

Appendix

Prepared By: Mattea Welch, Benjamin Grant, and Christopher Deutschman
In Consultation With: Clare McElcheran, Adam Badzynski, Jennifer A.H. Bell, Andrew Hope, Robert C. Grant, Tran Truong, Kelly Lane, Patti Leake, Divya Sharma, Ian Stedman, Laleh Seyyed-Kalantari, Mike Lovas, Jeremy Petch, Benjamin Haibe-Kains, and James A. Anderson

Table of Contents

Table of Contents	1
1. General.....	4
1.1. Fairness Definitions	4
Bias	4
Inequity.....	4
Responsibility.....	4
1.2. Common Methodological Biases.....	5
1.3. Financial Considerations	6
2. Problem Identification and Study Design	7
2.1. Define the problem and the solution.....	7
Assembling your cross-functional team.....	7
Patient engagement and participation.....	7
Non-technical solutions	7
2.2. Focus on Equity in the Problem Space	8
Explore and document baseline inequities	8
Literature review (social determinants of health, sources of bias and inequity).....	8
Consult with knowledgeable partners.....	9
Working with First Nations, Indigenous, and Métis Research and Data.....	9
2.3. Outcome Measurements and Data Requirements	10
Assess AI Output Integration into Clinic	10
Outcome Measurement.....	10
Data requirements and availability	11
3. Model Training and Development	12

3.2. Appropriateness of Retrospective Data.....	12
Identify current clinical benchmarks	12
Assessing dataset for disparities and ‘fit for purpose’	12
3.3. Defining Objectives and Metrics	12
Appropriate statistical measures of performance.....	12
Defining Equity Objectives.....	13
Measures of AI-Solution Fairness	15
Fairness metrics	15
Open-Source fairness metric libraries.....	16
3.4. Model Training and Testing.....	16
Acquisition of a third-party model.....	16
Data Representativeness and Quality	18
Mitigation methodologies.....	18
Data pre-processing	18
ML training methods	19
Prediction post-processing	20
4. Silent Deployment and Clinical Evaluation	21
4.1. Prospective Deployment Preparation	21
Approvals for prospective deployment	21
Retrospective to prospective mapping	21
Prospective Data Representativeness	22
AI-Solution Failures and Adverse Events	22
4.2. Prospective Model Assessment (Silent Mode)	23
Silent Deployment Considerations	23
Many approaches to silent mode.....	23
Threshold selection.....	23
Assessing Silent Mode Performance	23
Define auditing methods.....	23
4.3. Prospective Clinical Evaluation	24

Education materials.....	24
Recruit end-users for future clinical integration.....	25
5. Operationalization and Lifecycle Monitoring	26
5.1 Preparation and Documentation	26
Comprehensive documentation	26
Internal Documentation:	26
1. Model Development and Technical Specifications.....	26
2. Deployment Documentation	27
3. Regulatory and Compliance Documentation	27
4. Monitoring and Post-Deployment Validation.....	27
5.2. Communication and Education.....	28
Modifying trial educational materials	28
Communication plan	29
Communicating to Patients and Care Partners	30
5.3. AI-Solution Rollout and Monitoring	30
Regular monitoring and reporting.....	30
Gathering Feedback for Continuous Quality Improvement.....	31
5.4. Updating or Decommissioning of AI-Solution	31
6. Community Tools	33
References	34

1. General

1.1. Fairness Definitions

While ‘bias’, ‘inequity’, and ‘responsibility’ have myriad meanings, for the purposes of this document, the terms will be defined as follows:

Bias

An umbrella term for both social biases and methodological biases. Social bias, which can be conscious or unconscious, can be defined as “discrimination for, or against, a person or group, or a set of ideas or beliefs, in a way that is prejudicial”.¹ Methodological bias can be defined as “systematic errors stemming from choices made during the research process”², such as selection bias, observer bias, response bias, and publication bias.³ In the machine learning sphere, there is often an intersection between social and methodological biases (for example, data based on nursing assessments may have a methodological observer bias, where part of the inter-observer variability is due to each nurse’s individual social biases). Within the context of AI, a biased solution is one that has different performance for different subgroups of patients or inequitable impact on different subgroups of patients.

