



CLIAS

CENTRO DE INTELIGENCIA
ARTIFICIAL Y SALUD
PARA AMÉRICA LATINA
Y EL CARIBE

Practical guide for the ethical and regulatory assessment of artificial intelligence systems in health

Latin America and the Caribbean

**FOR DEVELOPERS AND IMPLEMENTERS
AI IN HEALTH**





CONTENT

<u>AUTHORS</u>	3
<u>1. ABOUT THIS DOCUMENT</u>	4
<u>LOGIC OF USE</u>	5
<u>2. WHO IS IT AIMED AT</u>	? 6
<u>3. METHODOLOGICAL FOUNDATION</u>	6
3.1 WHY ETHICS FIRST	7
3.2 HUMAN RIGHTS AS A COMMON LEGAL BASIS	9
<u>4. HOW TO USE THIS GUIDE</u>	11
<u>5. TRAFFIC LIGHT SYSTEM</u>	11
<u>6. IMPLEMENTER TEMPLATE</u>	12
<u>7. DEVELOPER TEMPLATE</u>	13
<u>8. CARE ACTION PLAN</u>	15
<u>9. GENERAL QUESTION BANK</u>	17
9.1 ETHICS QUESTIONS	17
9.2 REGULATORY QUESTIONS	18
<u>10. CONCLUSION</u>	20



AUTHORS

PAULA KOHAN: Lawyer. Graduate in Law and Innovative Technologies from Austral University, with a Postgraduate Diploma in Law and Artificial Intelligence from the University of Buenos Aires (UBA). Master's candidate in Civil Law at the National University of La Pampa (UNLPam). Graduate in Computer Procedural Law, Graduate in Digital Health from UBA. Former member of the technical team for the Postgraduate Program in Digital Health at UBA. Graduate in Public Policy and Smart Government from the National Technological University (UTN).

ROSA ANGELINA PACE: Medical Doctor (UNNE) and Master in Bioethics from the Complutense University of Madrid (UCM). Coordinator of the Bioethics Center at the Italian Hospital of Buenos Aires (HIBA), and Director of the Department of Human and Social Sciences at the Italian Hospital University Institute (IUHIBA). Member of the Ethics Council in Medicine, National Academy of Medicine. Recipient of the Velasco Suarez PAHO Bioethics Award. Consultant for CIIPS - IECS.

MARTIN SABAN: Pediatrician (University of Buenos Aires). Member of the Information and Communication Technologies Subcommittee of the Argentine Society of Pediatrics. Candidate for a Master's degree in Clinical and Health Effectiveness. CIIPS-IECS Fellow.

CINTIA CEJAS: Bachelor's degree in Political Science (UCA) and Master's degree in Social and Health Sciences (FLACSO - CEDES). Specialist in health project management. Coordinator of the Center for Implementation and Innovation in Health Policies (CIIPS - IECS) and the Center for Artificial Intelligence in Health for Latin America and the Caribbean (CLIAS).



1. About this document

This guide is a practical tool for those who design, adopt, or monitor AI systems in healthcare. Its purpose is to help make informed decisions from the outset of a project, integrating ethical and regulatory considerations at every stage.

This is not a theoretical framework or a collection of abstract principles. It is a working tool that translates those principles into actionable questions, criteria, and guidelines that allow us to evaluate, at each stage, what should be done, what risks exist, and what aspects need to be reviewed before moving forward.

In Latin America, teams often face regulatory fragmentation, with limitations in infrastructure, connectivity, and the availability of quality data. Furthermore, many AI systems used in the region were developed in other contexts, which can lead to mismatches with the characteristics of local populations. In this scenario, making decisions without a clear framework increases the risk of implementing inadequate, ineffective, or even harmful solutions.

Therefore, this guide aims to support teams at a crucial moment: when the idea of developing or adopting an AI system arises, and it's still possible to adjust course. Its purpose is to help answer, in an informed and structured way, whether a project should move forward, how it should do so, and under what conditions.

