolupament, desplegament i posterior monitorització.
- Qüestionar el procés de documentació i l'accés a informació rellevant necessària, en concret entorn dels conjunts de dades, els models d'IA i els resultats d'avaluacions existents.
- Fer preguntes sobre el procés de validació i testatge del sistema, així com la supervisió dels resultats en el context de desplegament i impacte previstos.
- Promoure l'autoavaluació sobre els punts d'interacció persona - IA, i la seva relació amb la presa de decisions i l'atribució de responsabilitats.

Més enllà de la necessitat d'avaluar l'impacte del sistema i els riscos associats tal com ja hem descrit anteriorment, en aquest procés de mapeig del sistema d'IA la selecció i processament de les dades utilitzades és essencial. Això inclou aspectes sobre el conjunt de dades emprades per l'entrenament dels algoritmes, la validació i els tests inicials, així com dades de funcionament durant el desplegament. L'accés a aquest darrer conjunt de dades de vegades no és possible i tampoc no són fàcilment supervisables. Per això es poden complementar amb alternatives de supervisió indirecta. L'èmfasi del Model PIO Salut no només es centra en l'estudi de les dades en sí, sinó també en tots els processos que tenen lloc al voltant de les dades, incloent-hi la seva documentació⁴⁷. Per això s'ha de tenir en compte el context en el seu conjunt, incloent-hi el coneixement clínic dels professionals mèdics i també el punt de vista de les persones pacients⁴⁸.

En qualsevol sistema d'IA, la metodologia de validació i el rigor científic en els processos de selecció de model, entrenament, validació i test del sistema són claus per una IA robusta i de confiança. La metodologia de validació està directament relacionada amb la possible atribució de responsabilitat a les persones creadores i desenvolupadores de la solució, així com amb la necessitat de transparència sobre tot el procés⁴⁹. El Model PIO Salut també proposa preguntes amb l'objectiu de qüestionar-se la qualitat d'aquest procés.

⁴⁶ WHO (2021), Ethics and governance of artificial intelligence for health. Disponible a: <https://www.who.int/publications/i/item/9789240029200> (consultada el 2 de maig de 2024).

⁴⁷ ⁴⁷ Liu, X. [Xiaoxuan], Glocker, B. [Ben], McCradden, M. [Melissa] M., Ghassemi, M. [Marzyeh], Denniston, A. [Alastair] K., i Oakden-Rayner, L. [Lauren] (2022). The medical algorithmic audit. *Op. cit.* 10.

⁴⁸ *Idem*.

⁴⁹ WHO (2021), Ethics and governance of artificial intelligence for Health. *Op. cit.* 42.

Seguint els principis, bona part de les recomanacions dirigides a usuaris clínics persegueixen millorar aspectes del procés de supervisió humana d'un sistema d'IA, ja sigui per mitigar errors, per habilitar una possible modificació del protocol de processament de dades, o fins i tot en casos extrems per limitar l'ús de dades i del sistema d'IA per grups específics, quan el rendiment dels sistemes sigui inacceptable⁵⁰.

Més endavant, s'exposaran els objectius concrets del Model PIO Salut enfocats en el principi transparència i explicabilitat, justícia i equitat, seguretat i no maleficència, responsabilitat i el retiment de comptes, privacitat, autonomia i sostenibilitat.

4. LA PERSONA PACIENT EN EL CENTRE

4.1. Els riscos derivats de la implementació de la IA

Estretament lligat als dilemes ètics que comporta l'ús IA, existeix la necessitat d'adreçar la implementació de la intel·ligència artificial en l'àmbit clínic situant a les persones pacients en el centre. Això és, cal posar la intel·ligència artificial al servei de les persones, i encarar els riscos que els sistemes d'IA comporten per als seus drets.

Certament, la salut és un aspecte que incideix profundament a multitud d'aspectes de la dignitat humana, i ha estat reconeguda com a tal en els instruments legals internacionals més rellevants, com ho són la Declaració universal dels drets humans (article 25.1) o el Pacte internacional de drets econòmics, socials i culturals (article 12). Així mateix, ha estat declarat per la Constitució espanyola, en el seu article 43, de la mateixa manera que l'Estatut de Catalunya recull, en el seu article 23, el dret de totes les persones a accedir als serveis sanitaris. És per aquest motiu que el RIA contempla que els sistemes d'IA utilitzats en aquest àmbit han de ser considerats d'alt risc.

És molt estesa l'opinió que la IA pot ajudar a reduir temps d'espera, retallar costos i omplir algunes limitacions de cobertura existents, per exemple, en l'atenció de persones que resideixen en zones llunyanes a institucions sanitàries⁵¹. Certament, la intel·ligència artificial pot ser altament beneficiosa en l'execució de determinades tasques (per exemple, l'anàlisi dels factors de risc, monitoratge de la salut d'una població, l'anàlisi d'historials clínics, així com en la predicció de malalties, infeccions, efectes adversos o supervivència), deguda a la seva capacitat per a detectar patrons entre grans grups de

⁵⁰ Liu, X. [Xiaoxuan], Glocker, B. [Ben], McCradden, M. [Melissa] M., Ghassemi, M. [Marzyeh], Denniston, A. [Alastair] K., i Oakden-Rayner, L. [Lauren] (2022). The medical algorithmic audit. *Op. cit.* 10.

⁵¹ McLennan, S. [Stuart], Fiske, A. [Amelia], Tigard, D. [Daniel], Müller, R. [Ruth], Haddadin, S. [Sami], i Buyx, A. [Alena] (2022). Embedded ethics: a proposal for integrating ethics into the development of medical AI. *BMC Medical Ethics*, 23(1), 6. Disponible a: <https://bmcmethics.biomedcentral.com/articles/10.1186/s12910-022-00746-3> (Consultada el 20 de gener de 2025).

dades⁵². Però no per aquest motiu la seva resposta esdevé més objectiva o acurada que la que prendria una persona professional de la salut, ja que els resultats algorítmics són probabilístics i no sempre estableixen relacions causals clares, i depenen altament de la qualitat i diversitat de les dades, la qual cosa, sovint, condueix a resultats esbiaixats, inconclusius o poc acurats⁵³. Per tant, sense negar que la intel·ligència artificial té un potencial positiu, s'han de prendre en consideració els perills derivats del seu ús.

A conseqüència de les limitacions presents en els sistemes d'IA, sorgeixen diverses implicacions ètiques i socials, com ho poden ser els diagnòstics incorrectes a gran escala provocats per errors algorítmics o biaixos en les dades, conflictes sobre la privacitat de les dades de les persones pacients, així com la pèrdua de l'autonomia i la deterioració de la relació entre pacients i facultatius⁵⁴.

En aquest sentit, el Consell d'Europa adverteix que, mentre els sistemes d'IA poden incrementar l'eficiència, també poden comportar un empitjorament de la qualitat de l'atenció sanitària en comparació amb les interaccions humanes cara a cara, les quals permeten al personal clínic tractar als pacients de manera més pròxima, humana i individualitzada⁵⁵. Les persones professionals de la salut compten amb la seva expertesa per a detectar factors que fàcilment poden quedar marginats dels patrons detectats per una IA. És per aquest motiu que l'organització internacional destaca el risc que determinats usos de la IA, com ho són les tecnologies de teleassistència, suposen respecte als diagnòstics incorrectes o poc precisos, podent-se generar una lesió sobre el dret humà a la salut i a la integritat física i mental de les persones afectades. A la llarga, aquest factor també pot conduir a un deteriorament de la confiança de la població respecte als serveis de la salut.

De la mateixa manera, existeix també el risc que l'ús clínic de la intel·ligència artificial condueixi a una sobre confiança dels resultats generats per la IA o la dependència d'aquesta per part de les persones professionals de la salut i les pacients. Això succeeix, per exemple, quan es considera que la intel·ligència artificial és capaç de prendre decisions inequívokes sobre el diagnòstic o el tractament a seguir. Aquest fet té com a conseqüència que la usuària del sistema no verifiqui si la decisió és suficientment correcta o adequada al cas concret. D'aquesta manera, es condueix a diagnòstics erronis o tractaments inadequats⁵⁶.

⁵² Morley, J. [Jessica], Machado, C. [Caio] C. V., Burr, C. [Christopher], Cows, J. [Josh], Joshi, I. [Indra], Taddeo, M. [Mariarosaria], i Floridi, L. [Luciano] (2020). The ethics of AI in health care: a mapping review. *Social Science & Medicine*, 260, 113172. Disponible a: <https://pubmed.ncbi.nlm.nih.gov/32702587/> (consultada el 14 de gener de 2025).

⁵³ *Idem*.

⁵⁴ *Idem*.

⁵⁵ Mittelstadt, B. [Brent] (2021). The impact of artificial intelligence on the doctor-patient relationship. *Consell d'Europa*. Disponible a: <https://www.coe.int/en/web/bioethics/report-impact-of-ai-on-the-doctor-patient-relationship> (consultada el 28 de juliol de 2024).

⁵⁶ Quinn, T. [Thomas] P., Sendadeera, M. [Manisha], et al (2021). Trust and medical AI: the challenges we face and the expertise needed to overcome them *Journal of the American Medical Informatics Association*, 28(4): 890–894. Disponible a: <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC7973477/#> (consultada el 28 de juliol de 2024).

Per tal de protegir l'autonomia de pacients i professionals, cal garantir que les decisions mèdiques siguin preses, en última instància, per una persona humana, i que estigui en control del sistema d'IA en qüestió en tot moment. En conseqüència, també esdevé necessari que els sistemes utilitzats permetin modificar, revertir i anul·lar els seus resultats, de manera que en cap cas es forci una decisió concreta⁵⁷.

Així doncs, resulta important d'advertir sobre els potencials impactes negatius de l'ús de la IA amb finalitats clíniques per a les persones i la necessitat de prendre mesures per adreçar-los. Per això resulta cabdal incorporar les consideracions ètiques en totes les fases del desenvolupament del sistema d'IA, idealment amb l'acompanyament de persones capacitades en l'àmbit de l'ètica⁵⁸. La regulació que afecta aquest processos, si bé indispensable, no és sinó l'últim pas en el reconeixement d'aquests principis ètics, atorgant-li caràcter obligatori.

Per aquest motiu, és necessari comptar amb mesures que facilitin aplicar i verificar el compliment dels estàndards i les normes esmentades. Un bon exemple de recursos creats per la Fundació TIC Salut Social, en el marc del Programa d'IA en Salut, que inclouen guies sobre bones pràctiques per al desenvolupament d'eines d'IA generativa en Salut⁵⁹; l'aplicació del RIA en l'àmbit sanitari⁶⁰; la qualificació i classificació d'aplicacions informàtiques basades en IA com a productes sanitaris d'acord la normativa europea⁶¹; i l'obtenció del marcatge CE d'aplicacions informàtiques basades en IA⁶². Aquestes guies, a més, van acompanyades d'infografies que sintetitzen el seu contingut⁶³.

El Model PIO Salut persegueix aquesta mateixa finalitat, posant a disposició de la ciutadania, especialment de les persones professionals de la salut o les entitats que desitgen crear un sistema d'IA amb aplicacions clíniques, una eina d'autoavaluació que

⁵⁷ WHO (2021). Ethics and governance of artificial intelligence for Health. *Op. cit.* 11.

⁵⁸ McLennan, S. [Stuart], Fiske, A. [Amelia], Tigard, D. [Daniel], Müller, R. [Ruth], Haddadin, S. [Sami], i Buyx, A. [Alena] (2022) Embedded ethics: a proposal for integrating ethics into the development of medical. *Op. cit.* 51.

⁵⁹ Aussó, S. [Susanna], Berenguer, A. [Antoni], Aznar, J. [Júlia], Raventós, C. [Carolina], Gómez, V. [Vaneza], i Bretones, M. [Maria] (2025). Guia de Bones Pràctiques per al Desenvolupament d'Eines d'IA Generativa en Salut Grans Models de Llenguatge (LLM), Fundació TIC Salut Social. Disponible a: https://iasalut.cat/wp-content/uploads/2025/03/Bones_Practiques_IAGen_LLM_Guia-CAT.pdf (consultada el 25 de març de 2025).

⁶⁰ Aussó, S. [Susanna], Berenguer, A. [Antoni], Aznar, J. [Júlia], Raventós, C. [Carolina], Gómez, V. [Vaneza], i Bretones, M. [Maria] (2025). Guia per a l'Aplicació del Reglament Europeu d'IA en Salut. Fundació TIC Salut Social. Disponible a: https://iasalut.cat/wp-content/uploads/2025/03/AI-Act_Guia-CAT.pdf (consultada el 25 de març de 2025).

⁶¹ Aussó, S. [Susanna], Berenguer, A. [Antoni], Aznar, J. [Júlia], Raventós, C. [Carolina], Gómez, V. [Vaneza], i Bretones, M. [Maria] (2025). Guia per a la qualificació i classificació d'una aplicació informàtica basada en IA com a producte sanitari segons els Reglaments Europeus MDR / IVDR, Fundació TIC Salut Social. Disponible a: https://iasalut.cat/wp-content/uploads/2025/03/Qualificacio_Classificacio_MDSW_Guia-CAT.pdf (consultada el 25 de març de 2025).

⁶² Aussó, S. [Susanna], Berenguer, A. [Antoni], Aznar, J. [Júlia], Raventós, C. [Carolina], Gómez, V. [Vaneza], i Bretones, M. [Maria] (2025). Guia per a l'obtenció del marcatge CE d'una aplicació informàtica basada en IA com a producte sanitari segons els Reglaments Europeus MDR / IVDR, Fundació TIC Salut Social. Disponible a: https://iasalut.cat/wp-content/uploads/2025/03/CE_Marcatge_MDSW_Guia-CAT.pdf (consultada el 25 de març de 2025).

⁶³ Podeu trobar les infografies al següent enllaç: <https://iasalut.cat/suport-a-iniciatives/>.

reculli les diferents guies i normatives rellevants en la matèria, facilitant la comprovació del grau d'adequació dels sistemes revisats.

4.2. Els drets de les persones usuàries dels serveis de salut

Si bé resulta evident que la protecció i promoció del dret a la salut i a la integritat física i mental ha de ser un element central en qualsevol escenari relatiu al desenvolupament, desplegament i ús clínic de la IA, l'aplicació d'aquesta tecnologia en l'àmbit de la salut afecta una gran varietat de drets humans, sigui de manera individual o per la seva relació amb l'esmentat.

Dins d'aquest context, resulta especialment rellevant el Conveni del Consell d'Europa sobre drets humans i biomedicina, més conegut com a Conveni d'Oviedo⁶⁴. Aquest es tracta del primer tractat internacional que té com a objecte la protecció i tracte digne de totes les persones humanes i els seus drets en l'àmbit biomèdic en condicions igualitàries. Així ho declara el seu article 1, al dictaminar que *“Les Parts en el present Conveni protegiran l'ésser humà en la seva dignitat i la seva identitat i garantiran a tota persona, sense cap discriminació, el respecte a la seva integritat i als seus altres drets i llibertats fonamentals respecte a les aplicacions de la biologia i la medicina”*⁶⁵. El benestar de les persones, a més, sempre haurà de primar sobre els interessos científics, segons es contempla en l'article 2 del conveni, amb la qual cosa, la possible consecució d'un avenç tecnològic o científic no justifica en cap cas posar en perill la seguretat o els interessos de cap individu.

L'article 3 del Conveni d'Oviedo proclama el dret a l'accés equitatiu a una atenció sanitària de qualitat. Això és, només cal garantir que totes les persones podran gaudir dels serveis sanitaris, sinó que, a més cal que aquest sigui equitatiu, evitant qualsevol situació de discriminació, i procurant que l'atenció dispensada sigui de suficient qualitat.

En aquest sentit, resulta important ressaltar que la intel·ligència artificial aplicada en l'àmbit de la salut pot reforçar biaixos sistèmics i agreujar desigualtats a existents en el sistema de salut, resultant en situacions de discriminació de persones o grups vulnerables. Això es deu al fet que els sistemes d'intel·ligència artificial són entrenats amb dades històriques, que poden estar esbiaixades, contenir informació discriminatòria o no ser prou representatives.

Per exemple, la sobre representació d'un gènere, una ètnia o una franja d'edat pot conduir a què el diagnòstic que s'hagi realitzat amb l'assistència d'un sistema d'IA no sigui prou

⁶⁴ Oviedo Convention (1997). Convention on Human Rights and Biomedicine. Disponible a: <https://www.coe.int/en/web/bioethics/oviedo-convention> (consultada el 20 de febrer de 2024).

⁶⁵ *Idem*. Traducció pròpia.

fiable, ja que no s'han tingut en compte les seves característiques particulars⁶⁶. En conseqüència, el tractament al qual es sotmeti a aquesta persona no serà prou efectiu, perjudicant la seva salut. Aquest fet, més enllà de comportar una desigualtat, malmet el dret a la salut i a la integritat física de la persona afectada⁶⁷.

Així mateix, la intel·ligència artificial pot dificultar l'accés als serveis de salut a aquelles persones que no posseeixen suficients coneixements o competències relacionades amb l'ús de les noves tecnologies, incrementant així l'escletxa digital⁶⁸. Per aquest motiu és indispensable prendre mesures que permetin adaptar i assistir a les persones que utilitzi les aplicacions d'IA, de manera que es garanteixi l'accessibilitat per a totes les persones.

Un altre dels drets que poden veure's amenaçats pel desplegament d'aquestes tecnologies és el dret a la privacitat, protegit a l'article 10 de el Conveni d'Oviedo, i a la protecció de dades. És cert que la recol·lecció i el processament de dades permeten millorar la qualitat de les dades utilitzades pel sistema i, per tant, el seu funcionament i fiabilitat. Això no obstant, aquestes dades molts cops contenen informació delicada de les persones afectades que ha de ser protegida. En conseqüència, quan els sistemes d'intel·ligència artificial recullen i utilitzen dades que han de ser protegides, esdevé necessari adoptar mesures per a la protecció d'aquestes. Aquestes mesures han d'incloure garanties que aquestes dades no seran venudes ni cedides sense el consentiment de la persona afectada, així com mesures de seguretat que prevegin que les dades siguin robades o compromeses.

Els riscos que comporten els sistemes d'IA per a la privacitat i protecció de dades no són exclusius del seu ús clínic. No obstant, les dades relacionades amb la salut són especialment sensibles, i la seva desprotecció pot comportar conseqüències especialment nocives per a les afectades. Per exemple, un reportatge de *Politico* ha denunciat que molts estats d'EUA que han criminalitzat l'avortament utilitzen dades de geolocalització recollides per *Google* per a la persecució de delictes, per la qual cosa es sospita que aquesta pràctica es podria començar a aplicar a la identificació, vigilància i persecució de dones que han decidit avortar⁶⁹. Si bé l'exemple proposat no es tracta d'un cas d'ús clínic de la IA, sí que reflecteix la importància de la protecció de dades i la privacitat per a garantir el correcte exercici del dret a la salut.

En relació amb la protecció de dades personals, cal remarcar que, segons l'article 4.1 del Reglament general de protecció de dades de la UE, aquesta expressió es refereix a tota

⁶⁶ Business for Social Responsibility (2013). AI and Human Rights in Healthcare. Disponible a: <https://www.bsr.org/reports/BSR-AI-Human-Rights-Healthcare.pdf> (consultada el 24 de juliol de 2024).

⁶⁷ Chung, J. [Jamison] (2021). How Will Health Care Regulators Address Artificial Intelligence?, The Regulatory Review. Disponible a: <https://www.theregreview.org/2021/10/18/chung-how-will-health-care-regulators-address-artificial-intelligence/> (consultada el 24 de juliol de 2024).

⁶⁸ Saeed S. [Sy] A. [Atezaz] i Masters, R. [Ross] M. [MacRae] (2021). Disparities in Health Care and Digital Divide, Current Psychiatry Reports 23(9): 61. Disponible a : <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC8300069/> (consultada el 24 de juliol de 2024).

⁶⁹ Ng, A. [Alfred] (2022). A uniquely dangerous tool': How Google's data can help states track abortions. Politico. Disponible a: <https://www.politico.com/news/2022/07/18/google-data-states-track-abortions-00045906> (consultada el 24 de juliol de 2024).

informació sobre una persona identificativa o identificable. En el context dels serveis sanitaris, els historials clínics inclouen dades com el nom, la data de naixement o el lloc de residència de les persones usuàries dels serveis de salut, però també contenen molta altra informació que reflecteixen el seu estat de salut. Aquestes dades, que poden incloure, per exemple, si un individu ha passat per una intervenció quirúrgica, quin és el seu grup sanguini, o si compta amb una pròtesi o un òrgan trasplantat, poden permetre identificar a una persona de manera indirecta ⁷⁰, incomplint d'aquesta manera la normativa de protecció de dades, podent causar una lesió sobre el dret a la privacitat de les persones afectades.

Això comporta que les mesures de protecció de dades en totes les fases del cicle de vida de l'aplicació d'IA, incloent-hi el seu entrenament, no poden limitar-se, per exemple, a l'anonimització de les mateixes, sinó que resulta necessari prendre accions addicionals, de manera que no resulti possible inferir la identitat d'una persona a partir del conjunt de dades.

Finalment, i estretament vinculat a l'autonomia, es troba el dret a no ser sotmès a cap intervenció sense prestar consentiment lliure i informat, recollit en els articles 5 a 9 del conveni. Aquest consisteix en que, abans de qualsevol intervenció la persona afectada ha d'haver donat el seu consentiment de manera lliure, i aquest consentiment únicament serà vàlid quan es disposa de tota la informació sobre el seu estat de salut i les malalties que pateixen, així com dels riscos associats a les tècniques de diagnòstic i dels tractaments suggerits per part dels professionals de la salut, el seu cost i les possibles alternatives.

Per tal que el pacient estigui suficientment informat, el personal mèdic ha de tenir un coneixement suficient de totes les eines involucrades en el procés clínic. Per tant, quan un sistema d'IA participa en el procés clínic, el compliment del consentiment informat no només depèn del coneixement mèdic que tinguin els professionals sanitaris sinó que també depèn de la seva formació sobre l'ús de dades i els sistemes d'IA. Les persones professionals de la salut han d'estar informades sobre en quines situacions s'apliquen aquests sistemes, quins són els seus límits, com funcionen i, molt important, com es validen i verifiquen els resultats obtinguts pel sistema d'IA per prendre decisions informades.

No obstant això, estudis actuals demostren la mancança de la integració de la formació en IA en els plans d'estudi dels estudis en medicina⁷¹. Així doncs, per tal de garantir una atenció sanitària de qualitat i respectuosa amb els drets de les persones pacients, esdevé

⁷⁰ Dewitte, P. [Pierre] (2025). AI Meets the GDPR: Navigating the Impact of Data Protection on AI Systems. In N. A. Smuha (Ed.), *The Cambridge Handbook of the Law, Ethics and Policy of Artificial Intelligence*, 133–157. Disponible a: [AI Meets the GDPR \(Chapter 7\) - The Cambridge Handbook of the Law, Ethics and Policy of Artificial Intelligence](#) (Consultada el 20 de febrer de 2025).

⁷¹ Weidener, L. [Lukas], i Fischer, M. [Michael] (2024). Artificial intelligence in medicine: cross-sectional study among medical students on application, education, and ethical aspects. *JMIR Medical Education*, 10(1), e51247. Disponible a: <https://mededu.jmir.org/2024/1/e51247/PDF> (Consultada el 4 de febrer de 2025).

indispensable afrontar el repte que suposa capacitar all i les professionals de la salut en matèria d'IA i els seus usos ètics. Això no només comporta la incorporació d'aquesta matèria en els estudis relacionats amb l'àmbit sanitari, sinó també en tasques de formació contínua per a aquelles persones que ja exerceixen una professió en l'àmbit clínic.

5. A QUI VA DIRIGIT?

5.1. Els subjectes del PIO Salut

El Model PIO Salut forma part del conjunt de sistemes d'autoavaluació que prenen com a referència el Model PIO sobre Obligacions i Drets i van dirigits a totes les organitzacions públiques o privades que tinguin un interès en el disseny, desenvolupament o desplegament de tecnologies basades en l'ús de dades i sistemes d'IA, i a totes les persones que estiguin utilitzant, adquirint o siguin destinatàries d'aquestes tecnologies. Donada la inevitable interacció amb el sistema sanitari de la població i la proliferació accelerada de sistemes d'IA, creiem que aquest treball és prou rellevant per interpel·lar bàsicament a tothom.

Un dels principals reptes en el desenvolupament i l'ús de dades i sistemes d'IA al món de la salut és la presència de biaixos en el seu desplegament. Amplificar la participació en el procés d'autoavaluació dels múltiples i diversos agents involucrats en l'ús d'un sistema d'IA té un impacte directe en detectar i possibilitar la reducció d'aquests biaixos. De fet, la construcció del Model PIO Salut segueix aquest mateix principi, involucrant en el treball autors tant de l'àmbit acadèmic com de l'àmbit privat, i de diferents perfils. D'aquesta manera el Model PIO Salut pretén interpel·lar de forma més directa a les persones usuàries dels sistemes d'IA en el món de la salut, allunyant-se de la sobreinvolucració recurrent d'actors provinents només de l'àmbit acadèmic, públic i de recerca.

El Model PIO Salut persegueix com una prioritat la inclusió de la visió de la persona pacient en l'autoavaluació dels sistemes d'IA: la seva veu ha d'estar visible i ser activa. Això es fonamenta en la importància que els sistemes siguin contestables per part de les persones pacients. Aquesta és una ambició central en tot el procés d'autoavaluació per una IA de confiança coherent amb multitud de treballs existents a l'àmbit de la salut⁷². La

⁷² Ploug, T. [Thomas]; Holm, S. [Søren] (2020), The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI. *Op. cit.* 17. En el mateix sentit: Coeckelbergh, M. [Mark] (2020). Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability. *Science and engineering ethics*, 26(4), 2051–2068. Disponible a: <https://doi.org/10.1007/s11948-019-00146-8>; ⁷² Liu, X. [Xiaoxuan], Glocker, B. [Ben], McCradden, M. [Melissa] M., Ghassemi, M. [Marzyeh], Denniston, A. [Alastair] K., i Oakden-Rayner, L. [Lauren] (2022). The medical algorithmic audit. *Op. Cit.* 10; MedTech Europe (2019). Trustworthy Artificial Intelligence (AI) in healthcare. *Op. Cit.* 25; WHO (2021), Ethics and governance of artificial intelligence for Health. *Op. Cit.* 42; i

participació tant de la persona pacient, com de la persona usuària del món sanitari, és imprescindible per evitar la falta d'atenció en la separació que hi ha sovint entre l'operador del sistema de dades i d'IA i el principal afectat: la persona pacient. Entendre les conseqüències d'aquesta separació és clau en el camí cap a una IA de confiança en l'àmbit de la salut.

Fetes aquestes consideracions inicials, la participació en el procés d'autoavaluació que proposa el Model PIO Salut hauria d'incloure agents i persones usuàries rellevants. Sense ser exhaustius, això inclou:

- Professionals científics de dades i altre personal TIC, com per exemple desenvolupadors de programari i informàtics, a càrrec del disseny, desenvolupament, desplegament i manteniment dels sistemes d'IA;
- Persones usuàries i operadores dels sistemes, com per exemple personal clínic i proveïdors sanitaris, paramèdics, infermeria, i també estudiants dins de l'entorn sanitari;
- Asseguradores i pagadors del sistema sanitari;
- Personal de l'àmbit legal, comitès d'ètica, personal administratiu i hospitals;
- Personal investigador, estudiants i personal en pràctiques;
- I un cop més: les persones pacients.

5.2. Subjectes afectats pel Reglament europeu sobre IA

Resulta oportú recordar, però, que si bé el PIO Salut interpel·la pràcticament a tota la societat, com dèiem al principi d'aquesta secció, determinats subjectes poden tenir un interès especial en el seu ús. En concret, les persones sobre a les quals les diferents normes referenciades al llarg del qüestionari vesteixen la responsabilitat de complir amb els requisits legals, sota la possibilitat d'una quantiosa penalització, poden fer ús d'aquest model per a examinar si el desenvolupament, les característiques i l'ús del sistema d'IA, així com l'obtenció i el tractament de les dades utilitzades, compleixen amb els estàndards imposats.

En aquest sentit, el RIA estableix els subjectes afectats als requisits aplicables als sistemes d'alt risc són els proveïdors de les aplicacions d'IA i els representants autoritzats dels proveïdors que es trobin fora de territori comunitari, així com els importadors, els distribuïdors i els responsables del desplegament dels sistemes.

Per proveïdor s'entén qualsevol persona física o jurídica, autoritat pública, òrgan o organisme que desenvolupi un sistema d'IA, o bé l'introdueix al mercat o el posa en servei

Procter, R. [Rob], Tolmie, P. [Peter], i Rouncefield, M. [Mark] (2023). Holding AI to Account: Challenges for the Delivery of Trustworthy AI in Healthcare. *Op. Cit.* 12.

sota el seu propi nom o marca. D'aquesta manera, no només es considerarà proveïdora a la persona que crea el sistema, sinó també a aquelles que encarreguen el desenvolupament del sistema per després comercialitzar-ho. Aquest es tracten dels subjectes sobre els quals el Reglament sobre intel·ligència artificial imposa la majoria d'obligacions.

D'acord amb l'article 16 del reglament, a més de garantir que els sistemes d'alt risc compleixen amb els requisits legals, els proveïdors han de comptar amb un sistema de gestió de qualitat que reuneixi els estàndards previstos a l'article 17 del mateix text normatiu. A més d'això, són responsables de sotmetre el sistema a l'avaluació de conformitat que certifica que compleix amb les disposicions RIA. En últim lloc, el subjecte proveïdor també serà responsable d'assegurar que l'aplicació reuneix amb els requisits d'accessibilitat jurídicament exigibles.

El representant autoritzat és aquella persona que compta amb un mandat per actuar en representació d'un proveïdor situat en un territori que no forma part d'un estat membre de la UE, i és una figura indispensable perquè aquests proveïdors puguin introduir el seu producte en el mercat europeu. Al seu torn, l'importador del sistema és aquell subjecte que s'encarrega de la introducció dels sistemes desenvolupats fora de la UE dins del territori comunitari, mentre que el distribuïdor es tracta d'aquella persona que, sense entrar en la definició de proveïdora o importadora, forma part de la cadena de subministrament que permet que el producte estigui disponible dins el mercat europeu.

Les obligacions d'aquests, subjectes, previstes als articles 22, 23 i 24 del RIA, principalment consisteixen a verificar que els sistemes que comercialitzen compten han superat l'avaluació de conformitat i compten amb el certificat corresponent, garantir que les seves tasques no afecten el sistema de manera que comprometi el compliment dels requisits legals i disposar de la documentació pertinent. A més hauran d'informar amb les autoritats competents si sospiten que el sistema no ha superat l'avaluació de conformitat o no satisfà alguna obligació. Addicionalment, els representants autoritzats hauran de portar a terme qualsevol altre compromís indicat al mandat.

Finalment, per responsables del desplegament ens referim a tota persona física o jurídica, autoritat pública, òrgan o organisme que utilitza el sistema d'IA sota la seva pròpia autoritat. En cas dels sistemes emprats en l'àmbit sanitari, els subjectes responsables del desplegament serien les institucions i centres sanitaris, siguin públics o privats, que comptin amb eines que incorporin, o siguin en si mateixes, una aplicació d'IA, així com professionals de la salut que exerceixin la seva professió de manera autònoma i facin ús d'aquestes tecnologies.