Inequity

A term defined broadly as “unjust differences between populations in the access, use, quality, and outcomes of care.”⁴ The key difference between ‘inequities’ and ‘disparities’ is that ‘inequity’ is a value-laden term – the differences described are unjust or unfair – whereas ‘disparity’ merely describes that a difference is present. Serious scrutiny is required to determine in which cases a ‘disparity’ should or should not be considered an ‘inequity’. Inequities are *systematic* differences in the opportunities groups have to achieve optimal health.⁵ These systematic differences are influenced by social biases and by structural differential access opportunities.

Responsibility

A ‘responsible’ AI solution is one that is applied in a way that is the most appropriate (just, ethical) - but what is most appropriate for a given AI solution in a given situation is not static. The proceeding framework is adaptable to most definitions of ‘responsibility’ that could be chosen. An AI solution that is developed and deployed following all of the steps provided in this framework should implicitly follow, at minimum, the principals of equity, fairness, scientific rigor, and democratic engagement, but other values will be infused by choices made throughout the AI lifecycle (such as which patient populations are served or prioritized, how fairness is defined and to what extent the definition is enforced, how information about the AI is disseminated to end users, and so on). Identifying which values are important to the project at the outset is critical to making these decisions downstream.

An AI solution could be considered responsible if at the end, it functions in alignment with the core principles of clinical ethics⁶ (beneficence, non-maleficence, autonomy, and justice); or if it meets the standards of a pre-existing set of ethical standards such as those put forth by the Coalition for Health AI⁷ (usefulness, fairness, safety, transparency, and privacy). Alternatively, the guiding definition for responsibility for your solution may be that the AI contributes to the common good⁸. However, when using a broader philosophical definition of responsibility, care must be taken to operationalize it appropriately. For example, in the case of defining responsible AI as contributing to the common good, who is the “common” - the patients served by your institution? All residents in the city of your institution? The nation? Humanity? After defining “common”, what is the definition of “good”? Who gets to have input on that definition?

1.2. Common Methodological Biases

Below is a list of commonly found methodological biases that can be present at different times of an AI-Solutions development and deployment. This list is not meant to be a comprehensive overview of all potential biases, but to be an introduction to some of the most commonly found ones.

- 1) **Selection bias:** Occurs when the training data is not representative of the broader patient population, resulting in skewed results. This can lead to models that perform well on certain subgroups but fail on others, especially if certain patient demographics are underrepresented (e.g., minority groups).^{9,10} This can be the result of biased sampling, measurement errors, or incomplete data collection, leading to non-representative datasets. In healthcare, this can occur if electronic health records (EHRs) or datasets are incomplete or contain errors, leading to skewed models. For example, if data is collected only from a specific hospital type or region, the model may not generalize well to other settings. Missing data also causes biases when missing in a systematic way, leading to skewed results. For instance, if patients who are healthier are more likely to have missing records (i.e. less tests performed), models may underpredict positive outcomes.
- 2) **Confirmation bias:** The tendency to search for, interpret, or prioritize information that confirms one's preexisting beliefs or hypotheses, while disregarding contradictory evidence. e.g. a belief that a population of patients is not impacted by inequitable care, and therefore a thorough literature search is not conducted.^{11,12}
- 3) **Incorporation bias:** A subtype of verification bias, where the predictive model includes features or data points that are part of the outcome it is trying to predict.

For example, if an ML model predicts a diagnosis and uses test results that are part of the diagnostic process itself, it may artificially inflate model performance.¹³

- 4) **Survivorship bias:** Arises when only patients who have survived or remained in care and included in the study and those who have died or dropped out are ignored. In ML model predictions, this can lead to over-optimistic performance in predicting patient outcomes, as the model ignores negative outcomes. For qualitative research related to patient experiences, this will also, likely, result in over-optimistic evaluation of an AI-Solution.^{14,15}
- 5) **Model overfitting:** When a ML model is too complex, it may fit the noise or anomalies in the training data rather than the underlying patterns. This leads to excellent performance on the training set but poor generalization to new, unseen patient data.^{16,17}
- 6) **Spectrum effects:** Occur when a model's performance varies across different patient subgroups, such as those with different disease severities. This bias can result in a model performing well for a mild disease but poorly for severe cases, leading to misleading evaluations of model accuracy.¹⁸
- 7) **Automation bias:** This bias occurs when clinicians overly trust ML model outputs without critically evaluating them. Over-reliance on automation can lead to errors, especially if the model is flawed or the input data is incorrect. ¹⁹
- 8) **Unconscious bias:** Also known as implicit bias, unconscious bias can be introduced if the data or model development reflects the prejudices of the research team. For example, models might inadvertently favor certain patient groups over others if the underlying data contains implicit biases related to race, gender, or socioeconomic status.²⁰