It is important to bear in mind that AI systems in healthcare are not neutral. Their results can directly influence clinical decisions, resource allocation, and access to healthcare services. When these systems are poorly designed, poorly implemented, or poorly evaluated, the negative effects are not theoretical: they impact real people, often exacerbating existing inequalities. Therefore, it is essential **to incorporate ethics from the very conception or decision to adopt AI, which is known as "ethics in AI" by design** or **"Embedded ethics"**, which is nothing more than thinking in those terms from minute zero.

Although ethics and regulation are addressed separately to facilitate analysis, they should be understood as deeply interconnected dimensions. The ethical questions that arise in practice—what problem is being addressed, who might be affected, what risks are generated, what biases could be amplified, and what conditions must be met for responsible use—typically guide regulatory responses. In this sense, **regulation should not be considered merely an external requirement, but rather a necessary consequence of having clearly identified what needs to be protected**. On the one hand, this Guide should be used iteratively, that is, not necessarily sequentially, and it should also be revisited periodically: as technology, institutional culture, and social contexts evolve, the answers to these questions may change, and with them, the resulting



actions. Its practical value lies precisely in helping to distinguish which aspects have been resolved and which questions remain open. The existence of unanswered questions does not imply that the project is poorly conceived, but rather that these points represent risks that must be acknowledged, mitigated, and addressed through a care plan.

In this context, this guide proposes a simple approach adapted to the region: stop to evaluate, decide with criteria and move forward only when there are real conditions to do so safely, ethically and sustainably.

WHY IT MATTERS

AI systems in healthcare are not neutral. Their outcomes can influence clinical decisions, resource allocation, and access to services. A poorly designed or poorly evaluated system does not cause theoretical harm; it impacts real people, often exacerbating existing inequalities.

LOGIC OF USE

This guide is not meant to be read cover to cover and filed away. It is intended for use at the moment decisions are made, such as when a developer defines the data with which to train a model, or when an implementer is evaluating the adoption of an AI system for diagnostic support or healthcare management.

Its structure combines two types of questions. The first set comprises cross-cutting questions, applicable at any stage and by any stakeholder, addressing the ethical and regulatory foundations that must be present throughout the entire lifecycle of an AI system in healthcare. The second set consists of specific questions, organized by stage—design, development, implementation, monitoring, and evaluation—that guide the concrete decisions corresponding to each moment of the process.

Both sets can be consulted independently, depending on the user's needs. It is not necessary to complete one section to move on to the next, although reading the entire set provides a more complete understanding of how the ethical and regulatory dimensions are articulated.

The guide does not prescribe unique solutions: it offers tools so that developers and implementers can build their own responses based on their institutional context, their user population, and their available capabilities.

One final point: this guide is a starting point, not a final destination. The field of AI in healthcare is evolving rapidly, and any ethical framework must evolve with it. It is recommended that this guide be reviewed periodically and supplemented with regulatory developments as they arise in each country of Latin America and the Caribbean.

IMPORTANT

This guide should be reviewed periodically. The answers may change over time as data, technology, the institutional context, or applicable regulations evolve.



2. Who is it aimed at?

The guide is aimed at two main profiles, which may overlap in the same institution:

- Developers: those who design, build, or commission the construction of an AI system and put it into operation under their name or institutional brand. This is equivalent to the concept of provider under the European AI ¹Act .
- Implementers: those who adopt and operate a system developed by third parties. This is equivalent to the deployer in the AI Act ². In the region, a hospital can simultaneously be an implementer, if it adopts an external system, and a developer, if it adapts and operates it under its own brand.

No specialized training in ethics or law is required. What is required is a willingness to incorporate these dimensions as an integral part of the work, and not as an afterthought.

3. Methodological Foundations

This guide is based on complementary reference frameworks, specifically: CLIAS Technical Document No. 8 ³, the document “ Ethics and governance of artificial intelligence for health of PAHO/WHO ” ⁴, the “ Ethical Impact UNESCO ⁵Assessment ” , the “Artificial Intelligence in Public Health Readiness Assessment Toolkit developed ⁶by PAHO/WHO and the IDB and the “Human Rights Impact Assessment ” of the Danish Institute for Human Rights” ⁷. Each of the aforementioned documents contributes a specific layer that, together, allows for the structuring of a methodological architecture that integrates ethical and regulatory considerations into operational criteria applicable to decision-making in artificial intelligence projects in health.