En virtut de l'article 26 del RIA, els subjectes responsables del desplegament d'un sistema d'IA d'alt risc hauran d'adoptar mesures tècniques i organitzatives que s'utilitzen d'acord amb les instruccions d'ús que l'acompanyen, i designaran a les persones físiques responsables de les mesures de vigilància humana. També hauran de controlar que les

dades entrades en el sistema resultes pertinents i vigilar el seu funcionament. En cas de considerar que l'ús de l'eina pot suposar algun tipus de risc per a la salut, la seguretat o els drets fonamentals d'alguna persona, el responsable del desplegament ha de suspendre l'ús i informar a les autoritats competents, així com al proveïdor i als altres subjectes implicats en la cadena de subministrament del producte. En últim lloc, el subjecte responsable del desenvolupament també té l'obligació d'informar a totes les persones responsables de les treballadores, i a les persones treballadores afectades o exposades al seu ús d'aquesta circumstància abans que es comenci a utilitzar.

5.3. Les 10 preguntes inicials o bàsiques

El Model PIO Salut s'inicia amb unes anomenades preguntes zero o preguntes bàsiques que reflecteixen el contingut ètic essencial que conforma la llista de verificació del model i, en aquest sentit, permet fer una ràpida aproximació sobre les consideracions ètiques fonamentals a tenir en compte.

Es recomana a les persones que encarregades de portar a terme l'autoavaluació presentar atenció a aquestes preguntes inicials abans d'iniciar algun projecte relacionat amb el desenvolupament, l'adquisició o el desplegament d'una aplicació d'IA.

10 preguntes inicials
Sabeu quin és el problema que esteu intentant resoldre amb l'ús d'un sistema d'IA?
Sou conscient que un sistema d'IA utilitzat per al diagnòstic o la determinació del tractament d'una patologia és considerat un producte sanitari a efectes legals?
Us heu plantejat si el potencial benefici que comporta l'ús del sistema és superior als riscos associats?
El personal sanitari és capaç d'entendre com funciona el sistema d'IA?
Sou conscient que les dades de les persones afectades pel sistema són altament sensibles i heu de garantir la seva protecció?

Sou conscient del vostre grau de responsabilitat respecte dels possibles perjudicis d'utilitzar un sistema d'IA?

Sou conscient que els conjunts de dades utilitzats per entrenar els sistemes d'IA poden contenir biaixos, conduint a la discriminació dels col·lectius poc representants?

Sou conscient que per a utilitzar un producte sanitari que tingui incorporat, o sigui en si mateix, un sistema d'IA heu d'obtenir prèviament el consentiment informat de la persona pacient?

Sou conscient que un sistema d'IA poc precís pot provocar perjudicis en la salut de les persones amb les quals s'ha utilitzat?

Coneixeu l'impacte ambiental com la petjada de carboni associada a l'ús del sistema d'IA?

6. INSTRUCCIONS

Abans de mostrar les qüestions que integren el qüestionari, en aquest capítol s'exposaran les instruccions a seguir per tal de completar-ho de manera òptima.

A continuació, es mostren les diferents fases de les quals es compon l'autoavaluació:

1	Inici de l'avaluació A través d'una llista de verificació o " <i>checklist</i> " organitzada sobre la base de principis ètics, disposicions legals, recomanacions ètiques i exemples.
2	Recollida de dades Mitjançant la informació que proporcionen els usuaris del Model PIO Salut a través de la llista de verificació o " <i>checklist</i> " sobre la base ètica i legal.

3	Obtenció de mètriques Adequació a les disposicions legals i recomanacions ètiques del Model PIO Salut.
4	Visualització de resultats A través d'una insígnia que reflecteix el grau d'assoliment dels principis ètics i disposicions legals contingudes en el Model PIO Salut.

IL·LUSTRACIÓ 1 FUNCIONAMENT DEL MODEL PIO SALUT.

0) Abans d'iniciar el procés d'avaluacions pròpiament, hi ha les 10 preguntes inicials o preguntes bàsiques que reflecteixen el contingut ètic essencial de les qüestions que conformen el **Model PIO Salut**, i us permetran fer una primera aproximació, i així us situarà prenent consciència sobre els sistemes d'IA.

1) El procés s'inicia amb el formulari de **l'avaluació** que consisteix en un llistat de verificació o *checklist* organitzat seguint els set principis. Les preguntes que conformen el qüestionari són fruit de la suma de les obligacions extretes fonamentalment del **RIA**, i també d'altres textos normatius, així com de recomanacions ètiques i legals. A més, hi trobareu exemples a cada pregunta que us permetran concretar millor allò que s'està preguntant.

Cada pregunta s'ha de respondre en funció de si el vostre sistema compleix o no. Veureu que en algunes preguntes, també hi ha una tercera opció que és "no aplica" per aquells casos els quals en el vostre sistema d'IA no heu previst allò preguntant.

Les preguntes estan agrupades d'acord amb quin dels 7 principis que vertebreren el Model PIO. Així mateix, cada principi es divideix en taules segons la temàtica que aborden **(1)**, de manera que cada pregunta es troba inclosa en una casella **(2)**. El requadre de la qüestió a respondre ve acompanyat de dues altres caselles. La primera conté les referències legals, polítiques i acadèmiques en què es basa la pregunta **(3)**, mentre que la segona exposa un exemple amb finalitats il·lustratives **(4)**.

(1) 1.2. Interpretabilitat del sistema i informació al pacient

Nº	Preguntes
1.2.1.	Coneixeu si el vostre sistema d'IA utilitza el model més intel·ligible garantint el rendiment assolit en favor de la transparència, o quin ha estat el criteri d'interpretabilitat inclòs en la metodologia de selecció del model d'IA? Aquesta pregunta fa referència a l'article 13 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA], que versa sobre la transparència i comunicació d'informació als usuaris. De la mateixa manera, la pregunta es fonamenta en l'article 5 del Conveni d'Oviedo, que versa sobre el consentiment lliure i informat, i en l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret, sobre la transparència i la supervisió dels sistemes d'IA. Finalment, té relació amb l'apartat Transparency de l'ALTAI i a l'apartat 2.2.1. Transparency d'AIEIG. Exemple: Hem determinat que el model utilitzat pel sistema d'IA és, d'una banda, suficientment simple i intel·ligible perquè les persones que interactuen amb ell siguin capaces de comprendre el seu funcionament (sense comprometre's en comparació a models més complexos). I d'altra banda, el nombre de paràmetres lliures del model d'aprenentatge estadístic ha estat un criteri addicional al propi rendiment del sistema per a la selecció final del model en desplegament.

IL·LUSTRACIÓ 2. EXEMPLE DE PREGUNTA DEL MODEL PIO SALUT

2) La suma de respostes a totes les preguntes dona lloc a uns percentatges que atribueixen el grau d'adequació de cada principi i en general, l'avaluació dels sistemes i dades de la IA. Aquests percentatges es tradueixen en un conjunt de **mètriques** del sistema d'IA sotmès a l'avaluació pel **Model PIO Salut**.

Les mètriques obtingudes són el reflex de l'adequació o no del sistema d'IA que l'entitat ha sotmès a avaluació com que partim d'una confiança implícita en una declaració de responsabilitat sobre la informació donada.

3) En finalitzar l'avaluació s'atribueix una **insígnia** gràfica que és la suma de les mètriques obtingudes de les respostes donades la qual reflecteix l'adequació de manera percentual de l'entitat als requisits legals i estàndards ètics de cada principi del model. Aquests valors es recullen amb una escala de colors sent el color blanc representatiu del no aplica, el color verd l'assumpció total del principi, vermell la no assumpció i els tons verd clar, caqui, groc i taronja les opcions intermèdies.

A més de la insígnia, que es tracta de la resposta sintetitzada, es facilita la matriu de riscos que es tracta d'una eina per obtenir un resum visual del nivell de risc amb relació a la probabilitat de portar a terme una acció al respecte. Aquesta probabilitat referent al temps poden ser cinc opcions: molt alta, ja que, l'efectuaireu aviat; alta perquè l'efectuaireu en un termini breu en temps sense ser immediat; ni alta ni baixa, ja que, compteu en realitzar-la en un mitjà termini; baixa perquè no forma part de les vostre prioritats, tot i que sou conscients que caldria efectuar-la; i molt baixa, ja que, no teníeu pensat dur-la a terme. Per tant, la taula recull les respostes negatives proporcionades segons la normativa i la probabilitat de dur a terme la qüestió sol·licitada a la pregunta que resumeix els punts de millora.

Els punts de millora són totes les preguntes on heu contestat negativament i heu precisat el grau de severitat i probabilitat en la qual duríeu a terme aquella qüestió per tal, de contestar positivament. Novament, fem servir la mateixa gamma de colors abans explicada.

Finalment, tota aquesta informació detallada per a cadascun dels principis i el percentatge assolit d'aquest el rebreu a l'adreça de correu de registre mitjançant un informe i la insígnia descarregable.

7. QÜESTIONARI

Principi 1: Transparència i explicabilitat

Aquest principi cerca garantir en l'àmbit sanitari que el personal mèdic disposi de totes les eines per tal que les seves decisions clíniques estiguin tan informades com sigui possible. D'altra banda, també té com a objectiu garantir que el pacient disposi de la informació necessària del sistema d'IA, el seu ús previst, les seves limitacions i els riscos associats. Les preguntes d'aquest apartat provenen majoritàriament de les seccions relatives a la transparència i explicabilitat dels diferents instruments i recomanacions consultades.

1.1. Anàlisi preliminar de l'ús del sistema

Nº	Preguntes
1.1.1.	Sou capaços d'entendre com el sistema d'IA arriba a un resultat determinat i com aquest contribueix en la decisió clínica? Aquesta pregunta fa referència als articles 13, 14, 50 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA], relatius a la transparència i comunicació d'informació als usuaris, la vigilància humana, els requisits de transparència dels sistemes d'IA generativa i altres sistemes que interactuïn amb persones i les obligacions de transparència dels proveïdors de models de propòsit general. Així mateix, està relacionada amb l'article 32 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris, l'article 29 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro relatius a la seguretat i el rendiment clínic dels dispositius mèdics i també l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris sobre les garanties sanitàries dels productes. De la mateixa manera, la pregunta es fonamenta en l'article 5 de el Conveni d'Oviedo, que versa sobre el dret al consentiment lliure i informat, i en l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret, sobre la transparència i la supervisió dels sistemes d'IA. Finalment, té relació amb l'apartat Transparency de l'ALTAI i amb l'apartat Transparency and Explicability del AI Ethics Assessment de GSMA. Exemple: El sistema d'anàlisi automàtic de radiografies de tòrax del nostre hospital és una eina de suport al diagnòstic d'un pacient amb pneumotòrax i som capaços d'identificar i entendre quines parts de la imatge contribueixen al fet que l'algoritme recomani amb una certa probabilitat el diagnòstic d'aquesta patologia per cada pacient en concret (explicabilitat local).
1.1.2.	Sou coneixedors de les possibles limitacions del vostre sistema d'IA i les heu comunicat a totes les persones implicades, incloent-hi les persones usuàries i pacients?

Aquesta pregunta fa referència [als articles 10.2, 13.3, 50 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la governança de dades, la transparència, els requisits de transparència dels sistemes d'IA generativa i altres sistemes que interactuïn amb persones i les obligacions dels proveïdors de models de propòsit general.

També es relaciona amb [l'article 10.11 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.10 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre la informació de la qual han d'anar acompanyats els productes sanitaris. En el mateix sentit, es vincula amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#) relatius a les garanties sanitàries dels dispositius mèdics.

Així mateix, la pregunta guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 de el Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA.

Així mateix, la pregunta està relacionada amb [els articles 4 i 5 de la Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica, relatius al dret a la informació assistencial](#).

Finalment, té relació amb [l'apartat Transparency del toolkit d'ICO sobre riscos i protecció de dades en la IA](#).

1.1.3. **Heu donat informació sobre l'ús previst del sistema d'IA a les persones pacients i de com les seves dades són processades, quins són els resultats previstos i quina influència tenen en la decisió clínica final?**

Aquesta pregunta fa referència [als articles 10.2, 13.3, 50 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la governança de dades, la transparència, els requisits de transparència dels sistemes d'IA generativa i altres sistemes que interactuïn amb persones i les obligacions dels proveïdors de models de propòsit general.

La pregunta també està relacionada amb [l'article 13 del Reglament UE](#)

[2016/679 del Parlament i el Consell, de 27 d'abril de 2016](#), sobre la informació que ha de facilitar-se a les persones que cedeixen les seves dades.

En relació a la governança de dades destaquem també [l'article 21 del Reglament 2022/868 del Parlament Europeu i del Consell, de 30 de maig del 2022, relatiu a la governança europea de dades \[Data Governance Act\]](#), relatiu als requisits específics per a salvaguardar els drets i interessos dels interessats i titulars de dades en relació amb les seves dades. Així com, [els articles 10.4 i 110.1 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [els articles 10.4, 102, i 103 del Reglament 2017/746 del Parlament Europeu i del Consell, de 5 d'abril del 2017, sobre els productes sanitaris per a diagnòstic in vitro](#), sobre la informació relativa als productes sanitaris i el tractament de les dades dels pacients que estableixen els requisits relatius al sistema de gestió de qualitat

De la mateixa manera, la pregunta es fonamenta en [els articles 5 i 10 de el Conveni d'Oviedo](#), que versen sobre el consentiment lliure i informat i el dret a la vida privada i la informació, i en [els articles 8 i 11 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA i el dret a la privacitat i la protecció de dades personals.

Finalment, té relació amb [l'apartat Transparency del toolkit d'ICO sobre riscos i protecció de dades en la IA](#), així com amb la publicació [The ethics of AI in health care: A mapping review](#) de Morleya, Machadoa, et al.

Exemple: Abans de la signatura del consentiment informat, hem establert un protocol que explica de manera detallada el funcionament del sistema d'IA i de les opcions que té el pacient per fer efectiu el seu dret a informar-se sobre un sistema d'IA contestable per part de la persona pacient.

1.1.4. **Totes les persones implicades en l'ús de dades i el sistema d'IA, en particular les usuàries clíniques, coneixen per a què s'utilitza el sistema d'IA i les seves limitacions?**

Aquesta pregunta fa referència [als articles 12, 13, 50 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la transparència i comunicació d'informació als usuaris, la documentació tècnica, els requisits de transparència dels sistemes d'IA generativa i altres sistemes que interactuïn amb persones i les obligacions dels proveïdors dels models de propòsit general.

Així mateix, la pregunta es relaciona amb [l'article 10.4 del Reglament](#)

[2017/745 del Parlament Europeu i del Consell, de 5 d'abril de dispositius mèdics \(MDR 2017\) 2017 sobre els productes sanitaris](#) i [l'article 10.4, i 59 del Reglament 2017/746 del Parlament Europeu i del Consell, de 5 d'abril del 2017](#), sobre els productes sanitaris per a diagnòstic in vitro, que versen sobre la informació relativa a l'ús previst del producte. Similarment, es basa en [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), sobre les garanties sanitàries dels dispositius mèdics.

També guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [els articles 5 i 10 de el Conveni d'Oviedo](#), que versen sobre el consentiment lliure i informat i el dret a la vida privada i la informació, i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA.

Finalment, té relació amb les guies establertes per la [WHO a l'apartat *Ensure transparency, explainability and intelligibility* del document *Ethics and governance of artificial intelligence for health*](#), i fa referència a [l'apartat *Transparency* de l'ALTAI](#) i està relacionada amb [l'apartat *Transparency and Explicability* del *AI Ethics Assessment* de GSMA](#).

Exemple: Hem establert que tot l'equip que envolta el procés clínic en què està involucrat el sistema d'IA ha de ser coneixedor del funcionament d'aquest al llarg de diverses jornades de formació obligatòries sobre l'ús del sistema.

1.1.5. **Heu tingut en compte l'existència del sistema d'IA en el procés del consentiment informat i n'heu avisat de manera clara i amb llenguatge comprensible?**

Aquesta pregunta fa referència [als articles 12, 13 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la transparència i comunicació d'informació als usuaris, la documentació tècnica i a les obligacions dels proveïdors dels models de propòsit general.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 de el Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA.

També guarda relació amb [els articles 10.4 i 63 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de dispositius mèdics \(MDR 2017\) 2017 sobre els productes sanitaris](#). Així com [els articles 10.4, i 59 del Reglament 2017/746 del Parlament Europeu i del Consell, de 5 d'abril del 2017, sobre els productes sanitaris per a diagnòstic in vitro](#), que versen sobre la informació relativa a l'ús previst del producte i el consentiment informat. En el mateix sentit, es vincula amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), sobre les garanties sanitàries dels dispositius mèdics.

A més, la pregunta es basa en [l'article 4 del Reial decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics](#) i els articles [2.2 i 4 Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica](#) sobre la informació als pacients i el consentiment informat.

La pregunta es fonamenta, també, amb el dret de les persones a ser informades sobre els tractaments mèdics, reconegut a [l'article 23.3 de l'Estatut d'Autonomia de Catalunya](#).

Finalment, té relació amb les guies establertes per la [WHO a l'apartat *Ethical Challenges to Use of Artificial Intelligence for Health Care* del document *Ethics and governance of artificial intelligence for health*](#), i fa referència a [l'apartat *Transparency* de l'ALTAI](#) i està relacionada amb [l'apartat *Transparency and Explicability* del *AI Ethics Assessment* de GSMA](#). De la mateixa manera, la pregunta fa referència a la publicació i [Generative AI and medical ethics: the state of play, de Zohny, Mann, et al.](#)

Exemple: En el consentiment informat hem inclòs de forma clara que al procés clínic hi ha l'ús de dades i un sistema d'IA involucrat i s'ha tingut en compte que l'explicació sigui intel·ligible per a persones usuàries i pacients a través d'enquestes i estudis específics.

1.1.6. **Heu establert mecanismes de documentació per facilitar la traçabilitat enfocada en la detecció d'errors, el retiment de comptes i la capacitat de les persones pacients de qüestionar el funcionament del sistema d'IA?**

Aquesta pregunta es fonamenta en [l'article 12.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versa sobre l'enregistrament del funcionament dels sistemes d'IA.

La pregunta també es està relacionada amb [els articles 10.9 i 83 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), i [els articles 10.8 i 78 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), que versen sobre els sistemes de gestió de qualitat i seguiment post-comercialització dels productes sanitaris. També guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

Així mateix, la pregunta fa referència [als articles 5 i 21 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatius a les garanties sanitàries dels productes sanitaris i a la informació que ha de fer-se pública al Registre de responsables de la posada al mercat de productes sanitaris de l'Agència Espanyola de Medicaments i Productes Sanitaris.

La pregunta també es basa en [l'article 9 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la responsabilitat i el retiment de comptes pels impactes adversos dels sistemes d'IA.

Finalment, té relació amb l'article [Transparency of AI in Healthcare as a Multilayered System of Accountabilities: Between Legal Requirements and Technical Limitations](#), d'Anastasia Kiseleva *et al.* i en el [subapartat Traceability de l'ALTAI](#).

1.1.7. El procés de generació de resultats del vostre sistema d'IA està degudament documentat i és consultable per part dels usuaris clínics en un llenguatge comprensible i accessible?

Aquesta pregunta fa referència [als articles 11, 12 i 53.1, 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la documentació tècnica i els registres automàtics dels sistemes d'alt risc i les obligacions dels proveïdors de models de propòsit general.

La pregunta també està relacionada amb [l'article 10.4 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#). Així com, [l'article 10.4 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#) que versen sobre la documentació tècnica dels productes sanitaris i [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#) relatius a les garanties sanitàries dels dispositius mèdics. Així mateix, la pregunta guarda relació amb [l'article 7 de la Directiva sobre responsabilitat](#)

[pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 de el Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA.

Finalment, té relació amb [Transparency of AI in Healthcare as a Multilayered System of Accountabilities: Between Legal Requirements and Technical Limitations](#), d'Anastasia Kiseleva et al., i [l'apartat fairness del toolkit d'ICO](#) i a [l'apartat Reliability, Fall-back plans and Reproducibility de l'ALTAI](#).

Exemple: Hi ha mecanismes que enregistren aquesta informació de manera automàtica a cada ús, aquesta informació està digitalitzada, centralitzada i és accessible a les persones usuàries clíniques. Aquesta informació inclou les variables i aspectes rellevants que el sistema ha tingut en compte a l'hora de generar els resultats i enregistra una còpia d'entrada per poder accedir-hi i consultar-lo en un futur.

1.1.8. **Heu establert un procés de documentació i accés controlat al material rellevant (en concret el conjunt de dades), als models d'IA i als resultats de les avaluacions existents sobre el sistema d'IA?**

Aquesta pregunta fa referència [als articles 11, 12 i 17 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la documentació tècnica i els registres i al sistema de gestió de qualitat dels sistemes d'IA.

La pregunta també es fonamenta en [els articles 10.9 i 29 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), i [els articles 10.8 i 26 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), que estableixen els requisits relatius al sistema de gestió de qualitat i el registre dels dispositius mèdics.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 de el Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA.

Finalment, té relació amb l'article [*Transparency of AI in Healthcare as a Multilayered System of Accountabilities: Between Legal Requirements and Technical Limitations*](#), d'Anastasia Kiseleva *et al.*

Exemple: Hem elaborat documentació detallada sobre el funcionament, l'ús previst del sistema d'IA, l'històric de resultats i aquest és accessible tant als usuaris com als pacients en els supòsits que també hem detallat en el procés. A més, poden dirigir les seves preguntes i comentaris sobre el mateix sistema a les persones responsables.

1.1.9. **Heu documentat el tipus de dades, la seva procedència i els canvis efectuats a totes les fases del cicle de vida del vostre sistema d'IA?**

Aquesta pregunta fa referència [als articles 10.2, 13.3, 50 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la governança de dades, la transparència i comunicació d'informació als usuaris, els requisits de transparència dels sistemes d'IA generativa i altres sistemes que interactuïn amb persones i a les obligacions dels proveïdors de models de propòsit general.

Així mateix, la pregunta guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA.

Finalment, té relació amb els apartats de [Reliability](#), [Fall-back plans and Reproducibility](#) i [Privacy and Data Governance](#) de l'ALTAI.

Exemple: Hem elaborat un document on es recullen les dades utilitzades tant en l'entrenament com en la resta de fases del cicle de vida de l'ús de dades i el sistema d'IA, detallant les característiques de les mateixes i de com s'han obtingut i processat.

1.2. Interpretabilitat del sistema i informació al pacient

Nº Preguntes

1.2.1. Coneixeu si el vostre sistema d'IA utilitza el model més intel·ligible garantint el rendiment absolut en favor de la transparència, o quin ha estat el criteri d'interpretabilitat inclòs en la metodologia de selecció del model d'IA?

Aquesta pregunta fa referència a [l'article 13 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versa sobre la transparència i comunicació d'informació als usuaris.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA.

Finalment, té relació amb [l'apartat Transparency de l'ALTAI](#) i a [l'apartat 2.2.1. Transparency d'AIEIG](#).

Exemple: Hem determinat que el model utilitzat pel sistema d'IA és, d'una banda, suficientment simple i intel·ligible perquè les persones que interactuen amb ell siguin capaces de comprendre el seu funcionament (sense comprometre's en comparació a models més complexos). I d'altra banda, el nombre de paràmetres lliures del model d'aprenentatge estadístic ha estat un criteri addicional al propi rendiment del sistema per a la selecció final del model en desplegament.

1.2.2. Heu considerat quins coneixements tècnics i quins materials didàctics addicionals són necessaris per entendre i explicar el funcionament del sistema d'IA a les persones pacients?

Aquesta pregunta fa referència [als articles 12, 13 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la transparència i comunicació d'informació als usuaris, la documentació tècnica i a les obligacions dels proveïdors dels models de propòsit general.

També guarda relació amb [els articles 10.4, i 63 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i els articles [10.4, i 59 del Reglament 2017/746 del Parlament Europeu i del Consell, de 5 d'abril del 2017](#), sobre els productes sanitaris per a

diagnòstic in vitro, que versen sobre la informació relativa a l'ús previst del producte i el consentiment informat. I amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), sobre les garanties sanitàries dels dispositius mèdics.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [els articles 8 i 20 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA i l'alfabetització i les habilitats digitals.

A més, la pregunta es basa en [l'article 4 del Reial decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics](#). En el mateix sentit, està vinculat amb els articles [2.2](#) i [4 Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica](#) que versen sobre la informació als pacients i el consentiment informat.

Finalment, es fonamenta en el dret de les persones a ser informades sobre els tractaments mèdics, reconegut a [l'article 23.3 de l'Estatut d'Autonomia de Catalunya](#).

Exemple: Hem desenvolupat una infografia que explica a les persones pacients i de manera detallada com funciona el sistema d'IA, a més, aquesta informació pot ser ampliada amb material addicional proporcionat pel personal clínic expert en el sistema

1.2.3. **Considereu que el sistema d'IA i l'obtenció de resultats és comprensible per a les persones pacients tenint en compte llurs capacitats i competències?**

Aquesta pregunta fa referència [als articles 13.2, 16 i 95.2, del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la transparència i comunicació d'informació als usuaris, les obligacions dels proveïdors dels sistemes d'alt risc i la promoció de codis de conducta d'aplicació voluntària.

Les preguntes també estan relacionades amb [els articles 10.4 i 63 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [els articles 10.4, i 59 del Reglament 2017/746 del Parlament Europeu i del Consell, de 5 d'abril del 2017 sobre els productes sanitaris per a diagnòstic in vitro](#), que versen sobre la informació relativa a l'ús

previst del producte i el consentiment informat, així com [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris sobre les garanties sanitàries dels dispositius mèdics](#).

La pregunta també es fonamenta en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA. De la mateixa manera, la pregunta es basa en el dret de les persones a ser informades sobre els tractaments mèdics, reconegut a [l'article 23.3 de l'Estatut d'Autonomia de Catalunya](#).

Exemple: Hem procurat proporcionar a les persones pacients informació clara i detallada sobre el funcionament del sistema d'IA, el seu procés de generació de resultats i els resultats obtinguts, adaptant l'explicació quan resulta necessari per garantir que sigui comprensible per a qualsevol persona.

1.3. Explicabilitat del sistema

Nº	Preguntes
----	-----------

1.3.1.	En el cas que el sistema estigui basat en variables clíniques, heu implementat mecanismes per determinar de forma clara quines són les variables clíniques més rellevants que utilitza el vostre sistema d'IA quan fa una predicció?
--------	---

Aquesta pregunta fan referència [als articles 13, 14 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la transparència i comunicació d'informació als usuaris, la vigilància humana i les obligacions dels proveïdors dels models d'IA de propòsit general.

D'altra banda, les preguntes també estan relacionades amb [l'article 10.11 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), i [l'article 10.10 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), relatius a la informació que ha d'acompanyar els sistemes d'IA. Així com [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), que versa sobre les garanties sanitàries dels dispositius mèdics. Així mateix, la pregunta guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter

defectuós dels productes.

Finalment, també es basa amb [l'apartat Transparency de l'ALTAI](#) i està relacionada amb [l'apartat Transparency and Explicability del AI Ethics Assessment de GSMA](#).

Exemple: Hem aplicat tècniques basades en la importància estadística de les variables d'entrada, i acceptades per la comunitat científica en aprenentatge estadístic, per tal de determinar les variables clíniques més rellevants per al model d'aprenentatge automàtic que informa el sistema d'IA. Aquestes tècniques i els resultats derivats han estat contrastades pel personal clínic rellevant en totes les fases de validació.

1.3.2. **Sou coneixedors del tipus de model que utilitza el vostre sistema d'IA i quin és el seu grau d'explicabilitat?**

Aquesta pregunta fa referència [als articles 13, 14 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la transparència i comunicació d'informació als usuaris, la vigilància humana i les obligacions dels proveïdors dels models d'IA de propòsit general.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA.

La pregunta es relaciona també amb [l'apartat Transparency de l'ALTAI](#) i està relacionada amb [l'apartat Transparency and Explicability del AI Ethics Assessment de GSMA](#).

Exemple: El vostre sistema d'IA utilitza models d'aprenentatge automàtic basats en arbres de decisió que és comprensible pel personal clínic (encara que el personal no sigui expert en aprenentatge automàtic) segons enquestes realitzades durant la fase de validació.

1.3.3. **Considereu que heu establert un nivell de transparència i explicabilitat que no posa en risc la seguretat del vostre sistema d'IA?**

Aquesta pregunta fa referència [als articles 10 i 15 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la governança de dades i precisió, solidesa i ciberseguretat dels sistemes d'IA.

Així mateix, la pregunta guarda relació amb [l'article 7 de la Directiva sobre](#)

[responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

També es fonamenta en [l'article 12 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre la fiabilitat dels sistemes d'IA.

Finalment, la pregunta es relaciona amb [l'apartat Transparency de l'ALTAI](#) i està relacionada amb [l'apartat Transparency and Explicability del AI Ethics Assessment de GSMA](#).

Exemple: Hem establert un nivell de transparència prou gran per permetre una comprensió detallada del sistema i hem implementat les mesures adients per evitar possibles ciberatacs que posin en perill el rendiment del sistema, l'accessibilitat i les dades.

1.3.4.

S'ha validat a nivell clínic per part de les persones professionals que utilitzaran el sistema d'IA si els resultats que dona estan ben explicats i són comprensibles, són adequats i ajuden en la tasca a realitzar?

Aquesta pregunta es fonamenta en [els articles 13, 14.4 i 95.2, del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la transparència i comunicació d'informació als usuaris, la vigilància humana i la promoció de codis de conducta d'aplicació voluntària en matèria de participació multisectorial.

Així mateix, està relacionada amb [l'article 32 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), [l'article 29 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#) relatius a la seguretat i el rendiment clínic dels dispositius mèdics i també [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#) sobre les garanties sanitàries dels productes. Així mateix, la pregunta guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es basa en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA.

Exemple: Prèviament al desplegament del sistema, l'equip de professionals que l'utilitzarà ha pogut comprovar i validar que els resultat generats per

aquest són comprensibles i útils en realització a la rasca a realitzar.

1.3.5. Sou coneixedors de les hipòtesis fetes i les mètriques que s'han utilitzat per optimitzar i validar el sistema d'IA així com la metodologia de validació per desenvolupar el sistema?

Aquesta pregunta fan referència [als articles 9 i 17 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a les mesures de gestió de riscos i al sistema de gestió de qualitat dels sistemes d'IA.

La pregunta també es relaciona amb [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#) que versen sobre el sistema de gestió de qualitat dels productes sanitaris. També [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#) relatiu a les garanties sanitàries dels dispositius mèdics. Així mateix, la pregunta guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, es fonamenta en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre la transparència i la supervisió dels sistemes d'IA.

Finalment, també es basa amb [l'apartat Accountability de l'ALTAI](#).

Exemple: El sistema d'IA ha estat optimitzat per maximitzar l'àrea sota la corba ROC, establint a més una sensibilitat mínima del 90%. La metodologia de validació ha estat publicada i acceptada en revistes d'alt nivell en ciència de dades, i correspon a un mètode robust acceptat per la comunitat científica.

1.3.6. Heu desenvolupat eines que permeten identificar i traçar les diferents accions i persones involucrades en el sistema d'IA?

Aquesta pregunta fa referència [als articles 12 i 17 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius als registres del sistema d'IA i al sistema de gestió de qualitat.

La pregunta també es relaciona amb [els articles 25 a 34 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [els articles 22 a 30 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#) sobre la

identificació i traçabilitat de dispositius mèdics.

Així mateix, es fonamenta en [l'article 9 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre la responsabilitat i el retiment de comptes pels impactes adversos dels sistemes d'IA.

També es basa en [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics.

Finalment, la pregunta està relacionada amb la publicació [The ethics of AI in health care: A mapping review](#) de Morleya, Machadoa, et al.

Exemple: Hem elaborat mecanismes que enregistren cada interacció i resultat del sistema d'IA, així com les persones usuàries que han estat involucrades en el seu ús a través d'un sistema que inclou un procés d'identificació de l'usuari.

1.3.7. **Teniu constància de si els mecanismes de transparència i traçabilitat implementats segueixen el marc regulador corresponent de la Unió Europea?**

Aquesta pregunta fa referència [als articles 12, 17 i 43 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius als registres, al sistema de gestió de qualitat i l'avaluació de conformitat dels sistemes d'IA.