1.3. Financial Considerations

Addressing bias, fairness and inequity in AI and ML problem spaces requires additional resources. It is encouraged to consider this guideline in the early stages of problem development to develop a realistic plan for incorporating responsible AI strategies into the project. Responsible AI development resources can be incorporated into grant applications through budgeting and through collaborators. In the case that additional funds are not available for addressing bias, fairness and equity in development, some resources are available through open-source community tools ([Section 6](#)).

2. Problem Identification and Study Design

2.1. Define the problem and the solution

Assembling your cross-functional team

The cross-functional team composition and involvement of specific individuals should depend on the scope, topic and the current stage of the project. The team is not a formal review committee, but one that ensures each stage of development is considered by an appropriate expert. Over the lifetime of the project, the team should, at some point, include ML researchers, clinicians, data analysts, statisticians, equity and ethics experts, patient and/or caregiver representatives, clinical champions and end-users.

Engaging a cross-functional team to support bias and fairness evaluation of a project may increase the budget of a project. It is encouraged to include resourcing for this team in grant applications. Experts in ethics, equity, ML researchers, and clinicians can be included in grants as collaborators, and additional budget can be set aside for engaging clinical and patient partners.

Patient engagement and participation

Patient voices, represented by patient partners, family and caregivers, and/or professional patient advocates, should be present within and beyond the cross-functional team. Patient partnerships are critical to make sure that patients are guiding the research in identifying unmet needs and improving outcomes for those needs. Patient voices must be present in the whole spectrum of AI research, from problem identification and study design through to clinical implementation and lifecycle monitoring.

Many institutions and disease specific advocacy groups will have their own supports for patient partners, so it is important to look for local initiatives. In Canada, the Canadian Institutes of Health Research has a framework for patient-oriented research and engagement (<https://cihr-irsc.gc.ca/e/48413.html>) - similar initiatives exist in other countries. There are a number of validated tools to assess how well patients were engaged in the research process, such as the PEIRS-22 tool^{21,22} and the REST²³ survey. These can be helpful both to validate that patient and caregiver partners have been successfully brought into the cross-functional team, and to serve as a resource for thinking about how to engage partners in the first place.

Non-technical solutions

The allure of AI in healthcare often creates a "shiny object syndrome"²⁴ where research teams and organizations rush to implement cutting-edge technology simply because it's novel and exciting, rather than because it's the most effective solution. An AI-first mindset can lead to overlooking simpler, proven interventions that might better serve patients and

healthcare workers. When enthusiasm for AI's potential overshadows careful consideration of alternatives, organizations risk investing substantial resources into complex technical solutions for problems that might be better solved through basic process improvements, better staffing, or enhanced communication systems. It's worth remembering that sometimes what appears to be a technology problem is actually a human systems problem in disguise, and no amount of sophisticated AI can fix underlying issues with workflow, staffing, or resource allocation.

When addressing healthcare challenges, it's crucial to first explore fundamental non-technical and non-AI-Solutions that may be more effective and sustainable. These approaches may require fewer resources, face fewer implementation challenges, and can be more easily adapted to local contexts compared to AI-Solutions.

2.2. Focus on Equity in the Problem Space

Explore and document baseline inequities

A key component of responsible AI development and deployment is mitigating bias, fairness and inequity. In order to do this, there need to be clearly defined equity objectives and fairness metrics to measure success. These definitions will be specific to the population, problem space, and solution that is being addressed by a particular AI-Solution. Exploring and understanding baseline inequities is paramount to developing a fair and equitable model, and can be done in a few different ways (please see sections below). Taken together, these learnings lay the foundation for quantitative analysis of biases and inequities within your target population once you have access to local retrospective data to explore in [Section 3.2](#).