From a conceptual standpoint, the guide starts from a central premise: ethics constitutes the normative foundation upon which regulatory obligations are built. In this sense, the Ethical Impact UNESCO's Assessment offers a systematic framework for identifying, analyzing, and evaluating the ethical implications throughout the entire lifecycle of AI

¹ <https://artificialintelligenceact.eu/article/3/#:~:text=A%20provider%20is%20a%20person,%2C%20operator%2C%20and%20authorised%20representative> .

² <https://artificialintelligenceact.eu/es/what-the-act-means-for-staffing-businesses/>

³ <https://clias.iecs.org.ar/publicaciones/marcos-de-evaluacion-practicos-para-la-inteligencia-artificial-responsable/>

⁴ <https://www.who.int/publications/i/item/9789240029200>

⁵ <https://www.unesco.org/en/articles/ethical-impact-assessment-tool-recommendation-ethics-artificial-intelligence>

⁶ <https://www.paho.org/en/documents/artificial-intelligence-public-health-readiness-assessment-toolkit>

⁷ <https://www.humanrights.dk/tools/human-rights-impact-assessment-guidance-toolbox/introduction-human-rights-impact-assessment>



systems, incorporating a proactive rather than merely reactive approach. This approach recognizes the non-neutral nature of technical decisions, as choices regarding data, models, and optimization goals incorporate normative assumptions that must be made explicit and subjected to scrutiny. In the health sector, this requirement takes on particular significance, given that AI systems directly or indirectly influence clinical decisions, resource allocation, and access to services. Consequently, the ethical dimension must be operationalized from the early design stages and not relegated to later validation phases. CLIAS Technical Paper 8 and the PAHO/WHO framework complement this dimension by incorporating specific criteria for AI governance in health, such as transparency, explainability, safety, meaningful human oversight, and inclusion, which allow these principles to be translated into concrete requirements for design, development, and deployment.

In other words, **the ethical dimension must be interwoven into every phase of AI development**, from the question that leads to the conception of the project, through decision-making, handling of training data and validation, to its implementation and subsequent monitoring.

In parallel, **the regulatory dimension is structured on the basis of human rights as a minimum normative standard**, particularly relevant in contexts where specific regulations on AI are incipient or nonexistent. In Latin America and the Caribbean, this absence of specialized regulatory frameworks does not imply a normative vacuum, but rather a reconfiguration of the analysis toward pre-existing legal sources with constitutional and international standing. Rights such as the right to health, to equality and non-discrimination, to privacy, etc., constitute enforceable legal parameters that define the scope of action of AI systems in healthcare. The Human Rights Impact Assessment (HRIA) is used in this guide as a methodological framework to structure the analysis from a human rights perspective. As mentioned, in a context like Latin America and the Caribbean, where specific regulations on AI are still limited, human rights serve as the minimum normative framework that guides what is acceptable and what is not in the use of these technologies. In this sense, the HRIA is not incorporated as a tool to be applied in its entirety, but rather as a conceptual basis that allows for structuring the analysis and ensuring that decisions are made considering their impact on people, especially those in the most vulnerable situations.

3.1 WHY ETHICS FIRST

Any project related to healthcare demands particularly careful ethical consideration, because it involves not only technical or administrative processes, but also people, care relationships, and decisions that can influence the life, integrity, autonomy, and well-being of those receiving care. When AI is incorporated into this field, ethical questions become even more relevant, since these systems can participate, directly or indirectly, in clinical decisions, resource allocation, access to services, and how patients are assessed, prioritized, or supported within the healthcare system.

Ethics here serves a practical function: providing a framework for reflection to guide decisions surrounding the design, development, implementation, and use of AI systems in healthcare. Its purpose is not to provide automatic answers or replace human judgment, but rather to help identify the values at stake, the potential tensions, and the aspects that



require deliberation before proceeding. In this context, the **ethical question is not limited to whether an AI solution can be developed, but extends to understanding what need it seeks to address, what conception of healthcare it presupposes, who benefits, who might be excluded, and how it alters the relationships between patients, professionals, and institutions** .