Així mateix, es relaciona amb [els articles 25 a 34 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) sobre la identificació i traçabilitat de dispositius mèdics.

També es fonamenta en [els articles 19, 20 i 52 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), sobre la declaració de conformitat i el marcatge CE dels productes sanitaris, i [els articles 17, 18 i 48 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#) que versen sobre les avaluacions de conformitat i el marcatge CE dels productes sanitaris. Així mateix, la pregunta guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

Similarment, es basa en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa

sobre la transparència i la supervisió dels sistemes d'IA.

Finalment, la pregunta guarda relació amb [l'article 21 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a la informació que ha de fer-se pública al Registre de Responsables de la posada al mercat de productes sanitaris de l'Agència Espanyola de Medicaments i Productes Sanitaris.

1.3.8. **Per tal de protegir les dades de persones menors d'edat heu tingut en compte salvaguardes com la traçabilitat de les dades per verificar-ne l'edat?**

Aquesta pregunta fa referència [als articles 10 i 17 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la governança de dades i el sistema de gestió de qualitat.

També està relacionada amb [l'article 8 del Reglament 2016/679 del Parlament i del Consell, de 27 d'abril de 2016, relatiu a la protecció de les persones físiques pel que fa al tractament de dades personals i a la lliure circulació d'aquestes dades](#) que versa sobre el consentiment de les persones menors d'edat.

Així mateix, es fonamenta en [l'article 18 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre els drets de la infantesa.

Finalment, també a les possibles afectacions de drets fonamentals assenyalades per [l'apartat Transparency de l'ALTAI](#).

Exemple: Abans de recollir les dades utilitzades hem verificat i constatat l'edat de la persona de la qual s'extreuen i s'han aplicat els protocols corresponents en cas que sigui menor d'edat.

1.4. Política de participació i protocols de comunicació

Nº Preguntes

1.4.1. **Els protocols ètics preveuen l'ús de sistemes d'IA per assistir a la pràctica clínica?**

Aquesta pregunta fa referència a [l'article 95.2 del Reglament 2024/1689 del](#)

[Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció de codis de conducta d'aplicació voluntària en matèria de participació multisectorial.

Així mateix, la pregunta es relaciona amb les recomanacions contingudes a [l'apartat Diversity, Non-discrimination and Fairness: Stakeholder Participation de l'ALTAI](#) i a [l'apartat Injustice in Algorithmic Systems del Algorithmic Equity Toolkit d'ACLU](#).

Exemple: El Comitè d'ètica està conformat per persones de diferents grups d'interès i receptors del nostre sistema i tenen la tasca de revisar decisions crítiques de la IA i el seu impacte.

1.4.2. **Heu establert procediments per comunicar de manera regular informació sobre en quines decisions clíniques les persones participants són assistides per un sistema d'IA?**

Aquesta pregunta està relacionada amb [l'article 50 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versa sobre els requisits de transparència dels sistemes d'IA generativa i altres sistemes que interactuïn amb persones.

Per l'altra banda, es fonamenta en [els articles 5 i 10 del Conveni d'Oviedo](#), que versen sobre el consentiment lliure i informat i el dret a la vida privada i la informació, i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA.

Així mateix, la pregunta fa referència a les recomanacions contingudes a [l'apartat Diversity, Non-discrimination and Fairness: Stakeholder Participation de l'ALTAI](#), [l'apartat 2.2.2 Accountability d'AIEIG](#), així com amb [l'apartat Comunicació i Informació de les recomanacions sobre ètica en la Intel·ligència Artificial d'UNESCO](#) i [l'apartat accountability del toolkit d'ICO](#).

Exemple: A les reunions setmanals, l'equip mèdic proper al sistema d'IA transmet les decisions clíniques preses informades pel sistema, així com les possibles incidències que hagin pogut ocórrer.

1.4.3. **Heu establert eines de comunicació que permeten a les persones usuàries del vostre sistema d'IA realitzar comentaris sobre el seu funcionament?**

Aquesta pregunta fa referència a [l'article 95.2. del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció

de codis de conducta d'aplicació voluntària en matèria de participació multisectorial.

Així mateix, la pregunta es relaciona amb [l'article 5 de la Llei 33/2011, del 4 d'octubre, general de Salut Pública](#) i [l'article 10.10 de la Llei 14/1986, de 25 d'aril, General de Sanitat](#) sobre el dret a la participació en matèria de salut.

Finalment, també es basa en les recomanacions contingudes a [l'apartat Diversity, Non-discrimination and Fairness: Stakeholder Participation de l'ALTAI](#), [l'apartat 2.2.2 Accountability d'AIEIG](#), [l'apartat Comunicació i Informació de les recomanacions sobre ètica en la Intel·ligència Artificial d'UNESCO](#), i [l'apartat accountability del toolkit d'ICO](#).

Exemple: Hem creat una plataforma en línia on les persones usuàries i les pacients poden realitzar comentaris respecte al sistema d'IA els quals són accessibles per a altres persones usuàries, les desenvolupadores i responsables del sistema.

1.4.4. **Heu consultat amb d'altres organitzacions que hagin dissenyat o implementat un sistema d'IA semblant per conèixer antecedents i experiències?**

Aquesta pregunta fa referència a [l'article 95.2. del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció de codis de conducta d'aplicació voluntària en matèria de participació multisectorial.

Finalment, també es basa en les recomanacions contingudes a [l'apartat Diversity, Non-discrimination and Fairness: Stakeholder Participation de l'ALTAI](#).

Principi 2: Justícia i equitat

Aquest principi té com a finalitat garantir que els sistemes d'IA no provoquin situacions que puguin resultar injustes o discriminatòries envers a les persones pacients i usuàries del sistema de salut que són afectades pels mateixos, especialment pel que respecta a persones, comunitats o grups vulnerables. Les preguntes d'aquest apartat provenen majoritàriament de les seccions relatives a la diversitat, no-discriminació i justícia dels diferents instruments i recomanacions consultades.

2.1. Identificació i mitigació de biaixos

Nº Preguntes

2.1.1. Coneixeu si el conjunt de dades amb què s'ha entrenat i validat el sistema d'IA és representatiu de la població a la qual s'aplica?

Aquesta pregunta es fonamenta en [l'article 10.3 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), sobre la representativitat de les dades utilitzades per sistemes d'IA d'alt risc.

Així mateix, la pregunta guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

D'altra banda, la pregunta es basa en [l'article 10 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre la igualtat i la no discriminació, i amb el dret a la igualtat i a la no discriminació, reconegut a [l'article 14 de la Constitució Espanyola](#), [l'article 15 de l'Estatut d'Autonomia de Catalunya](#) i [l'article 6 de la Llei 33/2011, d'octubre, general de salut pública](#).

Finalment, la pregunta està relacionada amb la publicació [The ethics of AI in health care: A mapping review](#) de Morleya, Machadoa, *et al* i el principi d'universalitat de la guia [FUTURE-AI](#).

Exemple: Hem determinat que el nostre sistema d'IA de triatge en atenció primària s'utilitza en un CAP amb un alt percentatge de gent d'edat avançada on els mecanismes de validació de les dades verifiquen que la població amb què es va entrenar l'algoritme n'inclou una part representativa.

2.1.2. Heu desenvolupat mecanismes per detectar i corregir possibles biaixos en el comportament del sistema d'IA?

Aquesta pregunta es fonamenta en [els articles 9, 10.3, 15.3, 53, 55.1, 72 i 95.2. del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a les mesures de gestió de riscos, la gestió de dades, la resiliència i robustesa dels sistemes d'IA, les obligacions dels proveïdors de sistemes de propòsit general, el seguiment post-comercialització dels sistemes d'alt risc, i la promoció de codis de conducta d'aplicació voluntària en matèria de grups vulnerables.

També està relacionada amb [els articles 10.2, 10.9 i 83 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els](#)

[productes sanitaris](#) i [els articles 10.2, 10.8 i 78 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#) i [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#) relatius als sistemes de gestió de riscos, de qualitat i de seguiment post-comercialització dels productes sanitaris. Juntament amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

D'altra banda, la pregunta es relaciona amb [els articles 10 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la igualtat i la no discriminació i sobre el marc de gestió de riscos i impactes dels sistemes d'IA. Juntament amb el dret a la igualtat i a la no discriminació, reconegut a [l'article 14 de la Constitució Espanyola](#), [l'article 15 de l'Estatut d'Autonomia de Catalunya](#) i [l'article 6 de la Llei 33/2011, d'octubre, general de salut pública](#).

Finalment, la pregunta està relacionada amb la publicació [The ethics of AI in health care: A mapping review de Morleya, Machadoa, et al](#) i el principi d'equitat de la guia [FUTURE-AI](#).

Exemple: Hem establert eines de monitoratge que periòdicament analitzen les dades recollides per detectar, a més de possibles variacions en el rendiment del sistema, l'aparició de biaixos respecte a les dades utilitzades per l'entrenament del sistema.

2.1.3. **Heu avaluat l'impacte dels possibles biaixos del vostre sistema d'IA per saber si podria reforçar discriminacions, estereotips o prejudicis envers persones pertanyents a grups vulnerables (ex. persones amb discapacitat, minories ètniques o individus d'edat avançada)?**

Aquesta pregunta fa referència [als articles 9, 10.3, 15.3, 27.1, 53, 55.1 i 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a les mesures de gestió de riscos, la governança de dades, la resiliència i robustesa dels sistemes d'IA, l'avaluació de l'impacte dels sistemes d'alt risc sobre drets fonamentals, i la promoció de codis de conducta d'aplicació voluntària en matèria de grups vulnerables.

També està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

D'altra banda, la pregunta està vinculada amb [els articles 10 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans](#),

[democràcia i estat de dret](#), que versen sobre la igualtat i la no discriminació i sobre el marc de gestió de riscos i impactes dels sistemes d'IA. Juntament amb el dret a la igualtat i a la no discriminació, reconegut a [l'article 14 de la Constitució Espanyola](#), [l'article 15 de l'Estatut d'Autonomia de Catalunya](#) i [l'article 6 de la Llei 33/2011, d'octubre, general de salut pública](#).

Així mateix, la pregunta està relacionada amb el principi d'equitat de la guia [FUTURE-AI](#).

2.1.4. **Heu realitzat una avaluació sobre l'impacte que pot tenir l'ús del sistema d'IA sobre els drets fonamentals de les persones pacients?**

Aquesta pregunta fa referència [als articles 9, 27.1, 55.1 i 95.2, del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a les mesures de gestió de riscos, l'avaluació de l'impacte dels sistemes d'alt risc sobre drets fonamentals, les obligacions dels proveïdors de models de propòsit general de risc sistèmic i la promoció de codis de conducta d'aplicació voluntària en matèria de grups vulnerables.

També està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

D'altra banda, la pregunta es relaciona amb [els articles 4 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la protecció dels drets humans i sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

2.1.5. **Heu adoptat mesures que permetin mitigar possibles biaixos cognitius, com per exemple els biaixos de confirmació o d'automatització, tant en la fase de disseny com en la implementació del sistema?**

Aquesta pregunta fa referència [als articles 9, 10.3, 15.3, 55.1 i 95.2, del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a les mesures de gestió de riscos, la governança de dades, la resiliència i robustesa dels sistemes d'IA, les obligacions dels proveïdors del models de propòsit general de risc sistèmic i la promoció de codis de conducta d'aplicació voluntària en matèria de grups vulnerables.

També està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

Així mateix, la pregunta es basa en [els articles 10 i 16 del Conveni marc del](#)

[Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la igualtat i la no discriminació i sobre el marc de gestió de riscos i impactes dels sistemes d'IA. Juntament amb el dret a la igualtat i a la no discriminació, reconegut a [l'article 14 de la Constitució Espanyola](#), [l'article 15 de l'Estatut d'Autonomia de Catalunya](#) i [l'article 6 de la Llei 33/2011, d'octubre, general de salut pública](#).

Finalment, la pregunta guarda relació amb l'article [The ethics of AI in health care: A mapping review de Morleya, Machadoa, et al.](#)

2.1.6. Heu adoptat mesures que permetin mitigar possibles biaixos que es puguin generar a partir de les dades d'entrada o problemes de control del sistema, tant en la fase de disseny com en la implementació del sistema?

Aquesta pregunta fa referència [als articles 9, 10.3, 15.3, 53.1, 55.1 i 95.2. del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a les mesures de gestió de riscos, la governança de dades, la resiliència i robustesa dels sistemes d'IA, les obligacions dels proveïdors del models de propòsit general i de propòsit general de risc sistèmic i la promoció de codis de conducta d'aplicació voluntària en matèria de grups vulnerables.

També està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

Així mateix, la pregunta es basa en [els articles 10 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la igualtat i la no discriminació i sobre el marc de gestió de riscos i impactes dels sistemes d'IA. Juntament amb el dret a la igualtat i a la no discriminació, reconegut a [l'article 14 de la Constitució Espanyola](#), [l'article 15 de l'Estatut d'Autonomia de Catalunya](#) i [l'article 6 de la Llei 33/2011, d'octubre, general de salut pública](#).

Finalment, la pregunta guarda relació amb l'article [The ethics of AI in health care: A mapping review de Morleya, Machadoa, et al](#) i el principi d'equitat de la guia [FUTURE-AI](#).

2.1.7. Heu establert una política de participació de diferents grups d'interès, incloent-hi els usuaris clínics i els pacients, del vostre sistema d'IA amb l'objectiu de mitigar possibles biaixos?

Aquesta pregunta fa referència a [l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció

de codis de conducta d'aplicació voluntària en matèria de participació multisectorial.

Així mateix, es relaciona amb [l'article 5 de la Llei 33/2011, del 4 d'octubre, general de Salut Pública](#) i [l'article 10.10 de la Llei 14/1986, de 25 d'aril, General de Sanitat](#) sobre el dret a la participació en matèria de salut.

La pregunta també es fonamenta en [l'article 6 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Finalment, es basa a les recomanacions contingudes a [l'apartat Diversity, Non-discrimination and Fairness: Stakeholder Participation de l'ALTAI](#) i a [l'apartat Injustice in Algorithmic Systems del Algorithmic Equity Toolkit d'ACLU](#). En el mateix sentit, guarda relació amb la publicació [The ethics of AI in health care: A mapping review de Morleya, Machadoa, et al.](#)

2.2. Inclusió i diversitat

Nº Preguntes

2.2.1. Heu avaluat si el vostre sistema s'adapta a un ampli espectre de preferències i habilitats individuals, especialment quan es tracti de tecnologia d'assistència que hagi de ser usada directament per alguna persona pacient?

Aquesta pregunta fa referència a [l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció de codis de conducta d'aplicació voluntària en matèria de grups vulnerables.

També està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

D'altra banda, la pregunta es fonamenta amb el dret a l'accés igualitari a la salut, reconegut a [l'article 3 del Conveni d'Oviedo](#), així com en el dret a la igualtat i a la no discriminació, reconegut a [l'article 14 de la Constitució Espanyola](#), [l'article 15 de l'Estatut d'autonomia de Catalunya](#) i [l'article 6 de la Llei 33/2011, d'octubre, general de salut pública](#). Juntament amb [l'article 18 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets](#)

[humans, democràcia i estat de dret](#), que versa sobre els drets de les persones amb discapacitat.

Finalment, fa referència a [l'apartat Diversity, Non-discrimination and Fairness: Accessibility and Universal Design de l'ALTAI](#), i està relacionada amb [l'apartat Fairness del AI Ethics Assessment de GSMA](#), així com amb el principi d'usabilitat de la guia [FUTURE-AI](#).

Exemple: Hem determinat que el nostre sistema d'IA de triatge en atenció primària s'utilitza en un CAP d'un barri amb un alt percentatge de gent d'edat avançada. El sistema basa les seves recomanacions en la anàlisi d'electrocardiogrames i anomalies comuns. Hem tingut en compte en la formació de les persones usuàries que el tipus d'anomalies que es reporten estan dintre de l'espectre de condicions clíniques que s'observen habitualment al CAP, i per tant el personal usuari té les habilitats individuals necessàries per interpretar la informació i informar als pacients en cas de dubtes.

2.2.2. **Heu valorat que dins de l'equip de desenvolupament i/o implementació del vostre sistema d'IA s'inclouin dones?**

Aquesta pregunta fa referència a [l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció de codis de conducta d'aplicació voluntària en matèria de diversitat dels equips de desenvolupament.

Així mateix, la pregunta es basa en [l'apartat Diversity, Non-discrimination and Fairness de l'ALTAI](#), [l'apartat fairness del toolkit d'ICO](#), i [el principi Participation del document A Human Rights-Based Approach To Data de l'ACNUDH](#). Finalment, la pregunta també guarda relació amb la publicació [Embedded ethics: a proposal for integrating ethics into the development of medical AI](#), de McLennan, Fiske, et al.

2.2.3. **Heu valorat que dins de l'equip de desenvolupament i/o implementació del vostre sistema d'IA hi hagi diversitat poblacional?**

Aquesta pregunta fa referència a [l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció de codis de conducta d'aplicació voluntària en matèria de diversitat dels equips de desenvolupament.

Així mateix, la pregunta es basa en [l'apartat Diversity, Non-discrimination and Fairness de l'ALTAI](#), [l'apartat fairness del toolkit d'ICO](#), i [el principi Participation del document A Human Rights-Based Approach To Data de l'ACNUDH](#). Finalment, la pregunta també guarda relació amb la publicació

[*Embedded ethics: a proposal for integrating ethics into the development of medical AI*](#), de McLennan, Fiske, *et al.*

2.2.4. Heu valorat que dins de l'equip de desenvolupament i/o implementació del vostre sistema d'IA s'incloguin de diferents disciplines de coneixement més enllà de la ciència de dades, l'enginyeria i les ciències de la salut?

Aquesta pregunta fa referència a [l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció de codis de conducta d'aplicació voluntària en matèria de diversitat dels equips de desenvolupament.

Així mateix, la pregunta es basa en [l'apartat Diversity, Non-discrimination and Fairness de l'ALTAI](#), [l'apartat fairness del toolkit d'ICO](#), i [el principi Participation del document A Human Rights-Based Approach To Data de l'ACNUDH](#). Finalment, la pregunta també guarda relació amb la publicació [*Embedded ethics: a proposal for integrating ethics into the development of medical AI*](#), de McLennan, Fiske, *et al.*

Principi 3: Seguretat i no maleficència

Aquest principi proposa que els sistemes d'IA haurien de funcionar de manera fiable, d'acord amb el propòsit previst, i assegurant que les persones usuàries clíniques disposen de les capacitats i els recursos per a la implementació segura dels dispositius, de manera que es garanteixi la seguretat de les persones pacients que siguin afectades pel mateix. Les preguntes d'aquest apartat provenen majoritàriament de les seccions relatives a responsabilitat, traçabilitat i robustesa tècnica dels diferents instruments i recomanacions consultades.

3.1. Guies d'utilització i ús segur del sistema

Nº	Preguntes
3.1.1.	Heu establert guies d'utilització i bones pràctiques dels sistemes d'IA que són necessàries i accessibles a les persones usuàries clíniques? Aquesta pregunta es basa en als articles 11.1, 13.2 i 3 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA], de la

Comissió Europea en relació a la documentació tècnica, la transparència i comunicació de la informació i les obligacions dels proveïdors de models fundacionals.

La pregunta també estan relacionades amb [l'article 10.4 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [10.4 del Reglament 2017/746 del Parlament Europeu i del Consell, de 5 d'abril del 2017, sobre els productes sanitaris per a diagnòstic in vitro](#) que versen sobre la informació dels productes.

Així mateix, guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat.

Finalment, la pregunta està relacionada amb el principi d'usabilitat de la guia [FUTURE-AI](#).

Exemple: Hem elaborat instruccions d'ús i cursos de formació sobre el sistema d'IA i aquesta documentació respectiva és accessible a través de la intranet de l'hospital.

3.1.2. **Considereu que els desenvolupadors del sistema d'IA han proporcionat suficient informació per fer-ne un ús segur i fiable que no posi en perill a les persones pacients?**

Aquesta pregunta es fonamenta en [l'article 11.1, 13.2, 3, 50 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), de la Comissió Europea amb relació a la documentació tècnica, la transparència i comunicació d'informació als usuaris i les obligacions dels proveïdors de models de propòsit general.

La pregunta també es basa en [l'article 10.4 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [10.4 del Reglament 2017/746 del Parlament Europeu i del Consell, de 5 d'abril del 2017, sobre els productes sanitaris per a diagnòstic in vitro](#) que versen sobre la informació dels productes.

Així mateix, guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

Exemple: Hem establert guies d'utilització del sistema on les persones usuàries disposen de documentació detallada i hem elaborat sessions de formació que permeten fer un ús del sistema d'IA amb seguretat i garanties, així com formació continuada sobre les actualitzacions del sistema.

3.1.3. **Heu desenvolupat mecanismes clars i ben definits per tal que les persones usuàries clíniques puguin contactar amb els desenvolupadors del sistema d'IA en cas de necessitat?**

Aquesta pregunta es fonamenta en [els articles 11.1, 13.2, 3 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la documentació tècnica, la transparència i comunicació d'informació als usuaris i les obligacions dels proveïdors de models de propòsit general.

La pregunta també fa referència a [l'article 10.4 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), i amb [l'article 10.4 del Reglament 2017/746 del Parlament Europeu i del Consell, de 5 d'abril del 2017, sobre els productes sanitaris per a diagnòstic in vitro](#), que versen sobre la informació que ha de constar a la documentació tècnica dels productes sanitaris.

De l'altra banda, la pregunta es fonamenta en [l'article 14 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre els recursos davant les violacions de drets derivades dels sistemes d'IA.

Així mateix, la pregunta es basa en [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics.

Exemple: L'equip desenvolupador ha habilitat un sistema de "tiquets digitals" que permet a les persones usuàries enviar preguntes i comentaris sobre el sistema d'IA.

3.1.4. **Heu desenvolupat mecanismes per actualitzar periòdicament la informació disponible sobre el rendiment del sistema d'IA?**

Aquesta pregunta fa referència [als articles 11.1, 13.3 i 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la documentació tècnica del sistema d'IA, la transparència i la comunicació d'informació als usuaris, i la promoció de codis de conducta

d'aplicació voluntària en matèria de participació multisectorial.

D'altra banda, la pregunta també es basa en [l'article 10.4 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), i [l'article 10.4 del Reglament 2017/746 de 5 d'abril de 2017, sobre els productes sanitaris de diagnòstic in vitro](#), que versen sobre la documentació tècnica dels productes sanitaris. Així com [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris sobre les garanties sanitàries dels dispositius mèdics](#).

Així mateix, guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

A més, es relaciona amb [els articles 4 i 5 de la Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica](#) relativa al dret a la informació assistencial.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 8 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió dels sistemes d'IA. Així com en el dret a la informació en l'àmbit de la salut, [reconegut a l'article 23 de l'Estatut d'Autonomia de Catalunya](#).

Finalment, la pregunta té relació amb les recomanacions contingudes a [l'apartat Diversity, Non-discrimination and Fairness: Stakeholder Participation de l'ALTAI](#) i amb el principi de traçabilitat de la guia [FUTURE-AI](#).

Exemple: Hem establert un procés d'actualització del rendiment del sistema d'IA amb responsables clars i una periodicitat setmanal. A més, hem establert vies de comunicació amb altres centres hospitalaris que utilitzen el mateix sistema i es celebren reunions periòdiques per compartir l'experiència entre les persones usuàries del sistema d'IA.

3.2. Gestió del risc

Nº Preguntes

3.2.1. El vostre sistema ha estat avaluat sobre els seus possibles impactes a la salut, la seguretat i els drets de les persones?

Aquesta pregunta fa referència [als articles 9.2, 27.1, 55.1 i 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la gestió de riscos i l'avaluació de l'impacte dels sistemes d'IA sobre els drets fonamentals, les obligacions dels proveïdors de sistemes de propòsit general de risc sistèmic i els codis de conducta d'aplicació voluntària.

La pregunta també està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

També es fonamenta en [l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Així mateix, la pregunta es relaciona amb [els apartats I punt 1.4 i II de l'Informe sobre avaluacions d'impacte dels Drets Fonamentals en relació amb la Llei de Serveis Digitals d'AN i ECNL](#), Juntament amb la publicació [The ethics of AI in health care: A mapping review](#) de Morleya, Machadoa, et al.

Exemple: Durant el procés de certificació com a dispositiu mèdic hem completat un anàlisi de risc exhaustiu. A més, hem elaborat instruccions d'ús, cursos de formació, així documentació necessària i accessible sobre el sistema de IA.

3.2.2. El vostre sistema ha passat per una avaluació per garantir que compleix amb els requisits de seguretat legals?

Aquesta pregunta es basa en [els articles 9.2, 27.1, 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la gestió de riscos, i l'avaluació de l'impacte dels sistemes d'IA sobre els drets fonamentals i les obligacions dels proveïdors de sistemes de propòsit general de risc sistèmic.

Així mateix, la pregunta està relacionada amb [l'article 62.4 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 58.5 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#) que versen sobre les

avaluacions clíniques i de funcionament dels productes sanitaris. Juntament amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [l'article 4 del Conveni d'Oviedo](#), relatiu als estàndards professionals, i en [l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Finalment, la pregunta es relaciona amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), sobre les garanties sanitàries dels productes i [l'article 3.b\) del Reial Decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics](#).

3.2.3. Heu implementat alguna metodologia d'avaluació o autoavaluació periòdica (interna o externa) reconeguda per identificar possibles funcionaments erronis i incrementar la confiança en l'ús del sistema d'IA?

La pregunta fa referència [als articles 9, 14, 17, 27.1, 55.1 i 72 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a les mesures de gestió de riscos, a la vigilància humana, al sistema de gestió de qualitat, a lavaluació de l'impacte dels sistemes d'alt risc sobre els drets fonamentals, a les obligacions dels proveïdors dels sistemes de propòsit general de risc sistèmic, i al seguiment postmercat dels sistemes d'alt risc.

Així mateix, la qüestió guarda relació amb [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre els sistemes de gestió de qualitat dels productes sanitaris. I també, [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#) relatiu a les garanties sanitàries dels dispositius mèdics. Juntament amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [l'article 4 del Conveni d'Oviedo](#), que versa sobre els estàndards professionals, i en [l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

La pregunta també es basa en les guies [The medical algorithmic audit de Liu, Glocker, et al.](#) i [Holding AI to Account: Challenges for the Delivery of Trustworthy AI in Healthcare](#), de Rob Procter, Peter Tolmie, i Mark Rouncefield, així com en el principi d'usabilitat de la guia [FUTURE-AI](#).

Exemple: Hem establert auditories externes de forma periòdica que permeten l'avaluació detallada, estructurada i constant del rendiment, els errors i altres possibles mancances del sistema d'IA.

3.2.4. **Heu seguit metodologies robustes, reconegudes i reproduïbles de validació per seleccionar i entrenar les components algorítmiques del sistema d'IA en totes les etapes del seu desenvolupament?**

La pregunta fa referència a [l'article 15.4 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la precisió, la robustesa i la ciberseguretat del sistemes d'alt risc.

La pregunta també està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

També es fonamenta en [els articles 12 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la fiabilitat i el marc de gestió de riscos i impactes dels sistemes d'IA.

Així mateix, la pregunta es basa en les guies [The medical algorithmic audit de Liu, Glocker, et al.](#) i [Holding AI to Account: Challenges for the Delivery of Trustworthy AI in Healthcare](#), de Rob Procter, Peter Tolmie, i Mark Rouncefield.

Exemple: El nostre sistema d'IA de triatge ha estat desenvolupat fent servir sistemes de validació estadística consistents amb les tècniques acceptades per la comunitat científica i, fins i tot, reconegut per a revistes d'alt nivell que han fet servir un procés de revisió entre iguals.

3.2.5. **Heu informat sobre els riscos residuals (aquells que no han pogut ser eliminats o corregits després d'haver pres totes les mesures possibles)?**

Aquesta pregunta es fonamenta en [l'article 9.4 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versa sobre la gestió de riscos.

També està relacionada amb [l'article 32 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), i [l'article 29 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#) que versen sobre el resum de seguretat i el rendiment clínic dels productes sanitaris. Juntament amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [els articles 8 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió i el marc de gestió de riscos i impactes dels sistemes d'IA.

Per l'altra banda, la pregunta es relaciona amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), sobre les garanties sanitàries dels productes i [l'article 4 del Reial Decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics](#).

La pregunta també guarda relació amb [l'article 10 de la Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica](#), relatiu al consentiment informat.

Finalment, la pregunta està relacionada amb el principi sobre traçabilitat de la guia [FUTURE-AI](#).

3.2.6. **Heu tingut en compte la possibilitat que s'involucri a persones d'una especialitat clínica diferent de l'actora principal durant el context de desplegament de la solució?**

Aquesta pregunta fa referència [als articles 9.5 i 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a l'avaluació de riscos i a les obligacions de les proveïdores de sistemes de propòsit general d'alt risc.

Així mateix, guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

La pregunta també es fonamenta en les guies establertes per la [WHO a l'apartat](#)

[*Governance of the private sector del document Ethics and governance of artificial intelligence for health*](#), sobre els recursos disponibles a l'hora d'operar un sistema d'IA.

Exemple: Un sistema de detecció d'anomalies en electrocardiogrames inclou, a més de la participació de cardiòlegs, la possibilitat que especialistes en medicina interna i urgències en facin ús si un cardiòleg no es troba disponible. Sempre indicant clarament les mesures formatives necessàries que han de complir aquests especialistes i amb relació a les preguntes anteriors.

Principi 4: Responsabilitat i retiment de comptes

El principi de Responsabilitat i retiment de comptes fa referència a l'actuació amb integritat en les accions, decisions i conseqüències de la gestió de dades i el desenvolupament i l'ús de sistemes d'IA. L'ús d'un sistema d'IA a l'àmbit de la salut ha d'estar recolzat i garantit per mecanismes clars, interns i externs, d'atribució de la responsabilitat i el retiment de comptes en relació amb el control de l'ús, els resultats i el rendiment del sistema en totes les fases del seu cicle de vida; en relació amb les conseqüències legals de possibles errors; i amb especial èmfasi en la supervisió humana del sistema quan aquesta sigui necessària.

4.1. Control

Nº	Preguntes
4.1.1.	Heu establert mecanismes de control amb participació humana durant les diferents fases de vida del sistema d'IA (en especial en la fase de recol·lecció i tractament de dades, en els pilots i els assajos clínics amb personal clínic i/o pacients)? La pregunta fa referència als articles 14.1 i 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA], que versen sobre la vigilància humana dels sistemes d'alt risc i les obligacions dels proveïdors dels models generals de risc sistèmic. La pregunta també es basa en l'article 86 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris, sobre l'actualització periòdica de l'informe de seguretat, així com

[l'article 62.4](#) del mateix text i [l'article 58.5](#) del [Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), que versen sobre les avaluacions clíniques i de funcionament dels productes sanitaris. Juntament amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [l'article 4 del Conveni d'Oviedo](#), relatiu als estàndards professionals, i en [els articles 8 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió i el marc de gestió de riscos i impactes dels sistemes d'IA.

Per l'altra banda, la pregunta es relaciona amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), sobre les garanties sanitàries dels productes i [l'article 3.f\) del Reial Decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics](#), que fa referència a la responsabilitat del personal sanitari.

Finalment, la pregunta també està relacionada amb els documents [The medical algorithmic audit](#) de Liu, Glocker, et al. i [Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability](#) de Coeckelbergh.

Exemple: En el desenvolupament i desplegament d'un nou sistema d'IA hem definit les diferents persones responsables de vetllar per la qualitat i representativitat de les dades emprades per entrenar l'algoritme, així com de comprovar que la validació clínica del sistema s'ha realitzat correctament, seguint les regulacions específiques, i sent coherent amb el processament de dades inicial per l'entrenament del sistema.