Literature review (social determinants of health, sources of bias and inequity)

The equitable and compassionate components of this framework rely on an extensive review of documented biases in healthcare literature. However, it should be acknowledged that biases identified in the literature are not guaranteed to be exhaustive, nor always fully representative in the collected data or healthcare system, particularly when the evolution of protected attributes over time is considered (e.g. newly adopted gender definitions that are more granular and representative, or new discoveries in a disease definition that increases specificity). Key patient factors associated with healthcare inequity include race^{25–29}, age^{30–33}, sex^{34–36}, geographical location of residence^{37,38}, patient support systems^{39,40}, primary language, income level, and other social determinants of health. Each of these elements independently may impact healthcare access, treatment, and/or outcomes, and often the effects compound when elements are present together, creating unique barriers for those with intersectional identities. Common intersections may include age/sex, and ethnicity/sex. Accounting for intersectional identities and intersections of other health-related groupings remains

challenging⁴¹; there are far more combinations of intersecting demographic factors than can be practically addressed. Broadening bias assessment to include intersectional identities is essential, but often constrained by data availability.

Consult with knowledgeable partners

There will be gaps in the healthcare literature regarding what biases exist in a given problem space. By consulting with individuals who understand the patient and caregiver experience in the specific problem space, as recommended in our framework, the team may begin to fill in these gaps. There are numerous stakeholders that can inform the team's understanding of these biases, including patients, patient partners, families/caregivers, patient navigators, and healthcare providers (see [Section 2.1 - Assembling your cross-functional team](#)). The goal of this consultation process is to identify any biases or inequities, real or perceived, that must be considered and either successfully mitigated or identified as a limitation to the work.

Working with First Nations, Indigenous, and Métis Research and Data

For AI-Solutions that are focused on First Nations, Indigenous, and/or Métis populations, or are known to utilize First Nations, Indigenous, or Métis data, special care should be taken to respect the data sovereignty of these groups. If the AI-Solution will use data from specific individual nations, the project team must seek approval from appropriate leaders from those nations prior to commencing the project.

To bolster the researcher's own capacity to respect/assert the principles of data sovereignty, it is recommended that all members of the research team complete The Fundamentals of OCAP® course that was developed by the First Nations Information Governance Centre. OCAP stands for Ownership, Control, Access and Possession, and the course provides valuable information on how to appropriately engage with these communities and their data.

Through this course, participants will learn about the history and motivation behind OCAP, as well as receive practical steps for participating in, or seeking to participate in relevant research studies. As examples, individuals should ask themselves the following questions, among others that are outlined in the course:

1. Was this data collected with the approval and knowledge of the First Nations?
2. Does the project align with the priorities of the First Nation(s) from whom the data is coming?
3. Could undertaking this project/using this data cause harm to the First Nation(s) or its members?
4. Are First Nation(s) members being included in the project from conception to analysis to implementation?

2.3. Outcome Measurements and Data Requirements

Assess AI Output Integration into Clinic

Prior to embarking on the lengthy and costly endeavor of AI-Solution development, testing and integration ensure that the solution will have a measurable impact on clinical processes. Two important questions to ask are: (1) What will be done differently in the clinic, either operationally or for patients, with the AI-Solution outputs; and (2) do pathways exist to act on AI-Solution outputs.

The clinical impact of an AI-Solution's output may be limited if systemic barriers exist that prevent recommended actions. These barriers may include lack of available interventions (e.g., an output recommending the patient switches treatment lines, when there are no treatments remaining that the patient is naïve to), resource limitations within the institution (e.g., an output that recommends referral to palliative care when the department does not have capacity for more patients), or resource limitations within the broader health system (e.g., an output that recommends discharging patients to long term care centres when these facilities do not have capacity). The utility of AI-Solutions in such situations must be thoughtfully considered.

Outcome Measurement

Defining the ideal outcome and determining what should be measured to assess that outcome are two separate steps. First, the ideal outcome should be defined based on the ideal clinical state after AI-Solution integration. For example, for an AI-Solution meant to predict and intervene in post-surgical pain crisis, the ideal clinical state post-deployment might be 'more patients with appropriately managed pain'. After the ideal clinical state has been defined, the most appropriate outcome measurement(s) can be identified. In this example, measurement may be simple, such as self-reported patient pain or pre- and post-deployment number of analgesic orders per patient.