Therefore, the ethical dimension cannot be incorporated only at the end of the process as a post-processing review or a formal validation mechanism. It must be present from the initial stages and accompany the entire system lifecycle. In AI applied to healthcare, technical decisions are never completely neutral: the problem definition, data selection, training criteria, performance metrics, integration into the clinical workflow, and the objectives pursued all reflect specific priorities and can have different consequences for different groups of people.

From this perspective, ethics allows us to analyze whether a solution effectively addresses a relevant need or is merely an available technological option. It also allows us to reflect on issues related to patient autonomy, privacy, equity, transparency, human oversight, and the impact these tools may have on the dynamics of care and professional practice. Similarly, it compels us to consider whether the systems adequately represent the populations for whom they will be used and whether they could reproduce or amplify pre-existing inequalities.

In the field of health, moreover, ethical evaluation cannot be limited to the technical performance of the system. A model may show good results in controlled environments and yet prove inadequate in real-world care settings if it fails to consider the cultural, social, or institutional characteristics of the population in which it will be used, if it hinders clinical practice, if it generates excessive dependence, or if it introduces additional barriers for certain groups.

Ethics, then, functions as a deliberative tool aimed at integrating technology in a way that is compatible with the goals of healthcare. Its role is to help build systems that are not only technically efficient, but also compatible with human dignity, respectful care relationships, and a fairer distribution of the benefits and burdens associated with the use of AI.

In this guide, the ethical foundation is distinguished from the regulatory foundation because they operate on different, though complementary, levels. **Ethics deals with reflecting on the values, purposes, and decisions that should guide the development and use of these technologies, even in situations where specific legal rules do not yet exist. Regulation, on the other hand, establishes obligations, limits, and enforceable legal mechanisms within a given legal framework** . This distinction is especially relevant in Latin America and the Caribbean, where specific regulatory frameworks for AI in health are still fragmented or incipient. In these contexts, ethical reflection remains essential to guide prudent and justifiable decisions, even when regulation has not yet developed specific responses for every possible situation.

**CENTRAL
PRINCIPLE**

Ethics cannot be added at the end as a formal review. It must be present from the beginning and accompany the entire life cycle of the system. This is what is called ethics. by design or embedded ethics.



3.2 HUMAN RIGHTS AS A COMMON LEGAL BASIS

In Latin America and the Caribbean, the development of specific legal frameworks for artificial intelligence is still in its early stages and heterogeneous. Unlike other regions, such as the European Union, the region currently lacks a common and uniform instrument that establishes a comprehensive regulatory basis for AI. Some countries have passed general laws, such as Peru ⁸ and El Salvador ⁹, while others, such as Chile ¹⁰ and Brazil, have legislative projects under discussion ¹¹. In other cases, such as Argentina ¹², there are various initiatives at different stages of development. However, the fact remains that all these instruments have different scopes, approaches, and levels of development, and most were not specifically designed for the health sector.

Given this scenario, and in the absence of a unified regional framework, this guide begins with the need to identify a common legal foundation that allows for the construction of a shared minimum standard for the analysis of AI systems in healthcare. To this end, human **rights are taken as the starting point**, since they constitute the common legal language of the region and enjoy constitutional and conventional recognition in most Latin American and Caribbean countries ¹³. In other words, **even in the absence of a specific law on AI, this does not imply an absence of rights. AI systems continue to be integrated into legal systems that already recognize and protect fundamental human rights.**

In this sense, the law does not start from a vacuum, but from prior and superior principles that mandate the protection of the human person even in the face of new technologies or those not expressly foreseen by the legislator. Therefore, any AI system applied to healthcare must be analyzed, at a minimum, in light of these fundamental rights.