4.1.2. **Heu establert mecanismes de control amb participació humana que permetin mitigar possibles errors del sistema i habilitar modificacions en el protocol clínic d'ús sense que això comprometi la qualitat del resultat?**

Aquesta pregunta es basa en [els articles 14.1 i 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la vigilància humana dels sistemes d'alt risc i les obligacions dels proveïdors dels models generals de risc sistèmic.

La pregunta també està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre

el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [l'article 4 del Conveni d'Oviedo](#), relatiu als estàndards professionals, i en [els articles 8 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió i el marc de gestió de riscos i impactes dels sistemes d'IA.

La pregunta també es basa en la guia [The medical algorithmic audit de Liu, Glocker, et al.](#) i l'article [Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability de Coeckelbergh](#).

Exemple: Hem establert mecanismes de control on el rendiment del sistema és constantment monitoritzat per detectar possibles anomalies i accions de mitigació així com per ajustar i adaptar el protocol clínic. En concret hem determinat que, per exemple, quan una de les variables a imputar sobre el pacient no està disponible, el sistema no es pot utilitzar amb garanties, i la persona usuària disposa d'un protocol clínic alternatiu per prendre la decisió.

4.1.3. **Heu identificat i definit clarament el nivell de control de cadascuna de les persones involucrades en la creació, el desenvolupament i l'ús del sistema d'IA?**

Aquesta pregunta es basa en [l'article 14 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la vigilància humana de sistemes d'IA.

La pregunta també està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

La pregunta també es fonamenta en [els articles 8 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió i el marc de gestió de riscos i impactes dels sistemes d'IA.

Finalment, la pregunta fa referència al [subapartat Human Oversight de l'ALTAI](#), i [Accuracy & Error in Algorithmic Systems del l'Algorithmic Equity Toolkit d'ACLU](#), així com amb [l'apartat Human Agency & Oversight del AI Ethics Assessment de GSMA](#).

Exemple: Hem preparat col·laborativament una infografia que estableix els límits i els nivells de confiança als resultats del sistema d'IA, identifica

situacions o decisions crítiques on cal la intervenció humana i proporciona informació amb relació al nivell de control del sistema d'IA per part de cada persona (tenint en compte llur aplicació en el procés de decisió clínica).

4.1.4. Heu desenvolupat eines que permetin validar i contrastar clínicament els resultats del sistema d'IA?

Aquesta pregunta fa referència [als articles 9, 14 i 17 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a les mesures de gestió de riscos, a la vigilància humana i al sistema de gestió de qualitat dels sistemes d'IA.

La pregunta també està relacionada amb [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre els sistemes de gestió de qualitat dels productes sanitaris. Així com [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics. Juntament amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

Així mateix, es fonamenta en [els articles 12 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la fiabilitat i el marc de gestió de riscos i impactes dels sistemes d'IA.

Finalment, també es basa en el principi sobre traçabilitat de la guia [FUTURE-AI](#).

Exemple: Hem establert mecanismes amb els quals els resultats del sistema d'IA són contrastats periòdicament amb els resultats obtinguts per altres mètodes i són revisats pel personal clínic rellevant.

4.1.5. Heu definit clarament quina és la contribució del sistema d'IA en la presa de decisions clínica i quina és la responsabilitat de les persones involucrades en el control i operació (en especial les usuàries clíniques) en la decisió clínica final?

Aquesta pregunta fa referència a [l'article 14.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versa sobre la vigilància humana dels sistemes d'alt risc.

La pregunta també està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre

el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [l'article 4 del Conveni d'Oviedo](#), que versa sobre es estàndards professionals, i en [els articles 9 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la responsabilitat i el retiment de comptes pels impactes adversos de la IA i el marc de gestió de riscos i impactes dels sistemes d'IA.

Així mateix, la pregunta està relacionada amb [l'apartat Accountability de l'ALTAI](#).

Exemple: Un nou software de triatge determina la probabilitat que la persona pacient pateixi una arrítmia, però és el metge qui, amb aquesta informació i un protocol clínic associat, determina en darrera instància el diagnòstic i decideix sobre el tractament.

4.1.6. **Heu definit i comprovat el nivell de formació, comprensió i habilitats necessàries per part del personal clínic involucrat en el control i operació del sistema d'IA?**

Aquesta pregunta es fonamenta en [l'article 14.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versa sobre la vigilància humana dels sistemes d'IA d'alt risc.

També està relacionada amb [l'article 32 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 29 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), relatius a la seguretat i el rendiment clínic dels dispositius mèdics. Juntament amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

Així mateix, la pregunta es basa en [els articles 8, 16 i 20 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la transparència i la supervisió, el marc de gestió de riscos i impactes dels sistemes d'IA i l'alfabetització i les habilitats digitals.

De la mateixa manera, la pregunta es relaciona amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#) sobre les garanties sanitàries dels dispositius mèdics.

Finalment, la pregunta guarda relació amb l'article [The ethics of AI in health care: A mapping review de Morleya, Machadoa, et al](#) i amb el principi sobre universalitat de la guia [FUTURE-AI](#).

Exemple: Un nou sistema d'anàlisi automatitzat d'imatges radiològiques s'ha implantat a l'hospital i el personal clínic de radiologia ha rebut sessions de formació sobre el seu funcionament i operació abans de fer-ne ús.

4.1.7. Heu establert mecanismes de control amb participació humana per limitar l'ús del sistema per part de grups de pacients determinats o en contextos clínics específics, quan el rendiment pugui esdevenir inacceptable, sense que això comprometi la qualitat del resultat?

Aquesta pregunta es basa en [els articles 14.1 i 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la vigilància humana dels sistemes d'alt risc i les obligacions dels proveïdors dels models generals de risc sistèmic.

La pregunta també està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

De la mateixa manera, la pregunta es fonamenta en [l'article 4 del Conveni d'Oviedo](#), relatiu als estàndards professionals, i en [els articles 8 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i la supervisió i el marc de gestió de riscos i impactes dels sistemes d'IA.

La pregunta també es basa en la guia [The medical algorithmic audit de Liu, Glocker, et al.](#) i l'article [Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability de Coeckelbergh](#).

4.1.8. Us heu assegurat que durant tot en el procés de creació, desenvolupament, desplegament i disseny dels mecanismes de control del sistema hi hagin participat totes les parts implicades? En particular, les expertes clíniques en el context de la solució, les representants dels pacients i les creadores de sistemes d'IA.

Aquesta pregunta fa referència a [l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció de codis de conducta d'aplicació voluntària en matèria de participació de les persones interessades.

També es fonamenta en [l'article 16 del Conveni marc del Consell d'Europa](#)

[sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Així mateix, la pregunta es basa en les guies establertes per la [WHO a l'apartat *Governance of the public sector* del document *Ethics and governance of artificial intelligence for health*](#) i [The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine](#), de Thomas Ploug i Søren Holm.

Exemple: En el disseny de mecanismes de control d'un sistema d'IA per la gestió de la insulina s'han inclòs persones desenvolupadores de la solució, representants de pacients diabètics i expertes en el seu tractament i de possibles comorbiditats més comunes.

4.1.9. **Heu establert mecanismes perquè les persones creadores i desenvolupadores del sistema d'IA tinguin coneixement del context clínic d'ús, del rendiment i la possibilitat de corregir i millorar aquests?**

Aquesta pregunta es fonamenta en [l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció de codis de conducta d'aplicació voluntària en matèria de participació de les persones interessades.

La pregunta també es relaciona amb [l'article 62.7 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 58.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre les avaluacions clíniques dels productes sanitaris i les avaluacions de funcionament dels productes sanitaris in vitro.

De la mateixa manera, fa referència a [l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Així mateix, la pregunta guarda relació amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#) que versa sobre les garanties sanitàries dels dispositius mèdics.

Finalment, la pregunta també es basa en les guies establertes per la [WHO a l'apartat *Autonomous decision-making* del document *Ethics and governance of artificial intelligence for health*](#) i [The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence](#)

[in Medicine](#), de Thomas Ploug I Søren Holm, així com el principi sobre universalitat de la guia [FUTURE-AI](#).

Exemple: Durant el procés de desenvolupament d'un sistema d'IA per la gestió de la insulina, les persones desenvolupadores han rebut formació específica del context clínic en què es desplegarà la solució per part d'experts en diabetis. A més a més, s'ha establert un protocol de comunicació periòdica entre les persones responsables del desplegament del sistema d'IA i les desenvolupadores per identificar possibles anomalies i planejar millores.

4.1.10. **En cas d'indisponibilitat temporal del personal que opera el sistema d'IA, heu tingut en compte de quins recursos disposeu per a assegurar el seu correcte funcionament?**

La pregunta fa referència a [l'article 8 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), sobre la llicència prèvia de funcionament d'instal·lacions.

Així mateix, es fonamenta en [l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Aquesta pregunta també es basa en les guies establertes per la [WHO a l'apartat Autonomous decision-making del document Ethics and governance of artificial intelligence for Health](#), sobre els recursos disponibles a l'hora d'operar un sistema d'IA.

Exemple: En el cas que les persones especialistes que operen el sistema d'IA no es trobin disponibles hem establert un protocol secundari que no necessita l'ús del sistema d'IA en el procés clínic i determina clarament l'especialista alternatiu.

4.1.11. **Heu establert quines són les persones responsables d'actualitzar el sistema d'IA que permeti garantir el seu funcionament òptim i assegurar-ne el control en tot moment?**

Aquesta pregunta fa referència a [l'article 17.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versa sobre els sistemes de gestió de qualitat dels sistemes d'alt risc i les obligacions de monitoratge post-comercialització.

La pregunta també està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

També es fonamenta en [els articles 8, 9 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la transparència i la supervisió, la responsabilitat i el retiment de comptes pels impactes adversos de la IA i el marc de gestió de riscos i impactes dels sistemes d'IA.

Així mateix, es basa en les guies establertes per [WHO a l'apartat *Autonomous decision-making* del document *Ethics and governance of artificial intelligence for health*](#) i les recomanacions presentades a [The medical algorithmic audit de Liu, Glocker, et al.](#) sobre auditories de sistemes d'IA en salut.

4.1.12. Heu establert mecanismes per revisar les recomanacions fetes pel sistema d'IA, les discrepàncies entre aquestes i la decisió clínica final i investigar proactivament els possibles errors (incloent-hi la planificació d'accions o protocols a dur a terme)?

Aquesta pregunta fa referència [als articles 9.2, 17, 55.1 i 72 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la gestió de riscos, sistema de gestió de qualitat dels sistemes, les obligacions dels proveïdors de sistemes de propòsit general de risc sistèmic, i la vigilància post mercat.

La pregunta també està relacionada amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

També es fonamenta en [els articles 8 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la transparència i la supervisió i el marc de gestió de riscos i impactes dels sistemes d'IA.

De l'altra banda, està relacionada amb [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre els sistemes de gestió de qualitat dels productes sanitaris.

Així mateix, la pregunta es relaciona amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics.

Finalment, la pregunta també es basa en [l'apartat Accountability de l'ALTAI](#), així com en les guies establertes per la [WHO a l'apartat Ensure transparency, explainability and intelligibility del document Ethics and governance of artificial intelligence for health](#) i les recomanacions presentades a [The medical algorithmic audit de Liu, Glocker, et al.](#) sobre auditories de sistemes d'IA en salut.

Exemple: Hem establert un sistema d'auditoria externa periòdica que avalua els resultats del monitoratge del sistema d'IA i aporta recomanacions i accions concretes per mitigar possibles errors i disfuncionaments.

4.2. Retiment de comptes

Nº Preguntes

4.2.1. Heu determinat qui és legalment responsable de les decisions derivades de l'ús del sistema d'IA durant tot el seu cicle de vida?

Aquesta pregunta fa referència [als articles 17.1 i 72 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), sobre el marc de responsabilitat inclòs en els sistemes de gestió de qualitat dels sistemes d'alt risc i les obligacions de monitoratge post comercialització.

També es fonamenta en [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre els sistemes de gestió de qualitat dels productes sanitaris.

De l'altra banda, la pregunta està relacionada amb [l'article 9 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre la responsabilitat i el retiment de comptes pels impactes adversos de la IA.

La pregunta també està relacionada amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics.

Finalment, la pregunta es basa en l'article [The ethics of AI in health care: A mapping review de Morleya, Machadoa, et al.](#)

4.2.2. Heu determinat en quins casos, i en quin grau, personal clínic usuari del sistema d'IA pot tenir responsabilitat legal respecte de les decisions clíniques derivades del seu ús?

Aquesta pregunta fa referència [als articles 17.1 i 72 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), sobre el marc de responsabilitat inclòs en els sistemes de gestió de qualitat dels sistemes d'alt risc i les obligacions de monitoratge post comercialització.

També es fonamenta en [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre els sistemes de gestió de qualitat dels productes sanitaris.

De l'altra banda, la pregunta està relacionada amb [l'article 9 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre la responsabilitat i el retiment de comptes pels impactes adversos de la IA.

La pregunta també està relacionada amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics.

Finalment, la pregunta es basa en l'article [The ethics of AI in health care: A mapping review de Morleya, Machadoa, et al.](#)

4.2.3. Heu determinat en quins casos, i en quin grau, l'organització sanitària sota el paraigua de la qual s'utilitza el sistema d'IA pot tenir responsabilitat legal respecte de les decisions clíniques derivades del seu ús?

Aquesta pregunta fa referència [als articles 17.1 i 72 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), sobre el marc de responsabilitat inclòs en els sistemes de gestió de qualitat dels sistemes d'alt risc i les obligacions de monitoratge post comercialització.

També es fonamenta en [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre els sistemes de gestió de qualitat dels productes sanitaris.

De l'altra banda, la pregunta està relacionada amb [l'article 9 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre la responsabilitat i el retiment de comptes pels impactes adversos de la IA.

La pregunta també està relacionada amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics.

Finalment, la pregunta es basa en l'article [The ethics of AI in health care: A mapping review de Morleya, Machadoa, et al.](#)

4.2.4. **Heu determinat en quins casos, i en quin grau, els creadors del sistema d'IA poden tenir responsabilitat legal respecte de les decisions clíniques derivades del seu ús?**

Aquesta pregunta fa referència [als articles 17.1 i 72 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), sobre el marc de responsabilitat inclòs en els sistemes de gestió de qualitat dels sistemes d'alt risc i les obligacions de monitoratge post comercialització.

També es fonamenta en [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre els sistemes de gestió de qualitat dels productes sanitaris.

De l'altra banda, la pregunta està relacionada amb [l'article 9 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre la responsabilitat i el retiment de comptes pels impactes adversos de la IA.

La pregunta també està relacionada amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics.

Finalment, la pregunta es basa en l'article [The ethics of AI in health care: A mapping review de Morleya, Machadoa, et al.](#)

4.2.5. **Heu establert mecanismes de revisió periòdica sobre responsabilitat i retiment de comptes amb els diferents agents involucrats en el desenvolupament i l'ús del vostre sistema d'IA?**

Aquesta pregunta fa referència [als articles 14 i 17 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), sobre la vigilància humana i la gestió de qualitat dels sistemes d'alt risc.

Així mateix, la pregunta està relacionada amb [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre els sistemes de gestió de qualitat dels productes sanitaris.

També es fonamenta en [els articles 9 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la responsabilitat i el retiment de comptes pels impactes adversos de la IA i el marc de gestió de riscos i impactes dels sistemes d'IA.

De la mateixa manera, la pregunta es relaciona amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics.

Finalment, té relació amb [l'apartat Accountability de l'ALTAI](#) i la guia [Holding AI to Account: Challenges for the Delivery of Trustworthy AI in Healthcare](#), de Rob Procter, Peter Tolmie, i Mark Rouncefield.

4.2.6. **El procés de retiment de comptes que heu dissenyat contempla que aquest sigui contínuament informat i revisat per les persones professionals en la pràctica clínica que interpreten les dades i segueixen el procés clínic habitual?**

Aquesta pregunta fa referència [als articles 13 i 14 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la transparència i la informació als desplegadors i a la vigilància humana de sistemes d'IA.

Així mateix, la pregunta està relacionada amb [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#) sobre els sistemes de gestió de qualitat dels productes sanitaris.

També fa fonamenta en [l'article 9 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre la responsabilitat i el retiment de comptes pels impactes adversos de la IA.

De la mateixa manera, es relaciona amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics.

Finalment, la pregunta es basa en [l'apartat Accountability de l'ALTAI](#) i a la guia [Holding AI to Account: Challenges for the Delivery of Trustworthy AI in Healthcare](#), de Rob Procter, Peter Tolmie, i Mark Rouncefield.

4.2.7. Heu tingut en compte en el desplegament del sistema d'IA la possible separació que existeix entre la persona usuària clínica i la persona pacient i l'efecte que pot tenir en l'atribució de responsabilitats i el retiment de comptes?

Aquesta pregunta fa referència a [l'article 17.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu al sistema de gestió de qualitat dels sistemes d'IA.

També està relacionada amb [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre els sistemes de gestió de qualitat dels productes sanitaris.

Així mateix, la pregunta es fonamenta en [els articles 9 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la responsabilitat i el retiment de comptes per els impactes adversos de la IA i el marc de gestió de riscos i impactes dels sistemes d'IA.

De la mateixa manera, la pregunta es relaciona amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels productes sanitaris.

Finalment, la pregunta també es basa en [The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine](#), de Thomas Ploug i Søren Holm i [Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability](#) de Coeckelbergh, sobre la contestabilitat dels sistemes d'IA per part dels pacients.

4.3. Figures de protecció i mesures de reparació

Nº	Preguntes
----	-----------

4.3.1.	Heu determinat quines figures de protecció (ex. responsable d'atendre les queixes relacionades amb el sistema d'IA) es poden establir per minimitzar o eliminar possibles impactes negatius derivats de l'ús del sistema d'IA?
--------	---

Aquesta pregunta es basa en [els articles 9 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la responsabilitat i el retiment de comptes pels impactes adversos de la IA i el marc de gestió de riscos i impactes dels sistemes d'IA.

També es fonamenta en [l'article 10.12 de la Llei 14/1986, de 25 d'aril, General de Sanitat](#), sobre els dret a utilitzar les vies de reclamació i de proposta de suggeriments i rebre resposta, i [l'article 5 de la Llei 33/2011, de 4 d'octubre, General de Salut Pública](#), relatiu al dret a a participació de la ciutadania en les accions de sanitat pública.

Exemple: Hem designat a una persona com a responsable d'atendre a les reclamacions relatives a les decisions preses pel sistema d'IA, i hem dissenyat un protocol d'avaluació i reparació o compensació dels possibles perjudicis causats. A més, hem informat a les persones usuàries i pacients sobre l'existència d'aquesta figura i les vies per contactar-hi, per tal que puguin derivar les seves preguntes o queixes sobre el sistema d'IA i aquestes es puguin gestionar degudament.

4.3.2.	Heu establert formes de reparació o compensació (ex. a través de pòlisses d'assegurances o altres mitjans monetaris adequats) en cas d'ocórrer algun mal funcionament del sistema o provocar un impacte advers a la persona pacient?
--------	---

Aquesta pregunta fa referència a [l'article 27.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a l'avaluació de l'impacte de sistemes d'alt risc sobre drets fonamentals.

També està relacionada amb [l'article 10.2 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), [l'article 10.15 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), i [l'article 32 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#).

De l'altra banda, la pregunta es fonamenta en [l'article 9 del Conveni marc del](#)

[Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre la responsabilitat i el retiment de comptes pels impactes adversos de la IA.

Així mateix, la pregunta es basa en [l'article 9 del Reial Decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics](#), sobre les indemnitzacions per danys i perjudicis causats pels productes sanitaris.

Finalment, es relaciona amb les guies establertes per la [WHO a l'apartat *Ensure transparency, explainability and intelligibility* del document *Ethics and governance of artificial intelligence for health*](#) i [l'*Ethics guidelines for trustworthy AI* de la Comissió Europea](#)

Exemple: Hem contractat una pòlissa d'assegurança que cobreix els possibles danys que el sistema d'IA pugui generar a les persones pacients.

4.3.3. **Heu considerat formes de reparació o compensació que incloguin també la responsabilitat per a les persones creadores del sistema d'IA perquè l'atribució de responsabilitats sigui recíproca?**

Aquesta pregunta fa referència a [l'article 9 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre la responsabilitat i el retiment de comptes pels impactes adversos de la IA.

Així mateix, la pregunta està relacionada amb [l'article 10.2 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), [l'article 10.15 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), i [l'article 32 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatius a les garanties sanitàries dels productes sanitaris.

També guarda relació amb [l'article 9 del Reial Decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics](#), sobre les indemnitzacions per danys i perjudicis causats pels productes sanitaris.

Finalment, es la pregunta també es basa en les guies establertes per la [WHO a l'apartat *Foster responsibility and accountability* del document *Ethics and governance of artificial intelligence for health*](#) i [l'*Ethics guidelines for*](#)

[*trustworthy AI* de la Comissió Europea.](#)

4.3.4. Heu establert mesures per facilitar la demostració del compliment de la legislació vigent, incloent-hi la Llei de Protecció de dades (GDPR), la regulació sobre dispositius mèdics (Reglament 2017/745) i/o de diagnòstic in vitro (Reglament 2017/746) i l'RIA?

Aquesta pregunta fa referència [als articles 11, 13, 17.1 i 43 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la documentació tècnica, la transparència i comunicació d'informació als usuaris, el sistema de gestió de qualitat dels sistemes d'alt risc i les avaluacions de conformitat.

Així mateix, la pregunta guarda vinculació amb [els articles 19, 20 i 52 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), sobre la declaració de conformitat i el marcatge CE dels productes sanitaris. I també es relaciona amb, [els articles 17, 18 i 48 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre les avaluacions de conformitat i el marcatge CE dels productes sanitaris.

Així mateix, la pregunta està relacionada amb [l'article 21 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a la informació que ha de fer-se pública al Registre de Responsables de la posada al mercat de productes sanitaris de l'Agència Espanyola de Medicaments i Productes Sanitaris.

Està relacionada també amb la [Secció 1 del Capítol IV del Reglament 2016/679, de 27 d'abril 2016 \[Reglament general de protecció de dades\]](#), que regula les obligacions legals dels responsables del tractament de dades i el seu règim de responsabilitat.

Finalment, la pregunta també es basa en el [subapartat Data Governance de l'ALTAI](#) i el [toolkit d'ICO sobre riscos i protecció de dades en la IA](#).

Exemple: Disposem de la certificació validada per la Comissió Europea que certifica que el nostre sistema d'IA està catalogat com a dispositiu mèdic d'una determinada classe seguint les regulacions dels dispositius mèdics (MDR) i és apte pel seu ús previst. A més a més, tenim clarament assignat i publicat qui és la persona responsable de respondre qüestions sobre regulació i la demostració del seu compliment.

4.3.5. Heu estudiat amb detall si el sistema d'IA pot ser considerat un producte

sanitari, d'acord amb la regulació sobre productes sanitaris i/o de diagnòstic in vitro, quina seria la seva classificació i quines conseqüències legals se'n deriven?

Aquesta pregunta està fonamenta amb [l'article 51 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 47 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre la classificació dels productes sanitaris.

Així mateix, la pregunta està relacionada amb la [Guia per a la qualificació i classificació d'una aplicació informàtica basada en IA com a producte sanitari segons els Reglaments Europeus MDR/IVDR elaborada per la Fundació TIC Salut Social](#).

Exemple: Hem avaluat en detall l'ús previst del sistema d'IA i la seva classificació segons la legislació vigent. El sistema està clarament definit i classificat com a dispositiu mèdic d'una determinada classe i s'ha certificat seguint les regulacions específiques del marc europeu.

4.3.6. **Heu estudiat amb detall quina és la classificació del vostre sistema d'acord amb el Reglament europeu sobre IA i quines conseqüències se'n deriven d'aquesta classificació?**

Aquesta pregunta fa referència [als articles 6 i 51 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la classificació dels sistemes d'alt risc i la classificació dels models de propòsit general.

Així mateix, la pregunta està relacionada amb la [Guia per a l'Aplicació del Reglament Europeu d'IA en Salut elaborada per la Fundació TIC Salut Social](#).

Exemple: El nostre sistema d'IA està pensat per facilitar el diagnòstic de patologies relacionades amb la pell a partir de fotografies. Aquesta descripció coincideix amb la definició de producte sanitari contemplada al Reglament de 2017/745 sobre dispositius mèdics. Aquest reglament es troba inclòs en llistat l'Annex I del Reglament europeu sobre IA. Segons l'article 6.1a) del RIA, els productes afectats pels reglaments que s'enumeren en el mencionat annex, seran considerats sistemes d'alt risc. En conseqüència, sabem que al nostre sistema li són d'aplicació els requisits legals que el RIA imposa als sistemes d'IA d'alt risc.

Principi 5: Privacitat

Aquest principi indica que el dret a la privacitat i la protecció de dades de les persones usuàries dels serveis sanitaris han de ser respectats en tot moment, de manera que l'ús de sistemes d'IA en l'àmbit sanitari s'ha de fer de manera que aquests siguin protegits. Les preguntes d'aquest apartat provenen majoritàriament de les seccions relatives a privacitat i governança de dades dels diferents instruments i recomanacions consultades.

5.1. Protecció de la privacitat

Nº	Preguntes
----	-----------

5.1.1.	Heu revisat si les dades de les persones pacients emprades en el sistema de dades o d'IA han estat anonimitzades tal com indica la Llei Orgànica de Protecció de dades (LO 3/18)?
--------	--

Aquesta pregunta fa referència [als articles 10.5, 17.1, 50.3 i 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la governança de dades i el sistema de gestió de qualitat dels sistemes d'alt risc, i els requisits dels sistemes de reconeixement d'emocions i de categorització biomètrica i dels models propòsit general de risc sistemàtic.

De la mateixa manera, la pregunta guarda relació amb [l'article 5 del Reglament 2022/868 del Parlament Europeu i del Consell, de 30 de maig del 2022, relatiu a la governança europea de dades \[Data Governance Act\]](#), que versa sobre les condicions de reutilització de les dades, així com [l'article 110.1 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [els articles 102, i 103 del Reglament 2017/746 del Parlament Europeu i del Consell, de 5 d'abril del 2017, sobre els productes sanitaris per a diagnòstic in vitro](#), sobre el tractament de les dades dels pacients que estableixen els requisits relatius al sistema de gestió de qualitat. En el mateix sentit, la pregunta es basa en [l'article 32 del Reglament 2016/679, de 27 d'abril 2016 \[Reglament general de protecció de dades\]](#) que versa sobre la seguretat de les dades personals.

Per l'altra banda, la pregunta es fonamenta en [l'article 10 del Conveni d'Oviedo](#), que versa sobre el dret a la vida privada i la informació, i en [l'article 11 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el dret a la privacitat i la protecció de dades personals. Juntament amb el dret fonamental a la intimitat recollit a [l'article 18 de la Constitució Espanyola](#) i el dret a la protecció de dades

personals contemplat a [l'article 31](#) de [l'Estatut d'Autonomia de Catalunya](#).

D'altra banda, la pregunta guarda vinculació amb [l'article 10.3](#) de la [Llei 14/1986, de 25 d'abril, General de Sanitat](#), [l'article 7](#) de la [Llei 33/2011, de 4 d'octubre, General de Salut Pública](#), i [l'article 7](#) de la [Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica](#), que versen sobre el dret a la privacitat de les persones pacients i la confidencialitat de les dades mèdiques.

Finalment, la pregunta també es basa en [l'apartat Privacy del document A Human Rights-Based Approach To Data](#) de l'ACNUDH, les recomanacions contingudes a [l'apartat 2.2.3. Privacy d'AIEIG](#) i a [l'apartat Privacy and Data Governance de l'ALTAI](#). Similarment, la pregunta guarda relació amb la publicació [Ethical Challenges and Solutions of Generative AI: An Interdisciplinary Perspective](#), d'Al-kfairy, Mustafa, et al.

Exemple: Hem establert mecanismes per assegurar que les dades dels pacients utilitzades pel sistema d'IA estan anonimitzades seguint els protocols d'anonimització establerts per la llei, i aquest procés es revisa per assegurar-ne la privacitat dels pacients.

5.1.2. **Heu tingut en compte que les mesures d'anonimització no poden limitar-se a ocultar el nom de la persona pacient, sinó que han d'impedir que pugui inferir-se la seva identitat a partir del conjunt d'informació?**

Aquesta pregunta fa referència [als articles 10.5, 17.1, 50.3 i 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la governança de dades i el sistema de gestió de qualitat dels sistemes d'alt risc, i els requisits dels sistemes de reconeixement d'emocions i de categorització biomètrica i dels models propòsit general de risc sistemàtic.

De la mateixa manera, la pregunta guarda relació amb [l'article 5 del Reglament 2022/868 del Parlament Europeu i del Consell, de 30 de maig del 2022, relatiu a la governança europea de dades \[Data Governance Act\]](#), que versa sobre les condicions de reutilització de les dades, així com [l'article 110.1 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [els articles 102, i 103 del Reglament 2017/746 del Parlament Europeu i del Consell, de 5 d'abril del 2017, sobre els productes sanitaris per a diagnòstic in vitro](#), sobre el tractament de les dades dels pacients que estableixen els requisits relatius al sistema de gestió de qualitat. En el mateix sentit, la pregunta es basa en [l'article 32 del Reglament 2016/679, de 27 d'abril 2016 \[Reglament general de protecció de dades\]](#) que versa sobre la seguretat de les dades personals.

Per l'altra banda, la pregunta es fonamenta en [l'article 10 del Conveni d'Oviedo](#), que versa sobre el dret a la vida privada i la informació, i en [l'article 11 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el dret a la privacitat i la protecció de dades personals. Juntament amb el dret fonamental a la intimitat recollit a [l'article 18 de la Constitució Espanyola](#) i el dret a la protecció de dades personals contemplat a [l'article 31 de l'Estatut d'Autonomia de Catalunya](#).

D'altra banda, la pregunta guarda vinculació amb [l'article 10.3 de la Llei 14/1986, de 25 d'abril, General de Sanitat](#), [l'article 7 de la Llei 33/2011, de 4 d'octubre, General de Salut Pública](#), i [l'article 7 de la Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica](#), que versen sobre el dret a la privacitat de les persones pacients i la confidencialitat de les dades mèdiques.

Finalment, la pregunta també es basa en [l'apartat Privacy del document A Human Rights-Based Approach To Data de l'ACNUDH](#), les recomanacions contingudes a [l'apartat 2.2.3. Privacy d'AIEIG](#) i a [l'apartat Privacy and Data Governance de l'ALTAI](#). Similarment, la pregunta guarda relació amb la publicació [Ethical Challenges and Solutions of Generative AI: An Interdisciplinary Perspective](#), d'Al-kfairy, Mustafa, et al.

Exemple: L'historial clínic de cada persona presenta particularitats diferents que poden permetre deduir la seva identitat fins i tot quan les seves dades identitàries han estat modificades (intervencions quirúrgiques realitzades, existència de pròtesis, patologies poc comunes, ingressos hospitalaris, etc.). En conseqüència, hem aplicat mesures d'anonimització de dades que impedeixen la reconstrucció total de l'historial clínic, de manera que no sigui possible inferir la identitat de la persona sobre la qual versen aquestes dades.

5.1.3. **En cas d'assistir-vos amb una IA generativa, heu procurat de no avocar al sistema les dades personals de les persones pacients i heu revisat la política de privacitat del mateix?**

Aquesta pregunta fa referència [als articles 50.3 i 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen els requisits dels sistemes de reconeixement d'emocions i de categorització biomètrica i dels models propòsit general de risc sistèmic.

De la mateixa manera, la pregunta guarda relació amb [l'article 32 del Reglament 2016/679, de 27 d'abril 2016 \[Reglament general de protecció de dades\]](#) que versa sobre la seguretat de les dades personals.