The challenge lies in making sure that the measurement(s) chosen is/are fair and equally assessable across all affected subpopulations. Patients of colour are often prescribed pain medication at a lower rate than white patients, despite these patients being in equal or higher levels of pain. When they are prescribed analgesics, patients of colour are also

more likely to be given a lower dose of the medication or given a different regiment entirely. If the AI-Solution is making predictions of the likelihood of pain crisis based in part on historic analgesic orders, the prediction will not be appropriately sensitive for patients of colour.

In this example, measuring outcomes such as ‘rate of analgesics ordered across target subpopulations’ or ‘average patient self-reported pain pre- and post-deployment per target subpopulation’ may be more appropriate to capture how the AI-Solution is affecting different groups. Without taking care to do measurement by subpopulation, there may still be an overall increase in orders or overall decrease in reported pain after surgery, but those trends may be primarily influenced by the experiences of white patients, and the unequal increase in orders/dosing would lead to an overall inequitable clinical state.

Data requirements and availability

Before moving on to model training and development, it is important to consider the data that is required to allow your model to perform effectively and equitably. This includes determining specific outcome measures, the features that need to be included in the model, and whether they are collected in a format, and for a purpose, that is adequate for your specific problem and population. Many clinicians, researchers, and AI specialists are not particularly familiar with the sources, types, structures, and availability of data, so it is worth doing this investigation up front to avoid significant problems at later stages. It is also important to consider whether PHI is required and if so, for exactly what purpose, or whether the data needs to be de-identified or anonymized ([Section 3.1](#)).

It is also important when identifying data elements and data sources to understand the context in which the data was gathered. How and why the data was collected, and the intended purpose/user of the data, is important information when assessing potential biases and limitations with the data.

3. Model Training and Development

3.2. Appropriateness of Retrospective Data

Identify current clinical benchmarks

Clinical benchmarks are essential for evaluating new machine learning models in healthcare because they ensure relevance and impact. By setting a comparison standard grounded in real-world clinical practice, these benchmarks help establish whether a model provides tangible benefits over existing methods. This relevance ensures models address actual clinical needs rather than producing technically interesting but clinically irrelevant results.

Furthermore, benchmarks foster trust and safety. They align models with regulatory standards, reducing risks and proving the model's utility before deployment. Comparative benchmarking also helps stakeholders gauge the model's effectiveness against current methods, ensuring that any advancements are not just statistically significant but clinically meaningful, thereby justifying the model's use in healthcare settings.

Assessing dataset for disparities and 'fit for purpose'

Biases that should be considered during this phase of the project include sampling, convergence and participation bias. Additionally, using known societal inequities and biases that were identified during the literature search of the problem space, assess your retrospective dataset to ensure it is properly representative of these known issues. This can be achieved most simply by looking at dataset distributions, and minority and majority class representations. Furthermore, for protected or minority groups, a statistical power analysis should be performed to ensure there are enough individuals in the dataset to garner meaningful information.

3.3. Defining Objectives and Metrics

Appropriate statistical measures of performance

Defining appropriate statistical measures of performance for an AI-Solution is crucial to ensuring its effectiveness and reliability. The choice of metrics should align with the specific problem domain, objectives, and nature of the dataset. For instance, in binary classification tasks, measures like AUC-ROC (Area Under the Receiver Operating Characteristic Curve)^{42,43} and PR-AUC (Precision-Recall Area Under the Curve)^{44,45} are often used. However, selecting between these metrics depends on the data characteristics and known imbalances. AUC-ROC provides an overall view of the model's ability to discriminate between classes, but it may not be as informative in cases of class imbalance. PR-AUC, on the other hand, focuses on the precision and recall trade-off, making it a better choice when the positive class is rare or when false positives and false negatives have differing consequences. Failure to select the right metric can result in a

misleading evaluation of the model's performance and potentially biased performances for the majority class.

Some metrics also have known dependencies that must be considered to ensure accurate interpretation of performance. For example, the Dice Similarity Coefficient (DICE), commonly used in medical image segmentation tasks, is influenced by the size or volume of the segmentation^{46,47}. Larger segmentations may artificially inflate the DICE score, while smaller ones might unfairly penalize it. Understanding these dependencies is essential to avoid obscuring the true performance of the model. Practitioners should complement such metrics with additional measures or normalization techniques to account for these biases and provide a more holistic evaluation.