The human rights already recognized in the constitutional and conventional legal frameworks of the region offer a useful interpretive framework for analyzing the impact of AI on health, even though many of these instruments were drafted in contexts prior to the development of these technologies. In this sense, the **right to health** ¹⁴ can guide the evaluation of whether these tools effectively contribute to improving access, quality, and continuity of care. The right **to equality** ¹⁵ and non-discrimination allows for an analysis of whether the systems could reproduce or amplify biases that disproportionately affect certain groups. The **right to privacy** ¹⁶ and the protection of private life offer criteria for defining the legitimate use of sensitive information, especially when it comes to health data. Similarly, the **right to information** ¹⁷ can serve as a basis for reflecting on the need for people to understand when AI systems are involved in processes related to their care and to be able to participate in an informed manner in decisions related to their use. Thus,

⁸ <https://www.gob.pe/110224-conoce-el-reglamento-de-la-ley-de-inteligencia-artificial-en-el-peru>

⁹ <https://www.asamblea.gob.sv/leyes-y-decretos/view/6137>

¹⁰ <https://www.senado.cl/comunicaciones/noticias/proyecto-que-regula-sistemas-de-inteligencia-artificial-sera-estudiado-por>

¹¹ <https://institutoautor.org/el-senado-de-brasil-aprueba-el-proyecto-de-ley-2338-2023-sobre-el-uso-de-la-inteligencia-artificial/>

¹² <https://cytdashboard.vercel.app/dashboard.html>

¹³ <https://www.un.org/es/about-us/udhr/history-of-the-declaration>

¹⁴ <https://www.ohchr.org/es/health/right-health-key-aspects-and-common-misconceptions>

¹⁵ <https://www.un.org/es/about-us/universal-declaration-of-human-rights>

¹⁶ <https://news.un.org/es/story/2018/11/1446671>

¹⁷ <https://www.unesco.org/es/right-information>



although these rights were not originally conceived to regulate AI, they continue to provide relevant legal principles for interpreting and guiding the use of these technologies in the health field.

In addition to this common framework, there are the applicable national regulations in each jurisdiction. Specifically, personal data protection laws regulate the processing of information, including health data, and establish specific obligations for those who use it. Furthermore, some countries, such as Argentina¹⁸, have specific patient rights frameworks that reinforce obligations regarding information, consent, confidentiality, and quality of care, and these remain fully applicable when AI systems are involved. Similarly, regulations concerning research involving human subjects must be considered¹⁹, especially when systems are trained, validated, or adjusted using data, records, or information from patients.

Finally, health regulations related to software as a medical device (SaMD) must be considered²⁰. Various regulatory agencies in the region, such as ANMAT²¹ in Argentina and ANVISA²² in Brazil, have developed frameworks for the evaluation, authorization, and monitoring of these technologies. When an AI system performs clinical functions, it may fall under these regulations, which entails specific validation, registration, and monitoring requirements. This stems from a fundamental legal principle: **the greater a technology's capacity to impact human health, the greater the need for control, oversight, and public guarantees regarding its operation.**

Ultimately, the main challenge posed by AI in healthcare in Latin America and the Caribbean is not the complete absence of regulations, but something more complex: understanding how existing legal frameworks should be reinterpreted in the face of technologies capable of intervening in decisions that historically belonged exclusively to human judgment. AI does not appear in a vacuum. It enters legal systems built over decades around the protection of the individual, human dignity, autonomy, privacy, and the right to health. Therefore, even when specific laws on AI are still limited, AI systems are already subject to principles, rights, and obligations that remain in force. The real challenge lies not only in creating new rules, but in deciding how to preserve what societies consider valuable when decisions begin to be shared with, or influenced by, systems capable of processing information, classifying people, and guiding behavior on an unprecedented scale.

4. How to use this guide

The guide combines two types of questions:

¹⁸ <https://servicios.infoleg.gob.ar/infolegInternet/anexos/160000-164999/160432/norma.htm>

¹⁹ <https://www.wma.net/policies-post/wma-declaration-of-helsinki/>

²⁰ <https://www.imdrf.org/working-groups/software-medical-device-samd>

²¹ <https://www.argentina.gob.ar/anmat>

²² <https://antigo.anvisa.gov.br/en/english>

- **General questions:** applicable at any stage and by any actor. They address ethical and regulatory foundations that are transversal to the system's life cycle.
- **Specific questions:** organized by profile, developers and implementers, and geared towards the specific decisions of each moment.