Per l'altra banda, la pregunta es fonamenta en [l'article 10 del Conveni d'Oviedo](#), que versa sobre el dret a la vida privada i la informació, i en [l'article 11 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el dret a la privacitat i la protecció de dades personals. Juntament amb el dret fonamental a la intimitat recollit a [l'article 18 de la Constitució Espanyola](#) i el dret a la protecció de dades personals contemplat a [l'article 31 de l'Estatut d'Autonomia de Catalunya](#).

Així mateix, la pregunta guarda vinculació amb [l'article 10.3 de la Llei 14/1986, de 25 d'abril, General de Sanitat](#), [l'article 7 de la Llei 33/2011, de 4 d'octubre, General de Salut Pública](#), i [l'article 7 de la Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica](#), que versen sobre el dret a la privacitat de les persones pacients i la confidencialitat de les dades mèdiques.

La pregunta també es basa en [l'apartat Privacy del document A Human Rights-Based Approach To Data de l'ACNUDH](#), i les recomanacions contingudes a [l'apartat 2.2.3. Privacy d'AIEIG](#) i a [l'apartat Privacy and Data Governance de l'ALTAI](#). De la mateixa manera, la pregunta guarda relació amb la publicació i [Generative AI and medical ethics: the state of play, de Zohny, Mann, et al.](#)

5.1.4. **Considereu que heu establert un nivell de transparència i explicabilitat que no posa en risc la privacitat de les persones pacients?**

Aquesta pregunta està fa referència [als articles 10.5, 17.1 i 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), sobre la governança de dades i el sistema de gestió de qualitat dels sistemes d'alt risc i els requisits dels models de propòsit general de risc sistemàtic.

De la mateixa manera, la pregunta es relaciona amb [l'article 5 del Reglament 2022/868 del Parlament Europeu i del Consell, de 30 de maig del 2022, relatiu a la governança europea de dades \[Data Governance Act\]](#), que versa sobre les condicions de reutilització de les dades, així com [l'article 110.1 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [els articles 102, i 103 del Reglament 2017/746 del Parlament Europeu i del Consell, de 5 d'abril del 2017, sobre els productes sanitaris per a diagnòstic in vitro](#), sobre el tractament de les dades dels pacients que estableixen els requisits relatius al sistema de gestió de qualitat. En el mateix sentit, la pregunta guarda relació amb [l'article 32 del Reglament 2016/679, de 27 d'abril 2016 \[Reglament general de protecció de dades\]](#), que versa sobre la seguretat de les dades personals.

Així mateix, la pregunta es fonamenta en [l'article 10 del Conveni d'Oviedo](#), que versa sobre el dret a la vida privada i la informació, i en [l'article 11 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el dret a la privacitat i la protecció de dades personals. Juntament amb el dret fonamental a la intimitat, recollit a [l'article 18](#) de la [Constitució Espanyola](#), el dret a la protecció de dades personals i contemplat a [l'article 31](#) de [l'Estatut d'Autonomia de Catalunya](#).

D'altra banda, la pregunta guarda vinculació amb [l'article 10.3](#) de la [Llei 14/1986, de 25 d'abril, General de Sanitat](#), [l'article 7](#) de la [Llei 33/2011, de 4 d'octubre, General de Salut Pública](#), i [l'article 7](#) de la [Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica](#), que versen sobre el dret a la privacitat de les persones pacients i la confidencialitat de les dades mèdiques.

Finalment, la pregunta es relaciona amb [l'apartat Right to Privacy, and Data Protection de les recomanacions sobre ètica en la Intel·ligència Artificial d'UNESCO](#).

5.1.5. Podeu per garantir que totes les parts implicades en el desenvolupament i implementació del sistema han establert mecanismes que permetin tractar les dades personals amb el grau de protecció legalment requerit i així evitar-ne el mal ús, alteració, pèrdua o sostracció per part de tercers?

La pregunta fa referència [als articles 10.5, 50.3 i 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la governança de dades i els requisits dels sistemes de reconeixement d'emocions i categorització biomètrica i dels models de propòsit general de risc sistemàtic.

Així mateix, la pregunta es relaciona amb [l'article 32 del Reglament 2016/679, de 27 d'abril 2016 \[Reglament general de protecció de dades\]](#), que versa sobre la seguretat de les dades personals i [l'article 21 del Reglament 2022/868 del Parlament Europeu i del Consell, de 30 de maig del 2022, relatiu a la governança europea de dades \[Data Governance Act\]](#), sobre els requisits específics per a protegir els drets i els interessos de les persones interessades i titulars de dades perquè respecta a les seves dades.

De la mateixa manera, guarda vinculació amb [l'article 10.3](#) de la [Llei 14/1986, de 25 d'abril, General de Sanitat](#), [l'article 7](#) de la [Llei 33/2011, de 4 d'octubre, General de Salut Pública](#), i [l'article 7](#) de la [Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria](#)

[d'informació i documentació clínica](#), que versen sobre el dret a la privacitat de les persones pacients i la confidencialitat de les dades mèdiques.

Per l'altra banda, la pregunta es fonamenta en [l'article 10 del Conveni d'Oviedo](#), que versa sobre el dret a la vida privada i la informació, i en [l'article 11 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el dret a la privacitat i la protecció de dades personals. Juntament amb el dret fonamental a la intimitat, recollit a [l'article 18](#) de la [Constitució Espanyola](#), el dret a la protecció de dades personals, contemplat a [l'article 31](#) de [l'Estatut d'Autonomia de Catalunya](#).

Finalment, guarda relació amb el principi [Privacy del document A Human Rights-Based Approach To Data](#) de l'ACNUDH, les recomanacions contingudes al [toolkit d'ICO sobre riscos i protecció de dades en la IA](#) i l'apartat [Privacy and Security](#) de [l'AI Ethics Assessment](#) de GSMA.

5.1.6. Podeu garantir que totes les parts implicades en el desenvolupament i implementació del sistema d'IA han establert mecanismes per tal que les dades de les persones afectades no siguin venudes, llogades ni cedides, sense el seu consentiment?

La pregunta fa referència [als articles 10.5, 50.3 i 55.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius a la governança de dades i els requisits dels sistemes de reconeixement d'emocions i categorització biomètrica i dels models de propòsit general de risc sistemàtic.

També es relaciona amb [l'article 32 del Reglament 2016/679, de 27 d'abril 2016 \[Reglament general de protecció de dades\]](#), que versa sobre la seguretat de les dades personals i [l'article 21 del Reglament 2022/868 del Parlament Europeu i del Consell, de 30 de maig del 2022, relatiu a la governança europea de dades \[Data Governance Act\]](#), sobre els requisits específics per a protegir els drets i els interessos de les persones interessades i titulars de dades perquè respecte a les seves dades.

Per l'altra banda, la pregunta es fonamenta en [l'article 10 del Conveni d'Oviedo](#), que versa sobre el dret a la vida privada i la informació, i en [l'article 11 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el dret a la privacitat i la protecció de dades personals. Juntament amb el dret fonamental a la intimitat, recollit a [l'article 18](#) de la [Constitució Espanyola](#), el dret a la protecció de dades personals, contemplat a [l'article 31](#) de [l'Estatut d'Autonomia de Catalunya](#).

Així mateix, la pregunta està relacionada amb [l'article 10.3 de la Llei 14/1986, de 25 d'abril, General de Sanitat](#), [l'article 7 de la Llei 33/2011, de 4 d'octubre, General de Salut Pública](#), i [l'article 7 de la Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica](#), que versen sobre el dret a la privacitat de les persones pacients i la confidencialitat de les dades mèdiques.

Finalment, té relació amb el principi [Privacy del document A Human Rights-Based Approach To Data de l'ACNUDH](#), les recomanacions contingudes al [toolkit d'ICO sobre riscos i protecció de dades en la IA](#) i l'apartat [Privacy and Security de l'AI Ethics Assessment de GSMA](#).

5.2. Governança de dades

Nº Preguntes

5.2.1. **Heu establert un registre de la informació d'accés a les dades (personals o no) per identificar clarament quan, qui i amb quin propòsit hi ha accedit?**

Aquesta pregunta fa referència [als articles 10.5, i 50.3 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la governança de dades i els requisits dels sistemes de reconeixement d'emocions i categorització biomètrica.

També es relaciona amb [l'article 30 del Reglament 2016/679, de 27 d'abril 2016 \[Reglament general de protecció de dades\]](#), sobre el registre de les activitats de tractament.

Així mateix, la pregunta guarda relació amb [l'article 20 del Reglament 2022/868 del Parlament Europeu i del Consell, de 30 de maig del 2022, relatiu a la governança europea de dades \[Data Governance Act\]](#), sobre els requisits de transparència i el registre de les persones a les quals se les ha permès el tractament de les dades i [l'article 16.7 de la Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica](#).

Per l'altra banda, la pregunta es fonamenta en [l'article 10 del Conveni d'Oviedo](#), que versa sobre el dret a la vida privada i la informació, i en [l'article 11 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets](#)

[humans, democràcia i estat de dret](#), sobre el dret a la privacitat i la protecció de dades personals.

Finalment, la pregunta també es basa en les recomanacions contingudes a [l'apartat Privacy and Data Governance de l'ALTAI](#).

5.2.2. **Heu documentat el procés de recopilació i preparació de les dades per al vostre sistema d'IA?**

Aquesta pregunta fa referència [als articles 10, 17.1 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre la governança de dades i el sistema de gestió de qualitat dels sistemes d'alt risc i les obligacions dels proveïdors dels models de propòsit general.

Així mateix, la pregunta està relacionada amb la [Secció 1 del Capítol IV del Reglament 2016/679, de 27 d'abril 2016 \[Reglament general de protecció de dades\]](#), que regula les obligacions legals dels responsables del tractament de dades i el seu règim de responsabilitat.

Finalment, la pregunta també es basa en les recomanacions contingudes a [l'apartat Privacy and Data Governance de l'ALTAI](#).

5.2.3. **El conjunt de dades recopilades per entrenar el vostre sistema són adequades pel context en el qual serà utilitzat i són representatives de la població on s'utilitzarà el producte sanitari?**

La pregunta fa referència [als articles 10.4 i 17.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), sobre la governança de dades i el sistema de gestió de qualitat dels sistemes d'IA d'alt risc.

De la mateixa manera, la pregunta guarda vinculació amb [l'article 32 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#), [l'article 29 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#) relatius a la seguretat i el rendiment clínic dels dispositius mèdics. Així mateix, guarda relació amb [l'article 7 de la Directiva sobre responsabilitat pels danys causats per productes defectuosos](#), que versa sobre el caràcter defectuós dels productes.

La pregunta també es fonamenta amb [l'article 12 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el dret a la fiabilitat dels sistemes d'IA.

De l'altra banda, la pregunta està relacionada amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), sobre les garanties sanitàries dels productes.

Finalment, la pregunta també es basa en [el principi Participation del document A Human Rights-Based Approach To Data de l'ACNUDH](#), i en els articles [The ethics of AI in health care: A mapping review de Morleya, Machado, et al.](#), i [Generative AI and medical ethics: the state of play, de Zohny, Mann, et al.](#)

Principi 6: Autonomia

Aquest principi proposa garantir que els sistemes d'IA respectin l'autonomia de les persones, això és, que es prenguin les mesures necessàries per a garantir que les persones que actuïn com a usuàries clíniques del sistema d'IA puguin actuar d'acord amb la seva pròpia opinió professional, i que es respecti en tot cas el dret a l'autonomia de les persones pacients. Les preguntes d'aquest apartat provenen majoritàriament de les seccions relatives a agència humana i autonomia dels diferents instruments i recomanacions consultades.

6.1. Sistemes d'IA contestables i punt de vista del pacient

Nº	Preguntes
6.1.1.	Heu establert un procediment clar perquè qualsevol persona pacient pugui qüestionar el sistema d'IA, el seu ús previst, biaixos existents, els processos de validació i el rendiment esperat? Aquesta pregunta fa referència als articles 17.1 i 95.2, del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA], relatius al sistema de gestió de qualitat dels sistemes d'IA d'alt risc i a la promoció de codis de conducta d'aplicació voluntària en matèria de participació multisectorial. Així mateix, la pregunta es relaciona amb l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris i l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro, sobre els sistemes de gestió de qualitat dels productes sanitaris.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 7 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la dignitat humana i l'autonomia individual.

La pregunta també es relaciona amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics.

Finalment, la pregunta també es basa en les recomanacions contingudes a [Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability de Coeckelbergh](#), així com en [The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine, de Thomas Ploug I Søren Holm](#), sobre la contestabilitat dels sistemes d'IA per part dels pacients,

Exemple: Hem elaborat un formulari obert i un protocol mitjançant el qual la persona pacient pot obtenir més informació sobre l'ús de dades i el sistema d'IA, així com qüestionar estructuradament aspectes rellevants a través de preguntes i comentaris. Les persones responsables clínics s'encarreguen de recollir i processar les respostes.

6.1.2. **Heu establert un procediment clar perquè qualsevol persona pacient pugui qüestionar els aspectes concrets sobre l'ús de llurs dades?**

Aquesta pregunta es fonamenta en l'article [21 del Reglament 2016/679, de 27 d'abril 2016 \[Reglament general de protecció de dades\]](#), relatiu al dret d'oposició de les persones les dades de les quals són tractades.

Així mateix, està relacionada amb [l'article 10 del Conveni d'Oviedo](#), que versa sobre el dret a la vida privada i a la informació, i en [l'article 11 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la privacitat i la protecció de dades personals.

Finalment, la pregunta també es basa en les recomanacions contingudes a [Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability de Coeckelbergh](#), així com en [The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine, de Thomas Ploug I Søren Holm](#), sobre la contestabilitat dels sistemes d'IA per part dels pacients,

6.1.3. Heu establert un procediment clar perquè qualsevol persona pacient pugui qüestionar la divisió entre el sistema d'IA i la persona usuària clínica en la contribució a la presa de decisions clíniques?

Aquesta pregunta fa referència [als articles 17.1 i 95.2, del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius al sistema de gestió de qualitat dels sistemes d'IA d'alt risc i a la promoció de codis de conducta d'aplicació voluntària en matèria de participació multisectorial.

De la mateixa manera, la pregunta es fonamenta en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 7 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la dignitat humana i l'autonomia individual.

Finalment, la pregunta també es basa en les recomanacions contingudes a [Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability](#) de Coeckelbergh, així com en [The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine](#), de Thomas Ploug i Søren Holm, sobre la contestabilitat dels sistemes d'IA per part dels pacients,

6.1.4. En el desplegament del sistema d'IA heu tingut en compte la possible separació que existeix entre la persona usuària clínica i la persona pacient afectada per la decisió clínica i l'efecte que pot tenir en el seu ús?

Aquesta pregunta fa referència a [l'article 17.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu al sistema de gestió de qualitat dels sistemes d'IA.

També està relacionada amb [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), sobre els sistemes de gestió de qualitat dels productes sanitaris.

Així mateix, la pregunta es fonamenta en [els articles 9 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la responsabilitat i el retiment de comptes per els impactes adversos de la IA i el marc de gestió de riscos i impactes dels sistemes d'IA.

De la mateixa manera, la pregunta es relaciona amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels productes sanitaris.

Finalment, la pregunta també es basa en [The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine](#), de Thomas Ploug i Søren Holm i [Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability](#) de Coeckelbergh, sobre la contestabilitat dels sistemes d'IA per part dels pacients.

6.1.5. Heu tingut en compte en el desplegament del sistema d'IA la possible separació que existeix entre la persona usuària clínica i la persona pacient i l'efecte que pot tenir en l'anàlisi de beneficis, costos i riscos respecte al seu ús per a cadascuna de les persones involucrades?

Aquesta pregunta es fonamenta en [els articles 9 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versen sobre la responsabilitat i el retiment de comptes per els impactes adversos de la IA i el marc de gestió de riscos i impactes dels sistemes d'IA.

De la mateixa manera, la pregunta es basa en [The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine](#), de Thomas Ploug i Søren Holm i [Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability](#) de Coeckelbergh, sobre la contestabilitat dels sistemes d'IA per part dels pacients.

6.1.6. Heu tingut en compte en aquest anàlisi de beneficis i riscos la visió de la persona pacient, de la persona usuària mèdica, del centre associat i del sistema sanitari en general?

Aquesta pregunta fa referència [als articles 17.1 i 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatius al sistema de gestió de qualitat dels sistemes d'IA d'alt risc i a la promoció de codis de conducta d'aplicació voluntària en matèria de participació multisectorial.

També està relacionada amb [l'article 10.9 del Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) i [l'article 10.8 del Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#), que versen sobre els sistemes de gestió de qualitat dels productes sanitaris.

De la mateixa manera, la pregunta es fonamenta en [l'article 16 del Conveni](#)

[marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que versa sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Així mateix, la pregunta guarda relació amb [l'article 5 del Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#), relatiu a les garanties sanitàries dels dispositius mèdics.

Finalment, la pregunta també es basa en [The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine](#), de Thomas Ploug i Søren Holm i [Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability](#) de Coeckelbergh, sobre la contestabilitat dels sistemes d'IA per part dels pacients.

Exemple: En la fase de l'avaluació prèvia a l'adquisició d'un nou sistema d'IA (centrada a fer una anàlisi de costos, errors i conseqüències multidireccional) ha participat un comitè d'experts multidisciplinaris amb representants de les persones usuàries mèdiques, la gerència de l'hospital, així com representants del sistema sanitari català i de les persones pacients.

6.1.7. Heu establert un procés d'avaluació del vostre sistema d'IA per garantir que en cap cas disminueix el nombre d'eleccions possibles, ni tampoc l'obliga o força a prendre unes determinades decisions?

Aquesta pregunta fa referència a [l'article 13 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la transparència i comunicació d'informació als usuaris.

Per l'altra banda, la pregunta es fonamenta en [l'article 4 del Conveni d'Oviedo](#), que versa sobre els estàndards professionals, i en [l'article 7 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la dignitat humana i l'autonomia individual.

De la mateixa manera, la pregunta també està relacionada amb les recomanacions contingudes al [subapartat Human Agency and Autonomy de l'ALTAI](#), i [l'apartat Human Agency & Oversight del AI Ethics Assessment de GSMA](#).

6.1.8. Heu establert un protocol de comprovació que permeti certificar que el vostre sistema ajuda les persones a prendre millors decisions i més informades, d'acord amb els seus objectius i respectant la seva autonomia?

Aquesta pregunta fa referència a [l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), sobre els codis de conducta d'aplicació voluntària.

Així mateix, la pregunta està relacionada amb [els articles 2.3, 4 i 8 Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica](#), que versen sobre el dret a l'autonomia de les persones pacients.

Per l'altra banda, la pregunta es fonamenta en [l'article 5 del Conveni d'Oviedo](#), que versa sobre el consentiment lliure i informat, i en [l'article 7 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la dignitat humana i l'autonomia individual.

La pregunta també es basa en les recomanacions contingudes al [subapartat Human Agency and Autonomy de l'ALTAI](#) i l'article [The ethics of AI in health care: A mapping review](#) de Morleya, Machadoa, *et al.*

6.2. Precaució envers l'ús de sistemes d'IA generativa

Nº Preguntes

6.2.1. **Prèviament a realitzar una consulta de caràcter sanitari a un sistema d'IA generativa heu tingut en compte que existeix el risc que el resultat generat pel sistema sigui fruit d'una al·lucinació (resposta incorrecta)?**

Aquesta pregunta fa referència a [l'article 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a les obligacions dels proveïdors dels models de propòsit general.

De la mateixa manera, la pregunta es fonamenta en [l'article 4 del Conveni d'Oviedo](#), que versa sobre els estàndards professionals en les intervencions en l'àmbit de la salut, i en [l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), que sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Finalment, la pregunta també es basa en les recomanacions contingudes al [subapartat Human Agency and Autonomy de l'ALTAI](#), i en la *checklist* proposada a [Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist](#), de Ning, Teixayavongaixí, *et al.* Així com en

l'estudi [Artificial Intelligence in Medicine: Cross-Sectional Study Among Medical Students on Application, Education, and Ethical Aspects](#), de Weidener i Fischer, i la [Guia de Bones Pràctiques per al Desenvolupament d'Eines d'IA Generativa en Salut elaborada per la Fundació TIC Salut Social](#).

6.2.2. En cas de voler utilitzar, o haver utilitzat, un sistema d'IA generativa per a resoldre qüestions de caràcter sanitari, sou conscients que les recomanacions del sistema poden ser poc precises o poc adequades i que l'opinió d'una persona professional de la salut sempre serà més idònia?

Aquesta pregunta fa referència a [l'article 53.1](#) del [Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a les obligacions dels proveïdors dels models de propòsit general.

De la mateixa manera, la pregunta es fonamenta en [l'article 4](#) del [Conveni d'Oviedo](#), que versa sobre els estàndards professionals en les intervencions en l'àmbit de la salut, i en [l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

La pregunta també es basa en les recomanacions contingudes al [subapartat Human Agency and Autonomy de l'ALTAI](#), i en la [checklist](#) proposada a [Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist](#), de Ning, Teixayavongaixí, *et al.* Similarment, està relacionada amb l'estudi [Artificial Intelligence in Medicine: Cross-Sectional Study Among Medical Students on Application, Education, and Ethical Aspects](#), de Weidener i Fischer, i la [Guia de Bones Pràctiques per al Desenvolupament d'Eines d'IA Generativa en Salut elaborada per la Fundació TIC Salut Social](#).

6.2.3. En cas de voler utilitzar, o haver utilitzat, un sistema d'IA generativa per a assistir-vos en la realització d'una tasca o presa de decisió, sou conscients que heu de revisar que el resultat generat sigui correcte?

Aquesta pregunta fa referència a [als articles 50 i 53.1](#) del [Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a les obligacions dels proveïdors dels models de propòsit general.

De la mateixa manera, la pregunta es fonamenta en [l'article 4](#) del [Conveni d'Oviedo](#), que versa sobre els estàndards professionals en les intervencions en l'àmbit de la salut, i en [els articles 8 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i el control i el marc de gestió de riscos i impactes dels

sistemes d'IA.

La pregunta també es basa en les recomanacions contingudes al [subapartat Human Agency and Autonomy de l'ALTAI](#), i en la *checklist* proposada a [Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist](#), de Ning, Teixayavongaixí, *et al.* Finalment, està relacionada amb les publicacions [Artificial Intelligence in Medicine: Cross-Sectional Study Among Medical Students on Application, Education, and Ethical Aspects](#), de Weidener i Fischer, i [Generative AI and medical ethics: the state of play](#), de Zohny, Mann, *et al.*, així com amb la [Guia de Bones Pràctiques per al Desenvolupament d'Eines d'IA Generativa en Salut elaborada per la Fundació TIC Salut Social](#).

6.2.4. En cas de voler utilitzar, o haver utilitzat, un sistema d'IA generativa per a redactar un document, sou conscients que en el mateix document s'ha de fer constar que ha estat generat per un sistema d'IA?

Aquesta pregunta fa referència a [als articles 50 i 53.1 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a les obligacions dels proveïdors dels models de propòsit general.

De la mateixa manera, la pregunta es fonamenta en [l'article 4 del Conveni d'Oviedo](#), que versa sobre els estàndards professionals en les intervencions en l'àmbit de la salut, i en [els articles 8 i 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre la transparència i el control i el marc de gestió de riscos i impactes dels sistemes d'IA.

La pregunta també es basa en les recomanacions contingudes al [subapartat Human Agency and Autonomy de l'ALTAI](#), i en la *checklist* proposada a [Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist](#), de Ning, Teixayavongaixí, *et al.*

Finalment, està relacionada amb les publicacions [Artificial Intelligence in Medicine: Cross-Sectional Study Among Medical Students on Application, Education, and Ethical Aspects](#), de Weidener i Fischer, i [Generative AI and medical ethics: the state of play](#), de Zohny, Mann, *et al.*, així com amb la [Guia de Bones Pràctiques per al Desenvolupament d'Eines d'IA Generativa en Salut elaborada per la Fundació TIC Salut Social](#).

6.2.5. Dins de la vostra organització heu portat a terme activitats de formació i capacitat que permetin al personal clínic utilitzar eines d'IA generativa de forma conscient i segura?

Aquesta pregunta es fonamenta en [l'article 4 del Conveni d'Oviedo](#), que versa sobre els estàndards professionals en les intervencions en l'àmbit de la salut, i en [l'article 20 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre l'alfabetització en matèria d'IA i les habilitats digitals.

La pregunta també es basa en la [Guia de Bones Pràctiques per al Desenvolupament d'Eines d'IA Generativa en Salut elaborada per la Fundació TIC Salut Social](#).

Principi 7: Sostenibilitat

Aquest principi planteja que els productes sanitaris que incorporin la IA han de ser utilitzats de manera respectuosa per als drets de les persones, la societat i el medi ambient. Això inclou l'establiment de mecanismes que permetin als diferents actors socials participar en la definició dels paràmetres ètics per a l'ús de la IA, garantir els drets del personal laboral dels serveis sanitaris i prendre en consideració l'impacte mediambiental dels sistemes d'IA. Les preguntes d'aquest apartat provenen majoritàriament de les seccions relatives a l'impacte de la IA sobre els drets i el benestar ambiental dels diferents instruments i recomanacions consultades.

7.1. Benestar laboral i comunitari

Nº	Preguntes
7.1.1	Com a organisme públic heu tingut una participació activa en la discussió sobre principis ètics per a una IA de confiança? Aquesta pregunta fa referència a l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA], que versa sobre la promoció de codis de conducta d'aplicació voluntària en matèria de sostenibilitat ambiental. També es fonamenta en l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret, sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Així mateix, la pregunta es relaciona amb les guies establertes per la [WHO a l'apartat *Ensure inclusiveness and equity* del document *Ethics and governance of artificial intelligence for health*](#) sobre la participació activa en les qüestions ètiques que envolten a l'ús de la IA en l'entorn sanitari.

Finalment, la pregunta també esta relacionada amb la publicació [The ethics of AI in health care: A mapping review](#) de Morleya, Machadoa, et al.

Exemple: Hem establert protocols de comunicació regular amb el comitè ètic avaluador del nostre hospital així com, altres organismes públics dedicats a l'estudi dels principis ètics de la IA.

A més a més, participem activament en fòrums interdisciplinaris d'IA de confiança aplicada a l'àmbit de la salut.

7.1.2. **Considereu que el vostre sistema d'IA té un impacte positiu en els drets laborals i/o en la qualitat del treball de les persones treballadores del sector sanitari?**

Aquesta pregunta fa referència [als articles 27.1, 55.1 i l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre l'avaluació de l'impacte dels sistemes d'alt risc sobre els drets fonamentals, les obligacions dels proveïdors dels sistemes de propòsit general de risc sistèmic i la promoció de codis de conducta d'aplicació voluntària en matèria sostenibilitat.

També es fonamenta en [l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Finalment, la pregunta es basa en la publicació [Ethical Challenges and Solutions of Generative AI: An Interdisciplinary Perspective](#), d'Al-kfairy, Mustafa, et al i en les recomanacions generals de la guia [FUTURE-AI](#).

7.1.3. **Col·laboreu activament amb organitzacions laborals o sindicats per mitigar el possible desplaçament laboral causat per l'adopció de tecnologies d'IA a través de la requalificació de les treballadores afectades?**

Aquesta pregunta fa referència a [l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció de codis de conducta d'aplicació voluntària en matèria sostenibilitat.

També es fonamenta en [l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Així mateix, la pregunta es relaciona amb el [punt 50, i al títol IV, actuació 10 sobre economia i treball, de les recomanacions sobre ètica en la Intel·ligència Artificial d'UNESCO](#), relatius a l'avaluació ètica de l'impacte dels sistemes d'IA sobre els drets i el bon funcionament dels sistemes d'IA en matèria d'empresa i treball.

Finalment, la pregunta es basa en la publicació [Ethical Challenges and Solutions of Generative AI: An Interdisciplinary Perspective](#), d'Al-kfairy, Mustafa, *et al*, així com en les recomanacions generals de la guia [FUTURE-AI](#).

7.2. Impacte mediambiental

Nº Preguntes

7.2.1. Heu avaluat l'eficiència energètica del vostre sistema d'IA tenint en compte el seu cicle de vida?

Aquesta pregunta fa referència [als articles 53.1 i 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), que versen sobre les obligacions dels proveïdors dels sistemes de propòsit general i els Codis de conducta d'aplicació voluntària en matèria de sostenibilitat ambiental.

També es fonamenta en [l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Així mateix, la pregunta es relaciona amb les recomanacions contingudes a [l'apartat d'Environmental well-being de l'ALTAI](#) i amb la [Policy Area 5: Environment and Ecosystems](#) de les recomanacions sobre ètica en la Intel·ligència Artificial d'UNESCO. Coincideix també amb [l'apartat Enviromental Wellbeing del AI Ethics Assessment de GSMA](#) i amb les recomanacions generals de la guia [FUTURE-AI](#).

7.2.2. En el disseny i implementació de les vostres estratègies per a reduir la petjada ambiental heu tingut en compte la despesa energètica que comporten els sistemes d'IA o les heu pogut adaptar per incloure aquest aspecte?

Aquesta pregunta fa referència a [l'article 95.2 del Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#), relatiu a la promoció de codis de conducta d'aplicació voluntària en matèria de sostenibilitat ambiental.

Així mateix, la pregunta es fonamenta en [la Resolució 48/13 de 8 d'Octubre de 2021 del Consell de Drets Humans de l'ONU](#) relativa a un medi ambient saludable, i en [l'article 16 del Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#), sobre el marc de gestió de riscos i impactes dels sistemes d'IA.

Per l'altra banda, la pregunta troba relació amb les recomanacions i textos elaborats per la UNESCO en matèria mediambiental i IA com són [l'apartat *Environmental well-being* de l'ALTAI](#), que ho preveu al requisit sisè, i [el títol IV, actuació 5 sobre medi ambient i ecosistemes de les recomanacions sobre ètica en la intel·ligència artificial d'UNESCO](#), que recull un conjunt de recomanacions per tal de garantir un bon funcionament dels sistemes d'IA. Finalment, també es relaciona amb les recomanacions generals de la guia [FUTURE-AI](#).

8. GUIA PER CONTRACTES EN MATÈRIA D'IA EN L'ÀMBIT DE LA SALUT

Un dels grans reptes que planteja la intel·ligència artificial és l'atribució de responsabilitats per les conseqüències derivades de la implantació d'aquesta tecnologia. És per aquest motiu que el RIA compta, entre els articles 24 i 29, amb disposicions en les quals es delimiten quins subjectes són els responsables de complir amb les exigències contemplades per la mateixa llei, ja siguin les persones fabricants dels productes, les importadores, les distribuïdores, les usuàries o terceres persones.

Tanmateix, els riscos associats a la implementació dels sistemes d'IA poden derivar a la potencial infracció d'una gran diversitat de textos legals, més enllà del reglament de la IA i qualsevol altra normativa específica d'aquests sistemes, en funció de quin sigui l'àmbit específic d'aplicació i les característiques i gravetat del dany causat. Per tant, per a

determinar qui és el subjecte responsable d'un determinat resultat, caldrà consultar tots els instruments legals aplicables en un context concret.

En aquest sentit, també resulten rellevants de cara a la determinació de responsabilitats les normes sobre responsabilitat civil extracontractual, això és, la determinació de la responsabilitat sobre els danys o perjudicis derivats de l'ús de la intel·ligència artificial quan es produeix un supòsit que no estava previst en un contracte signat entre les parts. En territori de l'estat espanyol, aquestes regles es recullen a partir de l'article 1902 del Codi Civil. Malauradament, el debat entorn com es podrien aplicar aquestes directrius en relació amb els supòsits en els quals ha estat involucrada una aplicació d'intel·ligència artificial encara no han estat resolts. Això es deu al fet que la capacitat que tenen els sistemes d'IA per a automatitzar decisions i per modificar el seu comportament segons aprèn del seu entorn, dificulta molt determinar el grau de culpa de les persones involucrades en el seu desenvolupament, comercialització i implementació. Així mateix, existeix la problemàtica de les *black box*, aquells espais dels sistemes d'IA en els que no és possible saber la raó per la qual el sistema ha pres una decisió en comptes d'una altra.