Both types can be viewed independently. It is not necessary to complete one section to advance to the next.

For each question, the team uses a traffic light system: Green if the answer is satisfactory, Yellow if it is partial or uncertain, Red if there is no answer or the risk is not managed. The yellow and red items become the Care Action Plan.

This guide is designed to be reviewed periodically to assess whether the Care Plans are providing answers to those responses that were marked yellow or red, and to evaluate whether they can be considered green or, conversely, whether due to regulatory changes or changes in the ecosystem itself, green actions now need to be examined in more detail.

METHODOLOGY

This guide is designed to be reviewed periodically to assess whether the Care Plans are providing answers to those responses that were marked yellow or red, and to evaluate whether they can be considered green or, conversely, whether due to regulatory changes or changes in the ecosystem itself, green actions now need to be examined in more detail.

5. Traffic light system

The operational traffic light system allows visualization of the level of preparedness or uncertainty associated with each critical question. Its purpose is not to create an illusion of absolute certainty, but rather to identify which aspects are resolved, which require attention, and which represent conditions for pause or halting.

STATE	MEANING	OPERATIONAL INTERPRETATION	ACTION REQUIRED
	Clear and satisfactory answer	The organization has adequate evidence or mechanisms.	Document and maintain.
	Partial or uncertain answer	There are doubts, insufficient evidence, or relevant unresolved limitations.	Generate responsible care action with a deadline.



No response or unmanaged risk

There is not enough information or the risk has no mitigation measure.

Resolve before proceeding or define pause condition.

CRITERION

The existence of unanswered questions does not imply that the project is poorly designed. It implies that these points represent risks that must be acknowledged, mitigated, and addressed through a care plan. A project can move forward with gray areas, provided they are identified, documented, and subject to continuous review.

6. Implementer Template

PROFILE: health authorities, health institutions, funders, etc.

FOCUS: deciding on adoption, establishing conditions of use, and building governance.

SUBGROUP	KEY DECISION QUESTION	WHAT TO EVALUATE?
NEED AND VALUE	Does the AI system solve a real and priority healthcare problem in our specific context?	<i>Clinical and institutional justification</i>
	Does AI add value compared to simpler, more accessible, or less risky alternatives?	<i>Documented comparative analysis</i>
EVIDENCE AND SECURITY	Is there sufficient evidence of performance in contexts similar to ours?	<i>Validation studies available</i>
	Do we clearly understand the limitations and situations where the system should not be used?	<i>Documented negative scope</i>
EQUITY AND IMPACT	Could the system harm or exclude certain groups of the population we serve?	<i>Subgroup bias analysis</i>
	Was it validated in populations comparable to ours (territory, language, socioeconomic profile)?	<i>Representativeness of the validation</i>
GOVERNANCE	Is it clear who oversees the system, who responds to incidents, and who can	<i>Defined responsibility structure</i>



DATA AND RIGHTS	suspend its use?	
	Does the institution have the real capacity—people, processes, systems—to monitor the system over time?	<i>Operational supervisory capacity</i>
	Does the system's use of data comply with applicable data protection and privacy regulations?	<i>Regulatory compliance verified</i>
	Will patients and professionals understand the role of AI in care and be able to question its recommendations?	<i>Transparency and information</i>

7. Developer Template

PROFILE: suppliers, startups, technical development teams, etc.

FOCUS: demonstrate evidence, traceability, limits and conditions for responsible implementation.

SUBGROUP	KEY DECISION QUESTION	WHAT TO EVALUATE?
PURPOSE AND NEED	Does the problem we're trying to solve really require AI, or can it be addressed with simpler solutions?	<i>Justification of the technological choice</i>
	Does the expected benefit for patients and professionals justify the potential risks introduced by the system?	<i>documented risk-benefit balance</i>
DATA AND BIASES	Do the training data adequately represent the population in which the system will be used?	<i>Representativeness and documented sources</i>
	Do we assess biases and performance differences between relevant subgroups (age, gender, territory, ethnicity)?	<i>Equity analysis incorporated</i>
VALIDATION AND SECURITY	Was the system validated with data independent of the training data and in scenarios representative of real use?	<i>Documented external validation</i>



	Do we know precisely in what contexts the system may fail, lose accuracy, or produce biased results?	<i>Identified failure conditions</i>
TRANSPARENCY	Does the technical documentation allow a third party to understand how the system was developed and under what conditions it should be used?	<i>Document traceability and legibility</i>
	Can the system be audited, monitored, and versioned continuously throughout its lifecycle?	<i>Auditability and version control</i>
RESPONSIBLE IMPLEMENTATION	Does the system significantly preserve human oversight and not unduly replace professional judgment?	<i>Explicit human control in design</i>
	Is there a concrete plan for reporting incidents, correcting problems, and updating or retiring the system if necessary?	<i>Incident management and lifecycle</i>

8. Care action plan

In AI projects in healthcare, identifying ethical, regulatory, or technical risks is only one part of the process. The real challenge arises when these observations must be translated into concrete decisions about how to design, implement, monitor, or even limit the use of a system. In practice, there isn't always an automatic link between recognizing a problem and having clear mechanisms to address it. Therefore, many of the issues identified during the initial assessment subsequently require specific monitoring, review, and care to ensure they can be properly managed throughout the technology's lifecycle.

Therefore, the questions posed in this guide should not be understood merely as an evaluation or compliance exercise. Their **true value emerges when they force the team to make operational decisions about how to address the uncertainties, limitations, or tensions that the project itself identifies**. In healthcare, recognizing a risk and failing to act on it is rarely neutral. Often, the most significant problems arise not from what was

unknown, but from what was known but not accompanied by adequate care, supervision, or follow-up.

In this sense, every AI project in healthcare should translate its ethical and regulatory findings into an explicit plan of care. The logic is not to create an illusion of absolute safety, something impossible in complex systems, but to recognize that all implementation occurs in real-world contexts, with real limitations, and that precisely for this reason it needs mechanisms for observation, correction, and continuous learning.

A care action plan transforms abstract concerns into concrete decisions. It requires defining which aspects need active monitoring, which situations demand enhanced human supervision, what uncertainties still exist, and what conditions should be met to move forward in a reasonable and proportionate manner.

In practical terms, the plan should be built upon the questions that remained open during the ethical or regulatory assessment. For each identified risk, uncertainty, or strain, the team should be able to define at least the following:

For each yellow or red item on the template, the team must define at least:

- Which dimension requires special care: security, equity, privacy, human oversight, governance, among others.
- Who could be affected: patients, professionals, institutions, vulnerable groups.
- What specific consequence are you trying to avoid or reduce?
- What measure will be implemented: monitoring, mandatory human review, additional validation, auditing, usage restrictions, specific consent, or incident reporting.
- Who is responsible for overseeing that measure?
- At what point should it be reviewed?
- What indicator or threshold will allow us to assess whether the actions are working?

PURPOSE

The care action plan is not an administrative appendix. It is the practical way to translate ethical and regulatory principles into concrete implementation decisions. It does not stifle innovation; rather, it helps ensure that technology remains aligned with its purpose in healthcare.

The basic template is presented below. It is recommended that you complete it with your own items identified during the evaluation using the template:

DIMENSION / RISK	STATE	ACTION REQUIRED	RESPONSIBLE	TERM	TERM
------------------	-------	-----------------	-------------	------	------



Validation in local context	●	Clear and satisfactory answer	The organization has adequate evidence or mechanisms.	6 months (estimated)	Sensitivity $\geq 90\%$ in local population
Underrepresentation bias		Evaluate differential performance by subgroups.	Data Team	3 months (estimated)	Documented equity report
Non-existent incident protocol	●	Create an institutional incident protocol.	Quality and patient safety	60 days (estimated)	Protocol approved by management
(Complete with your own items)					

9. General Question Bank

The following questions are relevant throughout the system lifecycle. It is recommended to review them during the design phase, before deployment, and periodically during operation. They complement the questions in the Business Model Canvas and are available as an expanded reference.