Per aquest motiu, resulta rellevant comptar amb directrius específiques sobre la responsabilitat extracontractual pels potencials danys i perjudicis provocats per un sistema d'IA. La UE va fer un pas en aquesta direcció amb la Proposta de Directiva sobre l'adaptació de les normes de responsabilitat civil extracontractual a la intel·ligència artificial (també coneguda com a *AI Liability Directive*). Aquesta es tractava d'una proposta legislativa que determinava la presumpció de culpa dels danys i perjudicis causats per un sistema d'IA (i en conseqüència, la responsabilitat dels mateixos) dels actors subjectes a les obligacions del RAI, als quals també els imposava l'exigència d'aportar les proves sobre el sistema d'IA necessàries per esclarir el cas. No obstant, la proposta va ser retirada per la Comissió Europea el passat 11 de febrer de 2025, per la qual cosa finalment no serà adoptada.

Alternativament, resulta d'aplicació la Directiva (UE) 2024/2853 del Parlament Europeu i del Consell, de 23 d'octubre de 2024, sobre responsabilitat pels danys causats per productes defectuosos i per la qual es deroga la Directiva 85/374/CEE del Consell. Aquesta norma, tot i ser d'aplicació més genèrica, també dona cobertura als sistemes d'intel·ligència artificial. De fet, una de les motivacions de la seva adopció és, precisament, l'adaptació de la normativa europea sobre productes defectuosos a aquestes tecnologies. Així es declara al tercer considerant de la directiva, en anunciar que "*La Directiva 85/374/CEE ha estat un instrument eficaç i important, però s'hauria de revisar a la llum dels avenços relacionats amb les noves tecnologies, inclosa la intel·ligència artificial (IA)*"⁷³.

⁷³ Directive (EU) 2024/2853 of the European Parliament and of the Council of 23 October 2024 on liability for defective products and repealing Council Directive 85/374/EEC. Disponible a: <https://eur-lex.europa.eu/legal-content/EN/ALL/?uri=CELEX%3A32024L2853> (Consultada el 20 de febrer de 2025).

D'acord amb l'article 7 de la directiva, un producte serà considerat defectuós quan no compleixi amb el nivell de seguretat que li és exigít en virtut del dret de la Unió Europea o d'una norma nacional, i que les persones afectades pel seu ús tenen el dret d'esperar. Per tant, si una aplicació d'IA no compleix amb els requisits que li són imposats pel RIA, el sistema tindrà caràcter defectuós d'acord amb la Directiva sobre productes defectuosos. Com a resultat, les persones que pateixin algun perjudici causat pel sistema, tindran dret a exigir una indemnització, tal com es reconeix en l'article 5 de la mateixa norma.

El mencionat article 7 també exposa un llistat de factors a tenir en compte de cara a determinar el caràcter defectuós del producte (tot i que no és una llista tancada, de manera que es podran tenir en compte circumstàncies no contemplades de manera expressa). Un d'aquests elements és la capacitat que tingui el producte per continuar aprenent o adquirir noves característiques després de la seva introducció al mercat o posada en funcionament (article 7.c). Per la qual cosa, la responsabilitat pel possible impacte negatiu de la IA no es limita al moment en què és introduïda al mercat o posada en servei, sinó que s'estén al llarg de tot el seu cicle de vida, de la mateixa manera que el RIA imposa l'obligació de mantenir mesures que permetin garantir que el sistema continua complint amb els requisits legals en tot moment.

En virtut de l'article 8 de la Directiva, les persones que poden ser considerades com a responsables, en funció del cas concret i de les circumstàncies que han determinat la defectuositat del producte, són: el fabricant del producte defectuós; el fabricant d'un component defectuós, sempre que aquest hagi causat que el defecte en el producte i que continuï sota el control d'aquest fabricant; en cas que el fabricant d'un producte o component defectuós estigui establert fora del territori de la Unió Europea, sense perjudici de la seva pròpia responsabilitat, també seran considerats responsables l'importador del producte o component, el representant autoritzat del fabricant, i, en absència de les dues figures anteriors, el prestador de serveis logístics; qualsevol persona (física o jurídica) que modifiqui substancialment un producte per a la seva posterior comercialització o posada en servei; finalment, en cas de no poder-se identificar a cap dels actors ja mencionats, podrà ser considerat responsable el distribuïdor del producte o el proveïdor d'una plataforma en línia que permeti als consumidors contractar formalitzar contractes a distància amb els comerciants, sempre que es compleixin amb una sèrie de condicions.

Si bé les persones afectades seran les responsables de demostrar el caràcter defectuós del producte, l'existència del dany causat, i la relació causal entre la defectuositat del sistema i les lesions causades, és la part demandada, això és, la persona contra la qual s'imposa la demanda, la que ha d'aportar les proves relatives al producte (articles 9 i 10 de la directiva). Aquesta exigència es fonamenta en que les persones consumidores o afectades per un producte molts cops no tenen la capacitat per accedir a la informació

que pot demostrar el caràcter defectuós del producte i, en conseqüència, el seu dret a rebre una indemnització, ja que aquesta està en possessió de la part demandada. Aquesta circumstància és especialment important en el cas de la intel·ligència artificial, car que moltes de les persones perjudicades per un sistema no tenen els coneixements tècnics per entendre el seu funcionament, la qual cosa dificulta la seva capacitat per exigir responsabilitats.

Per acabar, l'article 10 estableix una sèrie de circumstàncies en les quals es presumirà el caràcter defectuós del producte, aquestes són: 1) que la part demandada no hagi aportat les proves, tal com li és exigít; 2) que la persona perjudicada hagi pogut demostrar que el producte no compleix amb els requisits legals que li són aplicables; i 3) que la persona perjudicada demostrï que el dany ha estat causat per un mal funcionament manifest del producte.

Un cop determinat el caràcter defectuós del producte, es reconeixerà la responsabilitat de la demandada i el dret de la persona perjudicada a rebre una indemnització si aquesta última pot demostrar que els danys patits han estat provocats per la disfunció del sistema. En qualsevol cas, el punt 4 de l'article 10 especifica que, en cas que sigui extremadament complicat per a la persona que interposa la demanda demostrar el caràcter defectuós del producte o la seva relació amb els danys patits, especialment a causa de la seva complexitat tècnica o científica, es presumirà tant la defectuositat del producte com la seva relació amb el resultat lesiu.

Com es pot observar, la Directiva sobre responsabilitat pels danys causats per productes defectuosos suposa un gran increment de la seguretat jurídica en relació amb la implementació de sistemes d'IA i a la protecció dels individus respecte als potencials impactes negatius que se'n derivin del seu ús. Tanmateix, no representa una solució definitiva. Per un costat, aquesta norma únicament serà aplicable als productes introduïts al mercat o posats en serveis després del 9 de desembre de 2026, deixant desprotegides a les persones que pateixin danys i perjudicis pels sistemes d'IA comercialitzats en un moment anterior a aquesta data. Per l'altre, fins que no comencin a sorgir les primes sentències relacionades amb l'aplicació de la directiva, resulta com complicat comprovar fins a quin punt les seves previsions donen una resposta adequada a la realitat, o si bé el seu articulat esdevé insuficient.

Davant d'aquesta incertesa, esdevé cabdal determinar el grau de responsabilitat de totes les persones implicades en un negoci jurídic relacionat amb el desenvolupament, comercialització i desplegament d'un sistema d'IA en el contracte que dona lloc a aquesta situació. D'aquesta manera s'afavoreix el rendiment de comptes i la seguretat jurídica.

És per aquesta raó que en aquest últim capítol hem afegit una guia, sobre la base dels set principis establerts en el qüestionari, en la que sintetitzem els requisits ètic-legals formulats de manera que reflecteixi quins aspectes haurien de ser considerats en la

redacció de contractes sobre sistemes d'IA. Recaurà sobre les parts contractants determinar quina d'elles (sinó totes), seran responsables de cadascun d'aquests aspectes. D'aquesta manera, aspirem a fer que la guia serveixi com a una eina que permeti a les persones proveïdores i usuàries saber quin és el contingut esperable en contractes envers aquesta matèria.

Per a crear aquest nou apartat hem pres de referència la guia per a contractes introduïda al Model PIO sobre Obligacions i Drets, la qual es basà en l'informe fet per *Society for Computers & Law (SCL) AI Group* sobre *Artificial Intelligence Contractual Clauses*⁷⁴, i l'hem adaptada i actualitzada per donar resposta a les especificacions relatives a l'àmbit de la salut i les normes aplicables a aquesta matèria.

Principi 1: Transparència i explicabilitat

1.1. Anàlisi preliminar de l'ús del sistema

Nº	Recomanació de clàusula
1.1.1	El sistema anirà acompanyat d'informació precisa sobre el seu funcionament, les seves limitacions, com arriba a un resultat determinat i com aquest contribueix a la presa d'una decisió clínica, així com la identitat i les dades de contacte de la persona proveïdora o representants.
1.1.2	Es garantirà que el procés de generació de resultats del sistema compleix amb els següents requisits: • Està ben documentat en un llenguatge comprensible i accessible per a les persones usuàries clíniques per a detectar futurs problemes. • Utilitza el model més simple i intel·ligible en favor de la transparència. • Pren en consideració les competències i alfabetització en matèria d'intel·ligència artificial de totes les persones usuàries. • Realitza registres automàtics sobre el seu funcionament. • És reproducible.

⁷⁴ Society for Computers & Law (SCL) AI Group – AI Clauses Projecte (2023). *Artificial Intelligence Contractual Clauses*. (Informe guia per a detallar la responsabilitat de proveïdors i clients en les relacions contractuals en el marc de sistemes d'IA). Disponible a: <https://www.scl.org/news/13004-society-for-computers-and-law-ai-group-launches-artificial-intelligence-contractual-clauses> (consultada el 17 d'octubre 2023).

1.2. Interpretabilitat del sistema i informació al pacient

Nº Recomanació de clàusula

- 1.2.1. El sistema d'IA anirà acompanyat dels materials addicionals que siguin necessaris perquè les persones usuàries clíniques puguin entendre el funcionament del sistema d'IA i explicar-ho a les persones pacients.

1.3. Explicabilitat del sistema

Nº Recomanació de clàusula

- 1.3.1 Es documentarà i es farà pública la següent informació relativa al sistema:
- El seu procés de desenvolupament i les persones que hi han estat implicades.
 - La informació tècnica necessària per al seu ús.
 - Instruccions concises, clares i escrites en un llenguatge accessible, sobre el seu funcionament.
 - El procés pel qual el sistema arriba a un resultat determinat i els resultats previstos.
 - Les possibles limitacions del sistema.
 - Informació relativa a l'intercanvi i cessió de dades així com les vies i mecanismes que permeten a les persones interessades exercir els seus drets.
- 1.3.2 El sistema comptarà amb un mecanisme que permeti identificar el contingut que genera com a creació d'aquest.
- 1.3.3. El sistema comptarà amb un mecanisme que permeti identificar el quines són les variables clíniques més rellevants que utilitza quan fa una predicció.
- 1.3.4. S'informarà l'organització sanitària i les persones usuàries clíniques sobre les mètriques que s'han utilitzat per optimitzar i validar el sistema d'IA així com la metodologia de validació per desenvolupar el sistema.
- 1.3.5. Es prendran mesures per tal que el nivell de transparència i explicabilitat del sistema no suposi un risc per a la privacitat de les persones ni la mateixa seguretat del sistema.

- 1.3.6. Es garantirà que s'ha validat en l'àmbit clínic, per part de professionals de l'àmbit en el qual està previst que es faci servir el sistema, que els resultats que dona estan ben explicats i són comprensibles, són adequats i ajuden en la tasca a realitzar.

1.4. Protocols de comunicació

Nº Recomanació de clàusula

- 1.4.1. El sistema disposarà de les següents eines de comunicació:
- Una funció que permeti a les persones usuàries clíniques i les pacients realitzar comentaris sobre el seu funcionament.
 - Un sistema de notificacions en cas d'incident greu o error.
- 1.4.2. S'establirà un procediment per comunicar de manera regular informació sobre en quines decisions clíniques les persones participants són assistides per un sistema d'IA.

Principi 2: Justícia i equitat

2.1. Inclusió i diversitat

Nº Recomanació de clàusula

- 2.1.1. En relació amb les dades utilitzades pel sistema, s'assegurarà el compliment dels següents requisits:
- El sistema anirà acompanyat de documentació relativa al tipus de dades utilitzades, la seva procedència, i els canvis que han patit en cada fase del seu cicle de vida.
 - Es prendran mesures per verificar la representativitat i adequació de les dades utilitzades amb relació a la població a la qual s'aplicarà.
 - El sistema comptarà amb salvaguardes per verificar l'edat de les persones.
- 2.1.2. S'establirà una política de participació de diferents grups d'interès (*stakeholders*) on s'inclogui a les persones usuàries clíniques i les pacients.

2.1.3. El sistema comptarà amb mesures i mecanismes per a detectar, corregir i mitigar possibles biaixos durant tot el seu cicle de vida.

Especialment, les mesures aniran adreçades als biaixos presents en:

- La fase de disseny o implementació.
- Les dades d'entrenament i posteriors.
- Categories canviants o inadequades.
- Biaixos cognitius (ex. biaix d'automatització i biaix de confirmació).

2.1.4. Abans de la seva comercialització/utilització el sistema serà sotmès a una avaluació sobre l'impacte dels seus possibles biaixos, especialment aquells que podrien afectar els drets, salut o seguretat de les persones.

2.1.5. El sistema complirà amb els estàndards de classificació que tracten els biaixos relacionats amb l'edat, l'ètnica, la raça, el sexe o la discapacitat d'acord amb la normativa aplicable.

2.2. Identificació i mitigació de biaixos

Nº Recomanació de clàusula

2.2.1 El sistema d'IA comptarà amb mesures d'inclusió i equitat, especialment en relació amb els següents aspectes:

- Testatge dels resultats del sistema per a diferents grups afectats.
- Mesures d'augment de l'accessibilitat.
- Mesures que recolzin el seu ús per part de persones en situació de vulnerabilitat.

Mesures per a garantir que el sistema s'adapta a un ampli espectre de preferències i habilitats individuals.

2.2.2. El grup de desenvolupament del sistema d'IA ha d'incloure persones de gènere no masculí, persones de diversos orígens o cultures i de diverses disciplines de coneixement (més enllà de la ciència de dades, l'enginyeria i les ciències de la salut).

Principi 3: Seguretat i no maleficència

3.1. Gestió del risc

Nº	Recomanació de clàusula
----	-------------------------

3.1.1 El sistema anirà acompanyat de la següent informació:

- El nivell de precisió i seguretat del sistema.
- La manera i el grau d'influència que les accions de les persones usuàries poden tenir sobre el rendiment del sistema.
- Els possibles comportaments i escenaris que mai hauria de transgredir per garantir la seguretat i preservar la no maleficència del sistema.
- Els mecanismes de supervisió i control del sistema (tant intern com extern).
- Guies d'utilització i bones pràctiques dels sistemes d'IA redactades en un llenguatge accessible a les persones usuàries clíniques.

3.1.2. El sistema comptarà amb mecanismes clars i definits perquè les persones usuàries clíniques puguin contactar amb els desenvolupadors del sistema d'IA quan sigui necessari.

3.1.3. El sistema comptarà amb mecanismes que garanteixin l'actualització periòdica de la informació disponible sobre el rendiment del sistema periòdicament la informació disponible sobre el rendiment del sistema d'IA.

3.2. Traçabilitat per a la gestió del risc

Nº	Recomanació de clàusula
----	-------------------------

3.2.1. Prèviament a la seva comercialització/utilització el sistema serà sotmès a una avaluació que incorpori els següents aspectes:

- Els possibles impactes del sistema sobre la salut, la seguretat i els drets de les persones afectades.
- El risc que les dades recollides, generades o utilitzades pel sistema siguin emprades per discriminar a persones de manera negativa.
- Els possibles riscos i danys que provocaria el sistema en cas de realitzar prediccions inexactes en els escenaris previstos.

- Els riscos que comportaria utilitzar el sistema més enllà dels usos previstos.
- Els possibles riscos derivats del fet que s'involucrin persones d'una especialitat clínica diferent de l'actora principal durant el context de desplegament de la solució.

3.2.2. El sistema comptarà amb les següents mesures:

- Mesures de resistència davant errors sorgits per la seva interacció amb altres sistemes o persones.
- Mesures de resistència i mitigació d'errors sorgits d'intents de manipulació per persones no autoritzats.
- Metodologies d'avaluació o autoavaluació periòdica reconeguda per vetllar per al bon funcionament del sistema i preveure possibles funcionaments erronis.

3.2.3. Es garantirà que prèviament a la seva comercialització/utilització el sistema haurà estat sotmès a les avaluacions que corroboren la seva conformitat respecte als requisits legals que li són aplicables.

3.2.4. El sistema comptarà amb mitjans que permetin les següents funcions:

- Identificar, avaluar i adoptar mesures respecte als riscos sorgits del seu ús i de les dades recollides després d'entrar en funcionament.
- Informar sobre els riscos residuals del sistema que no hagin pogut ser corregits ni mitigats.

Principi 4: Responsabilitat i retiment de comptes

4.1. Control

Nº	Recomanació de clàusula
----	-------------------------

4.1.1.	El sistema comptarà amb eines que permetin validar i contrastar clínicament els seus resultats.
--------	---

4.1.2.	Prèviament a la utilització/comercialització del sistema es determinarà:
--------	--

- | | |
|--|---|
| | <ul style="list-style-type: none"> • Quina serà la contribució del sistema en la presa de decisions clínica. |
|--|---|

- El grau de responsabilitat de les persones involucrades en el control, operació i vigilància humana del sistema.
- Les persones responsables d'actualitzar el sistema.

4.1.3. Es proporcionarà les persones involucrades en el control de les decisions preses pel sistema, en especial de les persones usuàries clíniques, formació adequada perquè tinguin una comprensió adequada sobre el seu funcionament i comptin amb la capacitat necessària per exercir aquesta funció.

4.1.4. El sistema comptarà amb les eines que permetin a les persones encarregades del seu control:

- Intervenir, modificar, eliminar o revertir les decisions preses pel sistema.
- Habilitar modificacions en el protocol clínic de processament de dades i ús del sistema.
- Limitar l'ús del sistema per grups de pacients o contextos clínics específics quan el rendiment sigui inacceptable.

4.1.5. El sistema anirà acompanyat d'informació i documentació relativa als següents aspectes:

- Els límits dels mecanismes de control del sistema.
- Les mesures de supervisió humana amb les que compta el sistema.

4.1.6. Es garantirà que en el desenvolupament del sistema han participat persones expertes clíniques en el context de la solució, de manera que la desenvolupadora ha tingut coneixement d'aquest.

4.2. Retiment de comptes

Nº	Recomanació de clàusula
----	-------------------------

4.2.1. Prèviament a la utilització/comercialització del sistema es determinarà:

- La persona legalment responsable de les accions i/o decisions que pren durant tot el seu cicle de vida, particularment, quin és el grau de responsabilitat atribuït a la persona usuària clínica, l'organització sanitària a la qual pertany, i la creadora del sistema.

- Els mecanismes de revisió periòdica sobre responsabilitat i retiment de comptes amb els diferents agents involucrats en el desenvolupament i l'ús del sistema d'IA.

4.3. Figures de protecció i mesures de reparació

Nº Recomanació de clàusula

4.3.1. Prèviament a la comercialització/ús del sistema es prendran les següents mesures:

- Facilitar la demostració del compliment de la legislació vigent.
- Determinar formes de reparació o compensació en cas de perjudicis causats pel sistema, les quals inclouran la responsabilitat de les persones que han creat al sistema quan sigui pertinent.
- Determinar figures de protecció envers les persones usuàries i les persones afectades pel sistema.
- Determinar si el sistema d'IA pot ser considerat com a producte sanitari d'acord amb la regulació sobre productes sanitaris (MDR) o de diagnòstic in vitro (IVDR) i quina seria la seva classificació d'acord amb aquestes normes.

Principi 5: Privacitat

5.1. Protecció de la privacitat

Nº Recomanació de clàusula

5.1.1. El sistema comptarà amb mecanismes de protecció de dades presents durant tot el seu cicle de vida encaminats a garantir els següents aspectes:

- Que les dades de caràcter personal siguin tractades d'acord amb el grau de protecció legalment requerit per tal d'evitar-ne el mal ús, l'alternació, la pèrdua, o la sostracció per part de tercers.
- Que les dades subministrades per la persona interessada no siguin venudes, llogades, ni cedides sense el seu consentiment.

- Que el sistema compta amb processos per a mantenir o millorar la privacitat i protecció de les dades durant tot el seu cicle de vida.
- Que es fan servir tècniques d'anonimat, xifratge i agregació.
- Que existeix un protocol per a obtenir el consentiment de les persones afectades abans d'iniciar qualsevol processament de dades a través del sistema, i la informació relativa a aquestes accions és comprensible i accessible per a totes les possibles persones usuàries, inclosos aquelles amb necessitats especials.
- Que existeix un mecanisme a través del qual la persona interessada pot sol·licitar que la seva informació sigui eliminada.

5.2. Governança de dades

Nº Recomanació de clàusula

- 5.2.1. El sistema comptarà amb un registre de la informació d'accés a les dades que permeti identificar clarament qui, quan i amb quin propòsit hi accedeix.
- 5.2.2. Es prendran mesures per garantir que les persones responsables del processament i gestió de dades disposin de les competències necessàries per a realitzar aquestes funcions i coneixements actualitzats sobre la matèria.
- 5.2.3. De cara a garantir l'adequació i l'ús reposable de les dades utilitzades pel sistema es prendran les següents mesures:
- S'establirà un procediment per assegurar que els conjunts de dades usats siguin adients i rellevants pel rendiment i resultat previst del sistema i que siguin adequades pel context en el qual s'utilitzarà el sistema.
 - El procés de recopilació i preparació de les dades usades serà documentat.
 - El sistema utilitzarà mesures de privadesa (com per exemple la pseudonimització de les dades) i limitarà la reutilització de les dades personals sensibles.
- 5.2.4 El sistema està/ha estat dissenyat de manera es que pugui garantir que s'adhereix als principis de minimització de dades i protecció de dades des del disseny i per defecte.

Principi 6: Autonomia

6.1. Ús de dades i sistemes d'IA contestables i punt de vista del pacient

Nº Recomanació de clàusula

- 6.1.1. El sistema anirà acompanyat d'un procés d'avaluació per a garantir que no disminueix el nombre d'eleccions possibles i no obliga o força a prendre unes determinades decisions en general ni en cap fase del seu cicle de vida.
- 6.1.2. El sistema anirà acompanyat d'un procediment que permeti a les persones pacients qüestionar:
- Els resultats i l'ús previst del sistema.
 - Els aspectes concrets sobre l'ús de llurs dades.
 - Els biaixos existents, els processos de validació i el rendiment esperat del sistema d'IA.
 - La divisió en la contribució a la presa de decisions entre el sistema d'IA i la persona usuària clínica.
- 6.1.3. El sistema anirà acompanyat d'un protocol de comprovació que permeti certificar que ajuda les persones a prendre decisions millors i més informades.
- 6.1.4. El sistema comptarà amb mecanismes que garanteixin els següents aspectes:
- Que les persones usuàries clíniques siguin conscients que interactuen amb una intel·ligència artificial i estan prou capacitades i informades per a fer-ho.
 - Que s'informa les persones pacients sobre la seva exposició al sistema d'IA i com aquest funciona.
 - Que no posa en risc els nivells d'autonomia de les persones usuàries clíniques ni de les pacients i s'adapta a diferents preferències i necessitats.
- 6.1.5. El sistema comptarà amb els següents mecanismes:
- Mesures de supervisió humana que possibilitin gestionar el seu grau d'autonomia i adaptabilitat.
 - Mecanisme de limitació del grau d'autonomia del sistema per a evitar comportaments no desitjats o d'alt risc.
 - Adaptar-se i aprendre del seu entorn sense comprometre la seguretat ni estàndards ètics.

Principi 7: Sostenibilitat

7.1. Benestar laboral i comunitari

Nº	Recomanació de clàusula
----	-------------------------

7.1.1.	Es garantirà que en el disseny, la implementació i l'ús el sistema s'han pres en consideració les limitacions i obstacles derivats de la bretxa digital i l'accés social no igualitari.
--------	---

7.1.2.	Abans de la seva comercialització i/o utilització, el sistema serà sotmès a una avaluació sobre els seus impactes als drets de les persones, la seguretat, la salut i el medi ambient.
--------	--

Aquesta avaluació pren especialment en consideració els següents àmbits:

- Impacte sobre els drets laborals del personal de l'organització sanitària.
- Impacte i riscos relatius al desenvolupament sostenible.
- Mesures de compensació justa i tracte ètic a les persones implicades en totes les fases del cicle de vida del sistema i dels afectats per la seva implementació i ús.

7.2. Impacte mediambiental

Nº	Recomanació de clàusula
----	-------------------------

7.2.1.	El sistema comptarà amb mecanismes que li permetin enregistrar el seu consum d'energia i impacte ambiental durant tot el seu cicle de vida.
--------	---

7.2.2.	Abans de la seva comercialització i/o utilització, el sistema serà sotmès a una avaluació en la qual s'adreci, com a mínim, els següents aspectes:
--------	--

- L'eficiència energètica del sistema.
- Explorar l'ús de fonts d'energia renovables i pràctiques d'informàtica verda.

9. BIBLIOGRAFIA:

Ahmad, M. [Muhammad] A. [Aurangzeb], Yaramis, I. [Ilker], i Roy, T. [Taposh] D. [Dutta] (2023). Creating trustworthy llms: Dealing with hallucinations in healthcare ai. *arXiv Preprint arXiv:2311.01463*. Disponible a: <https://arxiv.org/abs/2311.01463>

AI Act, Comissió Europea (2021), Proposal for a Regulation of the European Parliament and of the Council laying down harmonised rules on Artificial Intelligence (Artificial Intelligence Act) and amending certain Union legislative acts. Disponible a: https://eur-lex.europa.eu/resource.html?uri=cellar:e0649735-a372-11eb-9585-01aa75ed71a1.0001.02/DOC_1&format=PDF

AI EI Group (2020) AI Ethics Impact Group: From Principles to Practice – An interdisciplinary framework to operationalise AI ethics. Disponible a: <https://www.ai-ethics-impact.org/en>

AI HLEG (2020) The Assessment List for Trustworthy Artificial Intelligence (ALTAI) for self-assessment, European Commission, Brussels. Disponible a: <https://futurium.ec.europa.eu/en/european-ai-alliance/pages/altai-assessment-list-trustworthy-artificial-intelligence>

Al-kfairy, M. [Mousa], Mustafa, D. [Dheya], Kshetri, N. [Nir], Insiew, M. [Mazen], i Alfandi, O. [Omar] (2024) Ethical challenges and solutions of generative AI: An interdisciplinary perspective. In *Informatics* (Vol. 11, No. 3, p. 58). MDPI. Disponible a: <https://www.mdpi.com/2227-9709/11/3/58>

Amann, J., Blasimme, A., Vayena, E., Frey, D., Madai, V. I., & Precise4Q consortium (2020). Explainability for artificial intelligence in healthcare: a multidisciplinary perspective. *BMC medical informatics and decision making*, 20(1), 310. Disponible a: <https://doi.org/10.1186/s12911-020-01332-6>

Aniek F. Markus, Jan A. Kors, Peter R. Rijnbeek (2021), The role of explainability in creating trustworthy artificial intelligence for health care: A comprehensive survey of the terminology, design choices, and evaluation strategies, *Journal of Biomedical Informatics*,

Volume 113, 2021, 103655, ISSN 1532-0464. Disponible a: <https://doi.org/10.1016/j.jbi.2020.103655>.

Assemblea General de les Nacions Unides (1948) Declaració Universal dels Drets Humans. Disponible a: <https://www.un.org/en/about-us/universal-declaration-of-human-rights>

Assemblea General de les Nacions Unides (1966) Pacte Internacional de Drets Civils i Polítics. Disponible a: https://treaties.un.org/doc/treaties/1976/03/19760323%2006-7%20am/ch_iv_04.pdf

Aussó, S. [Susanna], Berenguer, A. [Antoni], Aznar, J. [Júlia], Raventós, C. [Carolina], Gómez, V. [Vaneza], i Bretones, M. [Maria] (2025). Guia de Bones Pràctiques per al Desenvolupament d'Eines d'IA Generativa en Salut Grans Models de Llenguatge (LLM), Fundació TIC Salut Social. Disponible a: <https://iasalut.cat/wp-content/uploads/2025/03/Bones-Practiques-IAGen-LLM-Guia-CAT.pdf>

Aussó, S. [Susanna], Berenguer, A. [Antoni], Aznar, J. [Júlia], Raventós, C. [Carolina], Gómez, V. [Vaneza], i Bretones, M. [Maria] (2025). Guia per a l'Aplicació del Reglament Europeu d'IA en Salut. Fundació TIC Salut Social. Disponible a: <https://iasalut.cat/wp-content/uploads/2025/03/AI-Act-Guia-CAT.pdf>

Aussó, S. [Susanna], Berenguer, A. [Antoni], Aznar, J. [Júlia], Raventós, C. [Carolina], Gómez, V. [Vaneza], i Bretones, M. [Maria] (2025). Guia per a la qualificació i classificació d'una aplicació informàtica basada en IA com a producte sanitari segons els Reglaments Europeus MDR / IVDR, Fundació TIC Salut Social. Disponible a: <https://iasalut.cat/wp-content/uploads/2025/03/Qualificacio-Classificacio-MDSW-Guia-CAT.pdf>

Aussó, S. [Susanna], Berenguer, A. [Antoni], Aznar, J. [Júlia], Raventós, C. [Carolina], Gómez, V. [Vaneza], i Bretones, M. [Maria] (2025). Guia per a l'obtenció del marcatge CE d'una aplicació informàtica basada en IA com a producte sanitari segons els Reglaments Europeus MDR / IVDR, Fundació TIC Salut Social. Disponible a: <https://iasalut.cat/wp-content/uploads/2025/03/CE-Marcatge-MDSW-Guia-CAT.pdf>

Bærøe, K. [Kristine], Miyata-Sturm, A. [Ainar], & Henden, E. [Edmund] (2020). How to achieve trustworthy artificial intelligence for health. *Bulletin of the World Health Organization*, 98(4), 257–262. Disponible a: <https://doi.org/10.2471/BLT.19.237289>

Blease, C. [Charlotte] (2024) Open AI meets open notes: surveillance capitalism, patient privacy and online record access. *Journal of Medical Ethics*, 50(2), 84-89. Disponible a: <https://jme.bmj.com/content/50/2/84>.

Bharati, S., Mondal, M. R. H., & Podder, P. (2023). A Review on Explainable Artificial Intelligence for Healthcare: Why, How, and When?. *IEEE Transactions on Artificial Intelligence*. Disponible a: <https://arxiv.org/abs/2304.04780>

Business for Social Responsibility (2013) AI and Human Rights in Healthcare. Disponible a: <https://www.bsr.org/reports/BSR-AI-Human-Rights-Healthcare.pdf>

Catalonia.AI (2020), Estratègia d'Intel·ligència Artificial de Catalunya, Generalitat de Catalunya. Disponible a: <https://politiquesdigitals.gencat.cat/ca/economia/catalonia-ai/>

Chung, J. [Jamison] (2021) How Will Health Care Regulators Address Artificial Intelligence?, The Regulatory Review. Disponible a: <https://www.theregreview.org/2021/10/18/chung-how-will-health-care-regulators-address-artificial-intelligence/>

Coalition for Health AI (2022), Blueprint for Trustworthy AI Implementation Guidance and Assurance for Healthcare. Disponible a: <https://www.coalitionforhealthai.org/papers/Blueprint%20for%20Trustworthy%20AI.pdf>

Coeckelbergh M. Coeckelbergh, M. [Mark] (2020). Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability. *Science and engineering ethics*, 26(4), 2051–2068. Disponible a: <https://doi.org/10.1007/s11948-019-00146-8>

Comissió Europea (2019), Ethics guidelines for trustworthy AI. Disponible a: <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>

Constitució Espanyola de 1978, de 29 Diciembre 1978, BOE núm.311, de 29 de desembre de 1978.

Dewitte, P. [Pierre] (2025). AI Meets the GDPR: Navigating the Impact of Data Protection on AI Systems. In N. A. Smuha (Ed.), *The Cambridge Handbook of the Law, Ethics and Policy of Artificial Intelligence*, 133–157. Disponible a: [AI Meets the GDPR \(Chapter 7\) - The Cambridge Handbook of the Law, Ethics and Policy of Artificial Intelligence](#).

Directive (EU) 2024/2853 of the European Parliament and of the Council of 23 October 2024 on liability for defective products and repealing Council Directive 85/374/EEC. Disponible a: <https://eur-lex.europa.eu/legal-content/EN/ALL/?uri=CELEX%3A32024L2853>.

Farah L. [Line], Murris J. [Juliette] M., *et al.* (2023), Assessment of Performance, Interpretability, and Explainability in Artificial Intelligence–Based Health Technologies: What Healthcare Stakeholders Need to Know, Mayo Clinic Proceedings: Digital Health, Volume 1, Issue 2, 2023, Pages 120-138, ISSN 2949-7612. Disponible a: <https://doi.org/10.1016/j.mcpdig.2023.02.004>

Florien S van Royen *et al.* (2023), Five critical quality criteria for artificial intelligence-based prediction models, European Heart Journal, Volume 44, Issue 46, 7 December 2023, Pages 4831–4834. Disponible a: <https://doi.org/10.1093/eurheartj/ehad727>

Fundació TIC Salut Social (2022), Guia sobre l'explicabilitat en la Intel·ligència Artificial. Disponible a: <https://ticsalutsocial.cat/wp-content/uploads/2023/01/230226-Informe-Explicabilitat.pdf>

GDPR (2016) EU General Data Protection Regulation (GDPR): Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation), OJ 2016 L 119/1. Disponible a: <https://gdpr-info.eu/>

Govern of Australia (2019), The Artificial Intelligence (AI) Ethics Framework. Disponible a: <https://www.industry.gov.au/publications/australias-artificial-intelligence-ethics-framework>

GSMA (2022), The AI Ethics Playbook & Self-Assessment Questionnaire. Disponible a: <https://www.gsma.com/betterfuture/resources/ethicsplaybook>

Harrer S. (2023). Attention is not all you need: the complicated case of ethically using large language models in healthcare and medicine. *EBioMedicine*, 90, 104512. Disponible a: <https://doi.org/10.1016/j.ebiom.2023.104512>

ICO (2023), AI and data protection risk toolkit. Disponible a: <https://ico.org.uk/for-organisations/uk-gdpr-guidance-and-resources/artificial-intelligence/guidance-on-ai-and-data-protection/ai-and-data-protection-risk-toolkit/>

ISO 13485 (2016) Medical devices - Quality management systems - Requirements for regulatory purposes. Disponible a: <https://www.iso.org/standard/59752.html>

Karimian, G. [Golnar], Petelos, E. [Elena], i Evers, S. [Silvia] M.A.A. (2022) The ethical issues of the application of artificial intelligence in healthcare: a systematic scoping review. *AI Ethics* 2, 539–551 (2022) . Disponible a: <https://doi.org/10.1007/s43681-021-00131-7>

Kim, Y. [Yubin], Jeong, H. [Hyewon], *et al* (2025) Medical Hallucination in Foundation Models and Their Impact on Healthcare, *MedRxiv*, 2025-02. Disponible a: <https://www.medrxiv.org/content/10.1101/2025.02.28.25323115v1>

Kiseleva, A., Kotzinos, D., & De Hert, P. (2022). Transparency of AI in Healthcare as a Multilayered System of Accountabilities: Between Legal Requirements and Technical Limitations. *Frontiers in artificial intelligence*, 5, 879603. Disponible a: <https://doi.org/10.3389/frai.2022.879603>

Lee, H. P. [Hao-Ping] H. [Hank], Sarkar, A. [Advait], Tankelevitch, L. [Lev], Drosos, I. [Ian], Rintel, S. [Sean], Banks, R. [Richard], i Wilson, N. [Nicholas] (2025). The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and

Confidence Effects From a Survey of Knowledge Workers. *CHI Conference on Human Factors in Computing Systems (CHI '25)*. Disponible a: https://hankhplee.com/papers/genai_critical_thinking.pdf.

Lareo X. [Xabier] (2023). Large Language Models (LLM), a Supervisor europeu de protecció de dades (Ed.), TechSonar Report 2023-2024, 6-9. Disponible a: https://www.edps.europa.eu/data-protection/our-work/publications/reports/2023-12-04-techsonar-report-2023-2024_en

Lekadir, K. [Karim], Frangi, A. [Alejandro] F., Porras, A. [Antonio] R., et al. (2025). FUTURE-AI: international consensus guideline for trustworthy and deployable artificial intelligence in healthcare, *BMJ* 2025;388:e081554. Disponible a: <https://www.bmj.com/content/388/bmj-2024-081554>.

Llei 14/1986, de 25 d'abril, General de Sanitat, BOE núm. 102, de 29 d'abril de 1986.

Llei 33/2011, de 4 d'octubre, General de Salut Pública, BOE núm. 240, de 5 d'octubre de 2011.

Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica, BOE núm 274, de 15 de novembre de 2002.

Llei Orgànica 6/2006, de 19 de juliol, de reforma de l'Estatut d'Autonomia de Catalunya, BOE núm. 172, de 20 de juliol de 2006.

Liu, X., Cruz Rivera, S., Moher, D. et al. (2020) Reporting guidelines for clinical trial reports for interventions involving artificial intelligence: the CONSORT-AI extension. *Nat Med* 26, 1364–1374. Disponible a: <https://doi.org/10.1038/s41591-020-1034-x>

Liu, X. [Xiaoxuan], Glocker, B. [Ben], McCradden, M. [Melissa] M., Ghassemi, M. [Marzyeh], Denniston, A. [Alastair] K., i Oakden-Rayner, L. [Lauren] (2022). The medical algorithmic audit. *The Lancet. Digital health*, 4(5), e384–e397. Disponible a: [https://doi.org/10.1016/S2589-7500\(22\)00003-6](https://doi.org/10.1016/S2589-7500(22)00003-6)

Luo, Z. [Zining], Qiao, Y. [Yang], *et al.* (2025). Cross sectional pilot study on clinical review generation using large language models. *pj Digital Medicil*, 8, 170. Disponible a: <https://www.nature.com/articles/s41746-025-01535-z#citeas>

McLennan, S. [Stuart], Fiske, A. [Amelia], Tigard, D. [Daniel], Müller, R. [Ruth], Haddadin, S. [Sami], i Buyx, A. [Alena] (2022) Embedded ethics: a proposal for integrating ethics into the development of medical AI. *BMC Medical Ethics*, 23(1), 6. Disponible a: <https://bmcmethics.biomedcentral.com/articles/10.1186/s12910-022-00746-3>

MedTech Europe (2019). Trustworthy Artificial Intelligence (AI) in healthcare. Disponible a: https://www.medtecheurope.org/wp-content/uploads/2019/11/MTE_Nov19_Trustworthy-Artificial-Intelligence-in-healthcare-1.pdf

Mittelstadt, B. [Brent] (2021) The impact of artificial intelligence on the doctor-patient relationship, Consell d'Europa. Disponible a: <https://www.coe.int/en/web/bioethics/report-impact-of-ai-on-the-doctor-patient-relationship>

Mondal, H. [Himel] (2025). Ethical engagement with artificial intelligence in medical education. *Advances in Physiology Education*, 49, 163–165. Disponible a: <https://journals.physiology.org/doi/full/10.1152/advan.00188.2024>

Morley, J. [Jessica], Machado, C. [Caio] C. V., Burr, C. [Christopher], Cows, J. [Josh], Joshi, I. [Indra], Taddeo, M. [Mariarosaria], i Floridi, L. [Luciano] (2020) The ethics of AI in health care: a mapping review. *Social Science & Medicine*, 260, 113172. Dispibile a: <https://pubmed.ncbi.nlm.nih.gov/32702587/>

Ning, Y. [Yilin], Teixayavong, S. [Salinelat], *et al.* (2024) Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist. *The Lancet Digital Health*, 6(11), e848-e856. Disponible a: [https://www.thelancet.com/journals/landig/article/PIIS2589-7500\(24\)00143-2/fulltext](https://www.thelancet.com/journals/landig/article/PIIS2589-7500(24)00143-2/fulltext)

Ng, A. [Alfred] (2022) A uniquely dangerous tool': How Google's data can help states track abortions. Politico. Disponible a: <https://www.politico.com/news/2022/07/18/google-data-states-track-abortions-00045906>

OEIAC (2022) El Model PIO (Principis, Indicadors i Observables): Una proposta d'autoavaluació organitzativa sobre l'ús ètic de dades i sistemes d'intel·ligència artificial. Revisió 2024: Model PIO. Disponible a: <https://oeiac.cat/>

Oviedo Convention (1997), Convention on Human Rights and Biomedicine. Disponible a: <https://www.coe.int/en/web/bioethics/oviedo-convention>

Ploug, T. [Thomas]; Holm, S. [Søren] (2020), The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine, Volume 107, 2020, 101901, ISSN 0933-3657. Disponible a: <https://doi.org/10.1016/j.artmed.2020.101901>

Quinn, T. [Thomas] P., Sendadeera, M. [Manisha], *et al.* (2021) Trust and medical AI: the challenges we face and the expertise needed to overcome them, Journal of the American Medical Informatics Association, 28(4): 890–894. Disponible a: <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC7973477/#>

Reial Decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics, BOE núm 307 de 24 de desembre de 2015.

Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris, BOE núm. 69, de 22 de març de 2023.

Reglament (UE) 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017, sobre els productes sanitaris, pel qual es modifiquen la Directiva 2001/83/CE, el Reglament (CE) nº 178/2002 i el Reglament (CE) nº 1223/2009 i pel qual es deroguen les Directives 90/385/CEE i 93/42/CEE del Consell, DOUE núm. 117 de 5 de maig de 2017.

Reglament (UE) 2017/746 del Parlament Europeu i del consell de 5 d'abril de 2017 sobre els productes sanitaris per a diagnòstic in vitro i pel qual es deroguen la Directiva 98/79/CE i la Decisió 2010/227/UE de la Comissió, DOUE núm. 117, de 5 de maig de 2017

Reglament (UE) 2024/1689 del Parlament i del Consell, de 13 de juny de 2024, per la qual s'estableixen normes harmonitzades sobre intel·ligència artificial i es modifiquen els Reglaments (CE) n° 300/2008, (UE) n° 167/2013, (UE) n° 168/2013, (UE) 2018/858, (UE) 2018/1139 i (UE) 2019/2144 i les Directives 2014/90/UE, (UE) 2016/797 i (UE) 2020/1828 (RIA).

Procter, R. [Rob], Tolmie, P. [Peter], i Rouncefield, M. [Mark] (2023). Holding AI to Account: Challenges for the Delivery of Trustworthy AI in Healthcare. *ACM Trans. Comput.-Hum. Interact.* 30, 2, Article 31, 34. Disponible a: <https://doi.org/10.1145/3577009>

Saeed S. [Sy] A. [Atezaz] i Masters, R. [Ross] M. [MacRae] (2021), Disparities in Health Care and Digital Divide, *Current Psychiatry Reports* 23(9): 61. Disponible a : <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC8300069/>

Selecció de crides (2023), Calls for Papers on trustworthy AI for Health. Disponible a: <https://www.frontiersin.org/research-topics/51826/trustworthy-ai-for-healthcare>
https://www.mdpi.com/journal/sensors/special_issues/Healthcare_Human_Machine
<https://www.embs.org/jbhi/special-issues/trustworthy-and-collaborative-ai-for-personalised-healthcare-through-edge-of-things/>

Sepas, A., Bangash, A. H., Alraoui, O., El Emam, K., & El-Hussuna, A. (2022). Algorithms to anonymize structured medical and healthcare data: A systematic review. *Frontiers in bioinformatics*, 2, 984807. Disponible a: <https://doi.org/10.3389/fbinf.2022.984807>

UNESCO (2021) Draft text of the recommendation on the ethics of Artificial Intelligence. SHS/IGM-AIETHICS/2021/JUN/3Rev.2. Disponible a: https://unesdoc.unesco.org//in/documentViewer.xhtml?v=2.1.196&id=p::usmarcdef_000377897&file=/in/rest/annotationSVC/DownloadWatermarkedAttachment/attach_i mport_4bc33956-d665-48ad-92af-57ae532c98ef%3F_%3D377897eng.pdf&locale=es&multi=true&ark=/ark:/48223/pf000377897/PDF/377897eng.pdf

Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). Attention is all you need. *Advances in neural information processing systems*, 30, <https://doi.org/10.48550/arXiv.1706.03762>

Vetter, D. [Dennis], Amann, J. [Julia], Bruneault, F. [Frédéric] *et al.* (2023) Lessons Learned from Assessing Trustworthy AI in Practice. *DISO* 2, 35. <https://doi.org/10.1007/s44206-023-00063-1>

Vickers, A. [Andrew] J., i Elkin, E. [Elena] B. (2006). Decision curve analysis: a novel method for evaluating prediction models. *Medical decision making : an international journal of the Society for Medical Decision Making*, 26(6), 565–574. Disponible a: <https://doi.org/10.1177/0272989X06295361>

Weidener, L. [Lukas], i Fischer, M. [Michael] (2024) Artificial intelligence in medicine: cross-sectional study among medical students on application, education, and ethical aspects. *JMIR Medical Education*, 10(1), e51247. Disponible a: <https://mededu.jmir.org/2024/1/e51247/>

WHO (2021), Ethics and governance of artificial intelligence for health. Disponible a: <https://www.who.int/publications/i/item/9789240029200>

WHO (2023) WHO outlines considerations for regulation of artificial intelligence for health. Disponible a: <https://www.who.int/news/item/19-10-2023-who-outlines-considerations-for-regulation-of-artificial-intelligence-for-health>

Zicari, R. [Roberto] V., Brusseau, J. [James], *et al.* (2021). On Assessing Trustworthy AI in Healthcare. Machine Learning as a Supportive Tool to Recognize Cardiac Arrest in Emergency Calls. *Frontiers in Human Dynamics*, 3. Disponible a: <https://doi.org/10.3389/fhumd.2021.673104>

Zohny, H. [Hazem], Mann, S. [Sebastian] P. [Porsdam], Earp, B. [Brian] D., i McMillan, J. [John] (2024) Generative AI and medical ethics: the state of play. *Journal of medical ethics*, 50(2), 75-76. Disponible a: [Generative AI and medical ethics: the state of play, de Zohny, Mann, et al](#)

ANNEX I: Resum de les preguntes del qüestionari

Principi 1: Transparència i explicabilitat

1.1. Anàlisi preliminar de l'ús del sistema

Nº	Preguntes
1.1.1.	Sou capaços d'entendre com el sistema d'IA arriba a un resultat determinat i com aquest contribueix en la decisió clínica?
1.1.2.	Sou coneixedors de les possibles limitacions del vostre sistema d'IA i les heu comunicat a totes les persones implicades, incloent-hi les persones usuàries i pacients?
1.1.3.	Heu donat informació sobre l'ús previst del sistema d'IA a les persones pacients i de com les seves dades són processades, quins són els resultats previstos i quina influència tenen en la decisió clínica final?
1.1.4.	Totes les persones implicades en l'ús de dades i el sistema d'IA, en particular les usuàries clíniques, coneixen per a què s'utilitza el sistema d'IA i les seves limitacions?
1.1.5.	Heu tingut en compte l'existència del sistema d'IA en el procés del consentiment informat i n'heu avisat de manera clara i amb llenguatge comprensible?
1.1.6.	Heu establert mecanismes de documentació per facilitar la traçabilitat enfocada en la detecció d'errors, el retiment de comptes i la capacitat de les persones pacients de qüestionar el funcionament del sistema d'IA?
1.1.7.	El procés de generació de resultats del vostre sistema d'IA està degudament documentat i és consultable per part dels usuaris clínics en un llenguatge comprensible i accessible?
1.1.8.	Heu establert un procés de documentació i accés controlat al material rellevant (en concret el conjunt de dades), als models d'IA i als resultats de les avaluacions existents sobre el sistema d'IA?
1.1.9.	Heu documentat el tipus de dades, la seva procedència i els canvis efectuats a totes les fases del cicle de vida del vostre sistema d'IA?

1.2. Interpretabilitat del sistema i informació al pacient

Nº	Preguntes
1.2.1.	Coneixeu si el vostre sistema d'IA utilitza el model més intel·ligible garantint el rendiment absolut en favor de la transparència, o quin ha estat el criteri d'interpretabilitat inclòs en la metodologia de selecció del model d'IA?
1.2.2.	Heu considerat quins coneixements tècnics i quins materials didàctics addicionals són necessaris per entendre i explicar el funcionament del sistema d'IA a les persones pacients?
1.2.3.	Considereu que el sistema d'IA i l'obtenció de resultats és comprensible per a les persones pacients tenint en compte llurs capacitats i competències?

1.3 Explicabilitat del sistema

Nº	Preguntes
1.3.1.	En el cas que el sistema estigui basat en variables clíniques, heu implementat mecanismes per determinar de forma clara quines són les variables clíniques més rellevants que utilitza el vostre sistema d'IA quan fa una predicció?
1.3.2.	Sou coneixedors del tipus de model que utilitza el vostre sistema d'IA i quin és el seu grau d'explicabilitat?
1.3.3.	Considereu que heu establert un nivell de transparència i explicabilitat que no posa en risc la seguretat del vostre sistema d'IA?
1.3.4.	S'ha validat a nivell clínic per part de les persones professionals que utilitzaran el sistema d'IA si els resultats que dona estan ben explicats i són comprensibles, són adequats i ajuden en la tasca a realitzar?
1.3.5.	Sou coneixedors de les hipòtesis fetes i les mètriques que s'han utilitzat per optimitzar i validar el sistema d'IA així com la metodologia de validació per desenvolupar el sistema?
1.3.6.	Heu desenvolupat eines que permeten identificar i traçar les diferents accions i persones involucrades en el sistema d'IA?

- | | |
|--------|---|
| 1.3.7. | Teniu constància de si els mecanismes de transparència i traçabilitat implementats segueixen el marc regulador corresponent de la Unió Europea? |
| 1.3.8. | Per tal de protegir les dades de persones menors d'edat heu tingut en compte salvaguardes com la traçabilitat de les dades per verificar-ne l'edat? |

1.4 Política de participació i protocols de comunicació

Nº	Preguntes
1.4.1.	Els protocols ètics preveuen l'ús de sistemes d'IA per assistir a la pràctica clínica?
1.4.2.	Heu establert procediments per comunicar de manera regular informació sobre en quines decisions clíniques les persones participants són assistides per un sistema d'IA?
1.4.3.	Heu establert eines de comunicació que permeten a les persones usuàries del vostre sistema d'IA realitzar comentaris sobre el seu funcionament?
1.4.4.	Heu consultat amb d'altres organitzacions que hagin dissenyat o implementat un sistema d'IA semblant per conèixer antecedents i experiències?

Principi 2: Justícia i equitat

2.1. Identificació i mitigació de biaixos

Nº	Preguntes
2.1.1.	Coneixeu si el conjunt de dades amb què s'ha entrenat i validat el sistema d'IA és representatiu de la població a la qual s'aplica?
2.1.2.	Heu desenvolupat mecanismes per detectar i corregir possibles biaixos en el comportament del sistema d'IA?
2.1.3.	Heu avaluat l'impacte dels possibles biaixos del vostre sistema d'IA per saber si podria reforçar discriminacions, estereotips o prejudicis envers

	persones pertanyents a grups vulnerables (ex. persones amb discapacitat, minories ètniques o individus d'edat avançada)?
2.1.4.	Heu realitzat una avaluació sobre l'impacte que pot tenir l'ús del sistema d'IA sobre els drets fonamentals de les persones pacients?
2.1.5.	Heu adoptat mesures que permetin mitigar possibles biaixos cognitius, com per exemple els biaixos de confirmació o d'automatització, tant en la fase de disseny com en la implementació del sistema?
2.1.6.	Heu adoptat mesures que permetin mitigar possibles biaixos que es puguin generar a partir de les dades d'entrada o problemes de control del sistema, tant en la fase de disseny com en la implementació del sistema?
2.1.7.	Heu establert una política de participació de diferents grups d'interès, incloent-hi els usuaris clínics i els pacients, del vostre sistema d'IA amb l'objectiu de mitigar possibles biaixos?

2.2. Inclusió i diversitat

Nº	Preguntes
2.2.1.	Heu avaluat si el vostre sistema s'adapta a un ampli espectre de preferències i habilitats individuals, especialment quan es tracti de tecnologia d'assistència que hagi de ser usada directament per alguna persona pacient?
2.2.2.	Heu valorat que dins de l'equip de desenvolupament i/o implementació del vostre sistema d'IA s'incloguin dones?
2.2.3.	Heu valorat que dins de l'equip de desenvolupament i/o implementació del vostre sistema d'IA hi hagi diversitat poblacional?
2.2.4.	Heu valorat que dins de l'equip de desenvolupament i/o implementació del vostre sistema d'IA s'incloguin de diferents disciplines de coneixement més enllà de la ciència de dades, l'enginyeria i les ciències de la salut?

Principi 3: Seguretat i no maleficència

3.1. Guies d'utilització i ús segur del sistema

Nº	Preguntes
3.1.1.	Heu establert guies d'utilització i bones pràctiques dels sistemes d'IA que són necessàries i accessibles a les persones usuàries clíniques?
3.1.2.	Considereu que els desenvolupadors del sistema d'IA han proporcionat suficient informació per fer-ne un ús segur i fiable que no posi en perill a les persones pacients?
3.1.3.	Heu desenvolupat mecanismes clars i ben definits per tal que les persones usuàries clíniques puguin contactar amb els desenvolupadors del sistema d'IA en cas de necessitat?
3.1.4.	Heu desenvolupat mecanismes per actualitzar periòdicament la informació disponible sobre el rendiment del sistema d'IA?

3.2. Gestió del risc

Nº	Preguntes
3.2.1.	El vostre sistema ha estat avaluat sobre els seus possibles impactes a la salut, la seguretat i els drets de les persones?
3.2.2.	El vostre sistema ha passat per una avaluació per garantir que compleix amb els requisits de seguretat legals?
3.2.3.	Heu implementat alguna metodologia d'avaluació o autoavaluació periòdica (interna o externa) reconeguda per identificar possibles funcionaments erronis i incrementar la confiança en l'ús del sistema d'IA?
3.2.4.	Heu seguit metodologies robustes, reconegudes i reproduïbles de validació per seleccionar i entrenar les components algorítmiques del sistema d'IA en totes les etapes del seu desenvolupament?
3.2.5.	Heu informat sobre els riscos residuals (aquells que no han pogut ser eliminats o corregits després d'haver pres totes les mesures possibles)?
3.2.6.	Heu tingut en compte la possibilitat que s'involucri a persones d'una

especialitat clínica diferent de l'actora principal durant el context de desplegament de la solució?

Principi 4: Responsabilitat i retiment de comptes

4.1. Control

Nº	Preguntes
4.1.1.	Heu establert mecanismes de control amb participació humana durant les diferents fases de vida del sistema d'IA (en especial en la fase de recollida i tractament de dades, en els pilots i els assajos clínics amb personal clínic i/o pacients)?
4.1.2.	Heu establert mecanismes de control amb participació humana que permetin mitigar possibles errors del sistema i habilitar modificacions en el protocol clínic d'ús sense que això comprometi la qualitat del resultat?
4.1.3.	Heu identificat i definit clarament el nivell de control de cadascuna de les persones involucrades en la creació, el desenvolupament i l'ús del sistema d'IA?
4.1.4.	Heu desenvolupat eines que permetin validar i contrastar clínicament els resultats del sistema d'IA?
4.1.5.	Heu definit clarament quina és la contribució del sistema d'IA en la presa de decisions clínica i quina és la responsabilitat de les persones involucrades en el control i operació (en especial les usuàries clíniques) en la decisió clínica final?
4.1.6.	Heu definit i comprovat el nivell de formació, comprensió i habilitats necessàries per part del personal clínic involucrat en el control i operació del sistema d'IA?
4.1.7.	Heu establert mecanismes de control amb participació humana per limitar l'ús del sistema per part de grups de pacients determinats o en contextos clínics específics, quan el rendiment pugui esdevenir inacceptable, sense que això comprometi la qualitat del resultat?
4.1.8.	Us heu assegurat que durant tot en el procés de creació, desenvolupament, desplegament i disseny dels mecanismes de control

	del sistema hi hagin participat totes les parts implicades? En particular, les expertes clíniques en el context de la solució, les representants dels pacients i les creadores de sistemes d'IA.
4.1.9.	Heu establert mecanismes perquè les persones creadores i desenvolupadores del sistema d'IA tinguin coneixement del context clínic d'ús, del rendiment i la possibilitat de corregir i millorar aquests?
4.1.10.	En cas d'indisponibilitat temporal del personal que opera el sistema d'IA, heu tingut en compte de quins recursos disposeu per a assegurar el seu correcte funcionament?
4.1.11.	Heu establert quines són les persones responsables d'actualitzar el sistema d'IA que permeti garantir el seu funcionament òptim i assegurar-ne el control en tot moment?
4.1.12.	Heu establert mecanismes per revisar les recomanacions fetes pel sistema d'IA, les discrepàncies entre aquestes i la decisió clínica final i investigar proactivament els possibles errors (incloent-hi la planificació d'accions o protocols a dur a terme)?

4.2. Retiment de comptes

Nº	Preguntes
4.2.1.	Heu determinat qui és legalment responsable de les decisions derivades de l'ús del sistema d'IA durant tot el seu cicle de vida?
4.2.2.	Heu determinat en quins casos, i en quin grau, personal clínic usuari del sistema d'IA pot tenir responsabilitat legal respecte de les decisions clíniques derivades del seu ús?
4.2.3.	Heu determinat en quins casos, i en quin grau, l'organització sanitària sota el paraigua de la qual s'utilitza el sistema d'IA pot tenir responsabilitat legal respecte de les decisions clíniques derivades del seu ús?
4.2.4.	Heu determinat en quins casos, i en quin grau, els creadors del sistema d'IA poden tenir responsabilitat legal respecte de les decisions clíniques derivades del seu ús?
4.2.5.	Heu establert mecanismes de revisió periòdica sobre responsabilitat i retiment de comptes amb els diferents agents involucrats en el

	desenvolupament i l'ús del vostre sistema d'IA?
4.2.6.	El procés de retiment de comptes que heu dissenyat contempla que aquest sigui contínuament informat i revisat per les persones professionals en la pràctica clínica que interpreten les dades i segueixen el procés clínic habitual?
4.2.7.	Heu tingut en compte en el desplegament del sistema d'IA la possible separació que existeix entre la persona usuària clínica i la persona pacient i l'efecte que pot tenir en l'atribució de responsabilitats i el retiment de comptes?

4.3. Figures de protecció i mesures de reparació

Nº	Preguntes
4.3.1.	Heu determinat quines figures de protecció (ex. responsable d'atendre les queixes relacionades amb el sistema d'IA) es poden establir per minimitzar o eliminar possibles impactes negatius derivats de l'ús del sistema d'IA?
4.3.2.	Heu establert formes de reparació o compensació (ex. a través de pòlisses d'assegurances o altres mitjans monetaris adequats) en cas d'ocórrer algun mal funcionament del sistema o provocar un impacte advers a la persona pacient?
4.3.3.	Heu considerat formes de reparació o compensació que incloguin també la responsabilitat per a les persones creadores del sistema d'IA perquè l'atribució de responsabilitats sigui recíproca?
4.3.4.	Heu establert mesures per facilitar la demostració del compliment de la legislació vigent, incloent-hi la Llei de Protecció de dades (GDPR), la regulació sobre dispositius mèdics (Reglament 2017/745) i/o de diagnòstic in vitro (Reglament 2017/746) i el RIA?
4.3.5.	Heu estudiat amb detall si el sistema d'IA pot ser considerat un producte sanitari, d'acord amb la regulació sobre productes sanitaris i/o de diagnòstic in vitro, quina seria la seva classificació i quines conseqüències legals se'n deriven?
4.3.6.	Heu estudiat amb detall quina és la classificació del vostre sistema d'acord amb el Reglament europeu sobre IA i quines conseqüències se'n deriven d'aquesta classificació?

Principi 5: Privacitat

5.1. Protecció de la privacitat

Nº	Preguntes
5.1.1.	Heu revisat si les dades de les persones pacients emprades en el sistema de dades o d'IA han estat anonimitzades tal com indica la Llei Orgànica de Protecció de dades (LO 3/18)?
5.1.2.	Heu tingut en compte que les mesures d'anonimització no poden limitar-se a ocultar el nom de la persona pacient, sinó que han d'impedir que pugui inferir-se la seva identitat a partir del conjunt d'informació?
5.1.3.	En cas d'assistir-vos amb una IA generativa, heu procurat de no avocar al sistema les dades personals de les persones pacients i heu revisat la política de privacitat del mateix?
5.1.4.	Considerereu que heu establert un nivell de transparència i explicabilitat que no posa en risc la privacitat de les persones pacients?
5.1.5.	Podeu per garantir que totes les parts implicades en el desenvolupament i implementació del sistema han establert mecanismes que permetin tractar les dades personals amb el grau de protecció legalment requerit i així evitar-ne el mal ús, alteració, pèrdua o sostracció per part de tercers?
5.1.6.	Podeu garantir que totes les parts implicades en el desenvolupament i implementació del sistema d'IA han establert mecanismes per tal que les dades de les persones afectades no siguin venudes, llogades ni cedides, sense el seu consentiment?

5.2. Governança de dades

Nº	Preguntes
5.2.1.	Heu establert un registre de la informació d'accés a les dades (personals o no) per identificar clarament quan, qui i amb quin propòsit hi ha accedit?
5.2.2.	Heu documentat el procés de recopilació i preparació de les dades per al vostre sistema d'IA?

- 5.2.3. El conjunt de dades recopilades per entrenar el vostre sistema són adequades pel context en el qual serà utilitzat i són representatives de la població on s'utilitzarà el producte sanitari?

Principi 6: Autonomia

6.1. Ús de dades i sistemes d'IA contestables i punt de vista del pacient

Nº	Preguntes
6.1.1.	Heu establert un procediment clar perquè qualsevol persona pacient pugui qüestionar el sistema d'IA, el seu ús previst, biaixos existents, els processos de validació i el rendiment esperat?
6.1.2.	Heu establert un procediment clar perquè qualsevol persona pacient pugui qüestionar els aspectes concrets sobre l'ús de llurs dades?
6.1.3.	Heu establert un procediment clar perquè qualsevol persona pacient pugui qüestionar la divisió entre el sistema d'IA i la persona usuària clínica en la contribució a la presa de decisions clíniques?
6.1.4.	En el desplegament del sistema d'IA heu tingut en compte la possible separació que existeix entre la persona usuària clínica i la persona pacient afectada per la decisió clínica i l'efecte que pot tenir en el seu ús?
6.1.5.	Heu tingut en compte en el desplegament del sistema d'IA la possible separació que existeix entre la persona usuària clínica i la persona pacient i l'efecte que pot tenir en l'anàlisi de beneficis, costos i riscos respecte al seu ús per a cadascuna de les persones involucrades?
6.1.6.	Heu tingut en compte en aquest anàlisi de beneficis i riscos la visió de la persona pacient, de la persona usuària mèdica, del centre associat i del sistema sanitari en general?
6.1.7.	Heu establert un procés d'avaluació del vostre sistema d'IA per garantir que en cap cas disminueix el nombre d'eleccions possibles, ni tampoc l'obliga o força a prendre unes determinades decisions?
6.1.8.	Heu establert un protocol de comprovació que permeti certificar que el vostre sistema ajuda les persones a prendre millors decisions i més informades, d'acord amb els seus objectius i respectant la seva autonomia?