9.1 ETHICS QUESTIONS

Purpose and meaning of the system

- What specific need does this AI system seek to address?
- Is the problem being addressed truly a priority in the context where it will be used?
- Does the incorporation of AI respond to a real need or mainly to the availability of the technology?



- What value does the system add for patients, professionals, or institutions?
- Could the solution significantly change care relationships or clinical practice?

People and groups affected

- Who could benefit and who could be excluded?
- Were the cultural, social, and territorial characteristics of the population taken into account?
- Could the system create additional barriers for certain groups?
- Was an assessment made of how people in the most vulnerable situations might be affected?
- Did the potentially affected individuals have any opportunity to express concerns or needs?

Equity and bias

- Could the system reproduce existing inequalities in the health system?
- Were potential biases analyzed based on age, gender, disability, language, socioeconomic level, or territory?
- Could expected performance vary between different groups of people?
- Are there mechanisms in place to detect and review unequal impacts once the system is implemented?

Autonomy and care relationships

- Does the use of the system preserve a significant space for human professional judgment?
- Could AI unduly influence clinical decisions or create automatic dependency?
- Will people understand when an AI system intervenes in processes related to their care?
- Are there mechanisms to question, review, or deviate from the system's recommendations?

Transparency

- Can the system be reasonably understood by those who use it or are affected by it?
- Are its limitations, conditions of use, and margins of uncertainty clear?



- Does the system adequately communicate what type of support it provides and what decisions still depend on human judgment?

Context and sustainability

- Was the system designed taking into account the actual conditions of infrastructure, connectivity, and resources?
- Does the institution have sufficient capabilities to monitor the use of the system securely?
- Is there a provision in place to periodically review whether the solution remains adequate?

9.2 REGULATORY QUESTIONS

Applicable regulatory framework

- Did they identify which regulations might apply to the system in the country where it will be developed, implemented, or used?
- Were frameworks related to data protection, health regulations, patient rights, research in human subjects, or software as a medical device reviewed?
- Were the applicable legal obligations analyzed even though there is no specific law on AI?

Authorization and supervision

- Did they verify if the system requires authorization, registration, or evaluation by a health authority?
- Could the system be subject to regulations related to software as a medical device?
- Did you check if the regulatory authorities offer applicable guidance or consultation channels?

Data protection and privacy

- Does the use of the system comply with applicable regulations on privacy and personal data protection?
- Are adequate security, anonymization, or pseudonymization measures in place?
- Is it clear what purpose the data will be used for and how long it will be kept?
- Was it assessed whether the data could be reused, transferred, or used to retrain future systems?



- Was the person properly informed, in clear and understandable language, about the use of their data in the AI system, including the purpose of the processing, the responsible parties, possible secondary uses, relevant risks and their rights?
- When data processing requires informed consent, was this obtained in a prior, free, specific, express and documented manner?

Assignment of roles and functions

- Are the roles of the people, teams, or institutions involved clearly defined?
- Is it defined who manages the data and who makes decisions about its use?
- Are there clearly identified individuals responsible for responding to incidents or problems?

Continuous monitoring and review

- Was it defined how the system's regulatory compliance will be reviewed over time?
- Are there procedures in place for reviewing incidents, complaints, or questions related to the system?
- Does the organization have the capacity to suspend or modify the use of the system if relevant risks appear?

10. Conclusion

AI in healthcare presents us with a paradox: **the more powerful the technology becomes, the more human the important questions become.**

Algorithms can predict, classify, and automate. But they cannot decide which lives should be prioritized, what risks a society is willing to accept, which inequalities deserve to be corrected, or what it means to care for a sick person. Those decisions do not belong to machines. They belong to institutions, professionals, communities, and the values we choose to protect.

The greatest risk is not just using imperfect systems, but believing that all innovation is, in and of itself, progress. In healthcare, an efficient tool can be unfair. A statistically successful model can fail in practice. A technology designed to optimize processes can weaken the human connections essential to care.



CLIAS

CENTRO DE INTELIGENCIA
ARTIFICIAL Y SALUD
PARA AMÉRICA LATINA
Y EL CARIBE