6.2. Precaució envers l'ús de sistemes d'IA generativa

Nº	Preguntes
6.2.1.	Prèviament a realitzar una consulta de caràcter sanitari a un sistema d'IA generativa heu tingut en compte que existeix el risc que el resultat generat pel sistema sigui fruit d'una al·lucinació (resposta incorrecta)?
6.2.2.	En cas de voler utilitzar, o haver utilitzat, un sistema d'IA generativa per a resoldre qüestions de caràcter sanitari, sou conscients que les recomanacions del sistema poden ser poc precises o poc adequades i que l'opinió d'una persona professional de la salut sempre serà més idònia?
6.2.3.	En cas de voler utilitzar, o haver utilitzat, un sistema d'IA generativa per a assistir-vos en la realització d'una tasca o presa de decisió, sou conscients que heu de revisar que el resultat generat sigui correcte?
6.2.4.	En cas de voler utilitzar, o haver utilitzat, un sistema d'IA generativa per a redactar un document, sou conscients que en el mateix document s'ha de fer constar que ha estat generat per un sistema d'IA?
6.2.5.	Dins de la vostra organització heu portat a terme activitats de formació i capacitatció que permetin al personal clínic utilitzar eines d'IA generativa de forma conscient i segura?

Principi 7: Sostenibilitat

7.1. Benestar laboral i comunitari

Nº	Preguntes
7.1.1	Com a organisme públic heu tingut una participació activa en la discussió sobre principis ètics per una IA de confiança?
7.1.2.	Considereu que el vostre sistema d'IA té un impacte positiu en els drets laborals i/o en la qualitat del treball de les persones treballadores del sector sanitari?
7.1.3.	Col·laboreu activament amb organitzacions laborals o sindicats per mitigar el possible desplaçament laboral causat per l'adopció de tecnologies d'IA a través de la requalificació de les treballadores afectades?

7.2. Impacte mediambiental

Nº	Preguntes
----	-----------

7.2.1.	Heu avaluat l'eficiència energètica del vostre sistema d'IA tenint en compte el seu cicle de vida?
--------	--

7.2.2.	En el disseny i implementació de les vostres estratègies per a reduir la petjada ambiental heu tingut en compte la despesa energètica que comporten els sistemes d'IA o les heu pogut adaptar per incloure aquest aspecte?
--------	--

ANNEX II: Catàleg legal i de recomanacions

El present catàleg conté les normes jurídiques en les que s'han inspirat les preguntes del Model PIO Salut. Així mateix, s'enumeren textos d'altres tipologies que contenen recomanacions relatives a l'ús ètic de la IA en l'àmbit de la Salut

	ARTICLE	PREGUNTA MODEL PIO
1. TRANSPARÈNCIA I EXPLICABILITAT		
1.1. ANÀLISI PRELIMINAR DE L'ÚS DEL SISTEMA		
Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	8 , 9 i 11	1.1.1, 1.1.2, 1.1.3, 1.1.4, 1.1.5, 1.1.6, 1.1.7, 1.1.8. i 1.1.9.
Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina [Conveni d'Oviedo]	5 i 10	1.1.1, 1.1.2, 1.1.3, 1.1.4, 1.1.5, 1.1.7, 1.1.8 i 1.1.9.
Directiva 2024/2853 del Parlament Europeu i del Consell, de 23 d'octubre de 2024, sobre responsabilitat pels danys causats per productes defectuosos	7	1.1.2, 1.1.4, 1.1.6, 1.1.7. i 1.1.9.
Estatut d'Autonomia de Catalunya	23.3	1.1.5.
Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica	2 , 4 i 5	1.1.2. i 1.1.5.
Reglament 2022/868 del Parlament Europeu i del Consell, de 30 de maig del 2022, relatiu a la governança europea de dades [Data Governance Act]	21	1.1.3.
Reglament 2016/679, del Parlament Europeu i del Consell, de 27 d'abril de 2016, relatiu a la protecció de les persones físiques pel que fa al tractament de dades personals i a la lliure circulació d'aquestes dades [Reglament general de protecció de dades]	13	1.1.3.

Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris	10 , 29 , 32 , 63 , 83 i 110 ,	1.1.1, 1.1.2, 1.1.3, 1.1.4, 1.1.5, 1.1.6, 1.1.7. i 1.1.8.
Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro	10 , 26 , 29 , 59 , 78 , 102 i 103	1.1.1, 1.1.2, 1.1.3, 1.1.4, 1.1.5, 1.1.6, 1.1.7 i 1.1.8.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	10 , 11 , 12 , 13 , 14 , 17 , 50 , 53 i 55	1.1.1, 1.1.2, 1.1.3, 1.1.4, 1.1.5, 1.1.6, 1.1.7, 1.1.8 i 1.1.9.
Reial decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics	4	1.1.5.
Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris	5 i 21	1.1.1, 1.1.2, 1.1.4, 1.1.5, 1.1.6. i 1.1.7.
AI and data protection risk toolkit d'ICO		1.1.2, 1.1.3. i 1.1.7.
AI ethics assessment de GSMA		1.1.1, 1.1.4. i 1.1.5.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment		1.1.1, 1.1.4, 1.1.5, 1.1.6, 1.1.7. i 1.1.9.
Ethics and governance of artificial intelligence for health (WHO)		1.1.4. i 1.1.5.
Generative AI and medical ethics: the state of play (Zohny, Mann, <i>et al</i>)		1.1.5.
The ethics of AI in health care: A mapping review (Morleya, Machadoa, <i>et al</i>.)		1.1.3.

[Transparency of AI in Healthcare as a Multilayered System of Accountabilities: Between Legal Requirements and Technical Limitations, \(Anastasia Kiseleva *et al.*\)](#)

1.1.6, 1.1.7.
i 1.1.8.

1.2. INTERPRETABILITAT DEL SISTEMA I INFORMACIÓ AL PACIENT

[Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#)

8 i 20

1.2.1, 1.2.2.
i 1.2.3.

[Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina \[Conveni d'Oviedo\]](#)

5

1.2.1, 1.2.2.
i 1.2.3.

[Estatut d'Autonomia de Catalunya](#)

23.3

1.2.2. i 1.2.3.

[Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica, relatius al dret a la informació assistencial](#)

2 i 4

1.2.2.

[Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#)

10 i 63

1.2.2. i 1.2.3.

[Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#)

10 i 59

1.2.2. i 1.2.3.

[Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#)

12, 13, 16,
53 i 95

1.2.1, 1.2.2.
i 1.2.3.

[Reial decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics](#)

4

1.2.2.

[Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#)

5

1.2.2. i 1.2.3.

[An interdisciplinary framework to operationalise AI ethics d'AIEIG](#) 1.2.1.

[Assessment list for Trustworthy artificial Intelligence \(ALTAI\) for self assessment](#) 1.2.1.

1.3. EXPLICABILITAT DEL SISTEMA

[Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret](#) 8, 9, 12 i 18 1.3.2, 1.3.3, 1.3.4, 1.3.5, 1.3.6, 1.3.7. i 1.3.8.

[Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina \[Conveni d'Oviedo\]](#) 5 1.3.2. i 1.3.4

[Directiva 2024/2853 del Parlament Europeu i del Consell, de 23 d'octubre de 2024, sobre responsabilitat pels danys causats per productes defectuosos](#) 7 1.3.1, 1.3.3, 1.3.4, 1.3.5. i 1.3.7.

[Reglament 2016/679, del Parlament Europeu i del Consell, de 27 d'abril de 2016, relatiu a la protecció de les persones físiques pel que fa al tractament de dades personals i a la lliure circulació d'aquestes dades\[Reglament general de protecció de dades\]](#) 8 1.3.8.

[Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#) 10, 19, 20, 25-34, 32 i 52 1.3.1, 1.3.4, 1.3.5, 1.3.6. i 1.3.7.

[Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#) 10, 17, 18, 22 - 30, 29 i 48 1.3.1, 1.3.4, 1.3.5, 1.3.6. i 1.3.7.

[Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#) 9, 10, 12, 13, 14, 15, 17, 43, 53 i 95 1.3.1, 1.3.2, 1.3.3, 1.3.4, 1.3.5, 1.3.6, 1.3.7. i 1.3.8.

[Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#) 5 i 21 1.3.1, 1.3.4, 1.3.5, 1.3.6. i 1.3.7.

AI ethics assessment de GSMA	1.3.1, 1.3.2. i 1.3.3.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment	1.3.1, 1.3.2, 1.3.3, 1.3.5. i 1.3.8.
The ethics of AI in health care: A mapping review (Morleya, Machadoa, et al.)	1.3.6.

1.4. POLÍTICA DE PARTICIPACIÓ I PROTOCOLS DE COMUNICACIÓ

Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	8	1.4.2.
Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina [Conveni d'Oviedo]	5 i 10	1.4.2.
Llei 33/2011, del 4 d'octubre, general de Salut Pública	5	1.4.3.
Llei 14/1986, de 25 d'aril, General de Sanitat	10.10	1.4.3.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	50 i 95	1.4.1, 1.4.2, 1.4.3. i 1.4.4.
AI and data protection risk toolkit d'ICO		1.4.2. i 1.4.3.
Algorithmic Equity Toolkit d'ACLU		1.4.1.
An interdisciplinary framework to operationalise AI ethics d'AIEIG		1.4.2. i 1.4.3.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment		1.4.1, 1.4.2, 1.4.3. i 1.4.4.

[Recommendation on the Ethics of Artificial Intelligence de la UNESCO](#)

1.4.2. i 1.4.3.

	ARTICLE	PREGUNTA MODEL PIO
2. JUSTÍCIA I EQUITAT		
2.1. IDENTIFICACIÓ I MITIGACIÓ DE BIAIXOS		
Constitució Espanyola	14	2.1.1, 2.1.2, 2.1.3, 2.1.4, 2.1.5. i 2.1.6.
Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	4 , 6 , 10 , 16 .	2.1.1, 2.1.2, 2.1.3, 2.1.4, 2.1.5, 2.1.6. i 2.1.7.
Directiva 2024/2853 del Parlament Europeu i del Consell, de 23 d'octubre de 2024, sobre responsabilitat pels danys causats per productes defectuosos	7	2.1.1, 2.1.2, 2.1.3, 2.1.4, 2.1.5. i 2.1.6.
Estatut d'Autonomia de Catalunya	15	2.1.1, 2.1.2, 2.1.3, 2.1.4, 2.1.5. i 2.1.6.
Llei 33/2011, del 4 d'octubre, general de Salut Pública	5 , 6	2.1.1, 2.1.2, 2.1.3, 2.1.4, 2.1.5, 2.1.6. i 2.1.7.
Llei 14/1986, de 25 d'aril, General de Sanitat	10 i 10	2.1.5. i 2.1.7.
Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris	10 i 83	2.1.2.
Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro	10 i 78	2.1.2.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	9 , 10 , 15 , 27.1 , 53 , 55 , 72 i 95	2.1.1, 2.1.2, 2.1.3, 2.1.4,

		2.1.5, 2.1.6. i 2.1.7.
Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris	5	2.1.2.
Algorithmic Equity Toolkit d'ACLU		2.1.7.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment		2.1.7.
FUTURE-AI (Karim Lekadir, Alejandro F Frangi, <i>et al.</i>)		2.1.1, 2.1.2, 2.1.3. i 2.1.6.
The ethics of AI in health care: A mapping review (Morleya, Machadoa, <i>et al.</i>)		2.1.1, 2.1.2, 2.1.4, 2.1.5, 2.1.6. i 2.1.7.

2.2. INCLUSIÓ I DIVERSITAT

Constitució Espanyola	14	2.2.1.
Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	18	2.2.1.
Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina [Conveni d'Oviedo]	3	2.2.1.
Directiva 2024/2853 del Parlament Europeu i del Consell, de 23 d'octubre de 2024, sobre responsabilitat pels danys causats per productes defectuosos	7	2.2.1.
Estatut d'Autonomia de Catalunya	15	2.2.1.
Llei 33/2011, del 4 d'octubre, general de Salut Pública	6	2.1.1.

Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	95	2.2.1, 2.2.2, 2.2.3. i 2.2.4.
A Human Rights-Based Approach To Data d'AC-NUDH		2.2.2, 2.2.3. i 2.2.4.
AI and data protection risk toolkit d'ICO		2.2.2, 2.2.3. i 2.2.4.
AI ethics assessment de GSMA		2.2.1.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment		2.2.1, 2.2.2, 2.2.3. i 2.2.4.
Embedded ethics: a proposal for integrating ethics into the development of medical AI (McLennan, Fiske, et al)		2.2.2, 2.2.3. i 2.2.4.
FUTURE-AI (Karim Lekadir, Alejandro F Frangi, et al.)		2.2.1.

	ARTICLE	PREGUNTA MODEL PIO
3. SEURETAT I NO MALEFICÈNCIA		
3.1. GUIES D'UTILITZACIÓ I ÚS SEGUR DEL SISTEMA		
Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	8 i 14	3.1.3 i 3.1.4.
Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina [Conveni d'Oviedo]	5	3.1.1. i 3.1.4.
Directiva 2024/2853 del Parlament Europeu i del Consell, de 23 d'octubre de 2024, sobre responsabilitat pels danys causats per productes defectuosos	7	3.1.1, 3.1.2. i 3.1.4.

Estatut d'Autonomia de Catalunya	23	3.1.4.
Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica	4 i 5	3.1.4.
Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris	10	3.1.1, 3.1.2, 3.1.3 i 3.1.4.
Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro	10	3.1.1, 3.1.2, 3.1.3 i 3.1.4.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	11 , 13 , 50 , 53 i 95	3.1.1, 3.1.2, 3.1.3 i 3.1.4.
Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris	5	3.1.3 i 3.1.4.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment		3.1.4.
FUTURE-AI (Karim Lekadir, Alejandro F Frangi, <i>et al.</i>)		3.1.1 i 3.1.4.

3.2. GESTIÓ DEL RISC

Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	8 , 12 i 16	3.2.1, 3.2.2, 3.2.3, 3.2.4 i 3.2.5.
Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina [Conveni d'Oviedo]	4 i 5	3.2.2, 3.2.3 i 3.2.5.
Directiva 2024/2853 del Parlament Europeu i del Consell, de 23 d'octubre de 2024, sobre responsabilitat pels danys causats per productes defectuosos	7	3.2.1, 3.2.2, 3.2.3, 3.2.4, 3.2.5 i 3.2.6.
Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica	10	3.2.5.

Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris	10 , 32 i 62	3.2.2, 3.2.3. i 3.2.5.
Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro	10 , 29 i 58	3.2.2, 3.2.3. i 3.2.5.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	9 , 14 , 15 , 17 , 27 , 55 , 72 i 95	3.2.1, 3.2.2, 3.2.3, 3.2.4, 3.2.5. i 3.2.6.
Reial decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics	3.b) i 4	3.2.2. i 3.2.5.
Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris	5	3.2.2, 3.2.3. i 3.2.5.
Ethics and governance of artificial intelligence for health (WHO)		3.2.6.
FUTURE-AI (Karim Lekadir, Alejandro F Frangi, <i>et al.</i>)		3.2.3. i 3.2.5.
Holding AI to Account: Challenges for the Delivery of Trustworthy AI in Healthcare, de Rob Procter, Peter Tolmie, i Mark Rouncefield		3.2.3. i 3.2.4.
Informe sobre avaluacions d'impacte dels Drets Fonamentals en relació amb la Llei de Serveis Digitals d'AN i ECNL.		3.2.1.
The ethics of AI in health care: A mapping review (Morleya, Machado, <i>et al.</i>)		3.2.1.
The medical algorithmic audit (Liu, Glocker, <i>et al.</i>)		3.2.3. i 3.2.4.

	ARTICLE	PREGUNTA MODEL PIO
4. RESPONSABILITAT I RETIMENT DE COMPTES		
4.1. CONTROL		
Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	8 , 9 , 12 , 16 i 20	4.1.1, 4.1.2, 4.1.3, 4.1.4, 4.1.5, 4.1.6, 4.1.7, 4.1.8, 4.1.9, 4.1.10, 4.1.11. i 4.1.12.
Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina [Conveni d'Oviedo]	4	4.1.1, 4.1.2, 4.1.5 i 4.1.7.
Directiva 2024/2853 del Parlament Europeu i del Consell, de 23 d'octubre de 2024, sobre responsabilitat pels danys causats per productes defectuosos	7	4.1.1, 4.1.2, 4.1.3, 4.1.4, 4.1.5, 4.1.6, 4.1.7, 4.1.11. i 4.1.12.
Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris	10 , 32 , 62 i 86	4.1.1, 4.1.4, 4.1.6, 4.1.9. i 4.1.12.
Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro	10 , 29 i 58	4.1.1, 4.1.4, 4.1.6, 4.1.9. i 4.1.12.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	9 , 14 , 17 , 55 , 72 i 95	4.1.1, 4.1.2, 4.1.3, 4.1.4, 4.1.5, 4.1.6, 4.1.7, 4.1.8, 4.1.9, 4.1.11. i 4.1.12.
Reial decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics	3.f)	4.1.1.

Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris	5 i 8	4.1.1, 4.1.4, 4.1.6, 4.1.9, 4.1.10 i 4.1.12.
AI ethics assessment de GSMA		4.1.3.
Algorithmic Equity Toolkit d'ACLU		4.1.3.
Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability de Coeckelbergh		4.1.1, 4.1.2, i 4.1.7.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment		4.1.3, 4.1.5. i 4.1.12
Ethics and governance of artificial intelligence for health (WHO)		4.1.8, 4.1.9, 4.1.10, 4.1.11. i 4.1.12.
FUTURE-AI (Karim Lekadir, Alejandro F Frangi, <i>et al.</i>)		4.1.4, 4.1.6. i 4.1.9.
The ethics of AI in health care: A mapping review (Morleya, Machadoa, <i>et al.</i>)		4.1.6.
The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine (Thomas Ploug i Søren Holm)		4.1.8. i 4.1.9.
The medical algorithmic audit (Liu, Glocker, <i>et al.</i>)		4.1.1, 4.1.2, 4.1.7, 4.1.11. 4.1.12.

4.2. RETIMENT DE COMPTE

Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	9 i 16	4.2.1, 4.2.2, 4.2.3, 4.2.4, 4.2.5, 4.2.6. i 4.2.7.
--	--	--

Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris	10	4.2.1, 4.2.2, 4.2.3, 4.2.4, 4.2.5, 4.2.6. i 4.2.7.
Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro	10	4.2.1, 4.2.2, 4.2.3, 4.2.4, 4.2.5, 4.2.6. i 4.2.7.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	13 , 14 , 17 i 72	4.2.1, 4.2.2, 4.2.3, 4.2.4, 4.2.5, 4.2.6. i 4.2.7.
Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris	5	4.2.1, 4.2.2, 4.2.3, 4.2.4, 4.2.5, 4.2.6. i 4.2.7.
Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability de Coeckelbergh		4.2.7.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment		4.2.5. i 4.2.6.
Holding AI to Account: Challenges for the Delivery of Trustworthy AI in Healthcare, de Rob Procter, Peter Tolmie, i Mark Rouncefield.		4.2.5. i 4.2.6.
The ethics of AI in health care: A mapping review (Morleya, Machadoa, et al.)		4.2.1, 4.2.2, 4.2.3. i 4.2.4.
The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine (Thomas Ploug i Søren Holm)		4.2.7.

4.3. FIGURES DE PROTECCIÓ I MESURES DE REPARACIÓ

Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	9 i 16	4.3.1, 4.3.2. i 4.3.3.
--	--	------------------------

Llei 33/2011, de 4 d'octubre, General de Salut Pública	5	4.3.1.
Llei 14/1986, de 25 d'aril, General de Sanitat	10.12	4.3.1.
Reglament 2016/679, del Parlament Europeu i del Consell, de 27 d'abril de 2016, relatiu a la protecció de les persones físiques pel que fa al tractament de dades personals i a la lliure circulació d'aquestes dades[Reglament general de protecció de dades]	24 - 31	4.3.4.
Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris	10, 19, 20, 51 i 52	4.3.2, 4.3.3, 4.3.4. i 4.3.5.
Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro	10, 17, 18, 47 i 48	4.3.2, 4.3.3, 4.3.4. i 4.3.5.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	6, 11, 13, 17, 27, 43 i 51	4.3.2, 4.3.4. i 4.3.6.
Reial decret 1090/2015, de 4 de desembre, pel qual es regulen els assajos clínics amb medicaments, els Comitès d'Ètica de la Investigació amb medicaments i el Registre Espanyol d'Estudis Clínics	9	4.3.2. i 4.3.3.
Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris	21 i 32	4.3.2, 4.3.3. i 4.3.4.
AI and data protection risk toolkit d'ICO		4.3.4.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment		4.3.4.
Ethics and governance of artificial intelligence for health (WHO)		4.3.2. i 4.3.3.
Ethics guidelines for trustworthy AI (Comissió Europea)		4.3.2. i 4.3.3.

[Guia per a l'Aplicació del Reglament Europeu d'IA en Salut elaborada \(Fundació TIC Salut Social\)](#)

4.3.6.

	ARTICLE	PREGUNTA MODEL PIO
5. PRIVACITAT		
5.1. PROTECCIÓ DE LA PRIVACITAT		
Constitució Espanyola	18	5.1.1, 5.1.2, 5.1.3, 5.1.4., 5.1.5. i 5.1.6.
Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	11	5.1.1, 5.1.2, 5.1.3, 5.1.4, 5.1.5. i 5.1.6.
Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina [Conveni d'Oviedo]	10	5.1.1, 5.1.2, 5.1.3, 5.1.4, 5.1.5. i 5.1.6.
Estatut d'Autonomia de Catalunya	31	5.1.1, 5.1.2, 5.1.3, 5.1.4, 5.1.5. i 5.1.6.
Llei 14/1986, de 25 d'aril, General de Sanitat	10.3	5.1.1, 5.1.2, 5.1.3, 5.1.4, 5.1.5. i 5.1.6.
Llei 33/2011, de 4 d'octubre, General de Salut Pública	7	5.1.1, 5.1.2, 5.1.3, 5.1.4, 5.1.5. i 5.1.6.
Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica	7	5.1.1, 5.1.2, 5.1.3, 5.1.4, 5.1.5. i 5.1.6.
Reglament 2016/679, del Parlament Europeu i del Consell, de 27 d'abril de 2016, relatiu a la protecció de les persones físiques pel que fa al tractament de dades personals i a la lliure circulació d'aquestes dades [Reglament general de protecció de dades]	32	5.1.1, 5.1.2, 5.1.3, 5.1.4, 5.1.5. i 5.1.6.

Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris	110	5.1.1, 5.1.2. i 5.1.4.
Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro	102 i 103	5.1.1, 5.1.2. i 5.1.4
Reglament 2022/868 del Parlament Europeu i del Consell, de 30 de maig del 2022, relatiu a la governança europea de dades [Data Governance Act]	5 i 21	5.1.1, 5.1.2, 5.1.4, 5.1.5. i 5.1.6.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	10 , 17 , 50 i 55	5.1.1, 5.1.2, 5.1.3, 5.1.4, 5.1.5. i 5.1.6.
A Human Rights-Based Approach To Data d'AC-NUDH		5.1.1, 5.1.2, 5.1.3, 5.1.5. i 5.1.6.
AI and data protection risk toolkit d'ICO		5.1.5. i 5.1.6.
AI ethics assessment de GSMA		5.1.5. i 5.1.6.
An interdisciplinary framework to operationalise AI ethics d'AIEIG		5.1.1, 5.1.2. i 5.1.3.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment		5.1.1, 5.1.2. i 5.1.3.
Ethical Challenges and Solutions of Generative AI: An Interdisciplinary Perspective (Al-kfairy, Mustafa, <i>et al</i>)		5.1.1. i 5.1.2.
Generative AI and medical ethics: the state of play (Zohny, Mann, <i>et al</i>)		5.1.3.
Recommendation on the Ethics of Artificial Intelligence de la UNESCO		5.1.4.

5.2. GOVERNANÇA DE DADES

<u>Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret</u>	<u>11</u> i <u>12</u>	5.2.1. i 5.2.3.
<u>Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina [Conveni d'Oviedo]</u>	<u>10</u>	5.2.1.
<u>Directiva 2024/2853 del Parlament Europeu i del Consell, de 23 d'octubre de 2024, sobre responsabilitat pels danys causats per productes defectuosos</u>	<u>7</u>	5.2.3.
<u>Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica</u>	<u>16</u>	5.2.1.
<u>Reglament 2016/679, del Parlament Europeu i del Consell, de 27 d'abril de 2016, relatiu a la protecció de les persones físiques pel que fa al tractament de dades personals i a la lliure circulació d'aquestes dades [Reglament general de protecció de dades]</u>	<u>24 - 31</u> i <u>30</u>	5.2.1. i 5.2.2.
<u>Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris</u>	<u>32</u>	5.2.3.
<u>Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro</u>	<u>29</u>	5.2.3.
<u>Reglament 2022/868 del Parlament Europeu i del Consell, de 30 de maig del 2022, relatiu a la governança europea de dades [Data Governance Act]</u>	<u>20</u>	5.2.1.
<u>Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]</u>	<u>10, 17, 50</u> i <u>53</u>	5.2.1, 5.2.2. i 5.2.3
<u>Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris</u>	<u>5</u>	5.2.3.

A Human Rights-Based Approach To Data d'AC-NUDH	5.2.3.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment	5.2.1. i 5.2.2.
Generative AI and medical ethics: the state of play (Zohny, Mann, <i>et al</i>)	5.2.3.
The ethics of AI in health care: A mapping review (Morleya, Machadoa, <i>et al.</i>)	5.2.3.

	ARTICLE	PREGUNTA MODEL PIO
6. AUTONOMIA		
6.1. SISTEMES D'IA CONTESTABLES I PUNT DE VISTA DEL PACIENT		
Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	7 , 9 , 11 i 16	6.1.1, 6.1.2, 6.1.3, 6.1.4, 6.1.5, 6.1.6, 6.1.7. i 6.1.8.
Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina [Conveni d'Oviedo]	4 , 5 i 10	6.1.1, 6.1.2, 6.1.3, 6.1.7. i 6.1.8.
Llei 41/2002, de 14 de novembre, bàsica reguladora de l'autonomia del pacient i de drets i obligacions en matèria d'informació i documentació clínica	2.3 , 4 i 8	6.1.8.
Reglament 2016/679, del Parlament Europeu i del Consell, de 27 d'abril de 2016, relatiu a la protecció de les persones físiques pel que fa al tracta-	21	6.1.2.

[ment de dades personals i a la lliure circulació d'aquestes dades \[Reglament general de protecció de dades\]](#)

[Reglament 2017/745 del Parlament Europeu i del Consell, de 5 d'abril de 2017 sobre els productes sanitaris](#)

[10](#)

6.1.1, 6.1.4.
i 6.1.6.

[Reglament 2017/746 de 5 d'abril de 2017, sobre productes sanitaris per a diagnòstic in vitro](#)

[10](#)

6.1.1, 6.1.4.
i 6.1.6.

[Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial \[RIA\]](#)

[13](#), [17](#) i [95](#)

6.1.1, 6.1.3,
6.1.4, 6.1.6,
6.1.7. i
6.1.8.

[Reial Decret 192/2023, de 21 de març, pel qual es regulen els productes sanitaris](#)

[5](#)

6.1.1, 6.1.4.
i 6.1.6.

[AI ethics assessment de GSMA](#)

6.1.7.

[Artificial Intelligence, Responsibility Attribution, and a Relational Justification of Explainability de Coeckelbergh](#)

6.1.1, 6.1.2,
6.1.3, 6.1.4,
6.1.5. i
6.1.6.

[Assessment list for Trustworthy artificial Intelligence \(ALTAI\) for self assessment](#)

6.1.7. i
6.1.8.

[The ethics of AI in health care: A mapping review \(Morleya, Machadoa, *et al.*\)](#)

6.1.8.

[The four dimensions of contestable AI diagnostics - A patient-centric approach to explainable AI, Artificial Intelligence in Medicine \(Thomas Ploug i Søren Holm\)](#)

6.1.1, 6.1.2,
6.1.3, 6.1.4,
6.1.5. i
6.1.6.

6.2. PRECAUCIÓ ENVERS L'ÚS DE SISTEMES D'IA GENERATIVA

Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	8 , 16 i 20	6.2.1, 6.1.2, 6.2.3, 6.2.4. i 6.2.5.
Conveni per a la protecció dels Drets Humans i la Dignitat del Ser Humà pel que fa a les aplicacions de la Biologia i la Medicina, conegut també com a Conveni sobre Drets Humans i Biomedicina [Conveni d'Oviedo]	4	6.2.1, 6.2.2, 6.2.3, 6.2.4. i 6.2.5.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	50 i 53	6.2.1, 6.2.2, 6.2.3 i 6.2.4.
Artificial Intelligence in Medicine: Cross-Sectional Study Among Medical Students on Application, Education, and Ethical Aspects (Weidener i Fischer)		6.2.1, 6.2.2, 6.2.3 i 6.2.4.
Assessment list for Trustworthy artificial Intelligence (ALTAI) for self assessment		6.2.1, 6.2.2, 6.2.3 i 6.2.4.
Generative AI and medical ethics: the state of play (Zohny, Mann, <i>et al</i>)		6.2.3 i 6.2.4.
Generative artificial intelligence and ethical considerations in health care: a scoping review and ethics checklist (Ning, Teixayavongaixí, <i>et al</i>)		6.2.1, 6.2.2, 6.2.3 i 6.2.4.
Guia de Bones Pràctiques per al Desenvolupament d'Eines d'IA Generativa en Salut (Fundació TIC Salut Social)		6.2.1, 6.2.2, 6.2.3, 6.2.4. i 6.2.5.

	ARTICLE	PREGUNTA MODEL PIO
7. SOSTENIBILITAT		
7.1. BENESTAR LABORAL I COMUNITARI		
Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	16	7.1.1, 7.1.2. i 7.1.3.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	27 , 55 i 95	7.1.1, 7.1.2. i 7.1.3.
Ethics and governance of artificial intelligence for health (WHO)		7.1.1.
Ethical Challenges and Solutions of Generative AI: An Interdisciplinary Perspective (Al-kfairy, Mustafa, et al.)		7.1.2. i 7.1.3.
FUTURE-AI (Karim Lekadir, Alejandro F Frangi, et al.)		7.1.2. i 7.1.3.
Recommendation on the Ethics of Artificial Intelligence de la UNESCO		7.1.3.
The ethics of AI in health care: A mapping review (Morleya, Machadoa, et al.)		7.1.1.
7.2. IMPACTE MEDIAMBIENTAL		
Conveni marc del Consell d'Europa sobre intel·ligència artificial i drets humans, democràcia i estat de dret	16	7.2.1 i 7.2.2.
Reglament 2024/1689 del Parlament i del Consell sobre intel·ligència artificial [RIA]	53 i 95	7.2.1 i 7.2.2.
AI ethics assessment de GSMA		7.2.1.

[Assessment list for Trustworthy artificial Intelligence \(ALTAI\) for self assessment](#)

7.2.1. i 7.2.2.

[FUTURE-AI \(Karim Lekadir, Alejandro F Frangi, *et al.*\)](#)

7.2.1. i 7.2.2.

[Recommendation on the Ethics of Artificial Intelligence de la UNESCO](#)

7.2.1. i 7.2.2.

[Resolució 48/13 de 8 d'Octubre de 2021 del Consell de Drets Humans de l'ONU](#)

7.2.2.