Another key aspect is the validation strategy employed to assess the AI model. Techniques like k-fold cross-validation help to mitigate overfitting and provide a robust estimate of the model's generalizability^{48,49}. Using too few folds might lead to high variance in performance estimates, while excessively large numbers of folds can increase computational costs without significant gains in reliability. An inadequate validation process can lead to over-optimistic results on retrospective data that fail to generalize to new prospective data, undermining the solution's practical applicability.

Overfitting is another common pitfall that arises when statistical measures are not appropriately applied or understood¹⁶. Overfitting occurs when a model learns the training data too well, capturing noise instead of generalizable patterns. This often happens when performance metrics are optimized exclusively on the training data or when complex models are used without adequate regularization. For instance, reporting high accuracy on training data while neglecting poor performance on unseen data can give a false sense of success. Similarly, reliance on a single metric, such as accuracy, in imbalanced datasets can obscure significant deficiencies, such as a model's inability to correctly identify minority class instances.

Misusing statistical measures—whether by selecting inappropriate metrics, inadequately validating models, or overfitting to training data—can lead to incorrect conclusions and poor real-world performance. By thoughtfully addressing these considerations, practitioners can create AI-Solutions that are both accurate and trustworthy.

Defining Equity Objectives

While the specific equity objectives for a project are informed by the inequity(/inequities) present in the problem space, the overarching type of objective can be broadly categorized as:

- 1) Preserving the current clinical state (i.e. the AI-solution doesn't worsen inequities after it is implemented).
- 2) Improving the clinical state for all groups as equally as possible, following the principles of formal justice.
- 3) Preferentially improving the clinical state for a historically underserved group, following the principles of distributive justice.

To exemplify how an equity objective might be defined, consider the scenario presented in [Outcome Measurement](#), in which an AI-Solution is deployed to assist in managing post-surgical pain crises. In this theoretical example, there has been a historic disparity between white patients (whose pain has been treated appropriately most but not all of the time) and Black patients (who have more often been undertreated for more severe pain).

If the equity objective was to make sure that the clinical state stayed the same post AI-deployment, it might be expressed as an objective like “the rates of analgesic prescriptions for white patients and for Black patients does not decrease post-deployment”. This would ensure that the model did not worsen under-prescribing of pain medication for Black patients based on historic trends, but it may or may not close the gap between white and Black patients' pain management.

If the equity objective was to produce a clinical state where all groups are benefited as equally as possible, it may be expressed as an objective that “a higher proportion of patients receive analgesics *and* report lower pain post-deployment”. This objective captures individual benefit regardless of subgroup belonging, but it risks obscuring if the AI-solution is favoring white or Black patients unduly.

If the equity objective was to focus on benefiting the historically underserved group the most (in this case, Black patients), the objective might be expressed as “the rate of analgesics prescribed to Black patients increased, *and* self-reported pain of Black patients decreased, post-deployment”. This objective would ensure that medication is being prescribed to Black patients more often in the appropriate situations, but it would not specifically measure the effect of the AI-solution on white patients.

In addition to choosing a type of equity objective, it may be valuable for teams to choose a magnitude of impact to attach to their objective. For example, instead of phrasing their objective as “a higher proportion of patients receive analgesics *and* report lower pain post-deployment”, they may want to state their objective as “10% more patients receive analgesics, *and* the average patient treated with analgesics reports 20% lower pain post-deployment”. Having these concrete magnitudes attached to the objectives allows for a

clearer delineation of when the AI-solution is succeeding in deployment, and when it is in need of updating or decommissioning.

Decisions on what type of equity objectives to choose, and what (if any) magnitude to assign to them, should be made in consultation with the entire cross-functional team. While patient and caregiver partners should be involved in the entire lifecycle of the project, it is particularly important to engage them in deciding what the equity objectives are, as this measure is directly related to the magnitude of clinical impact the AI-solution will have.

Measures of AI-Solution Fairness

Selecting appropriate measures of fairness is imperative to the assessment of developed AI-Solutions. A few example measures are highlighted below, but investigators are encouraged to do a review of the literature to see if any newer and more relevant measures have been developed.^{50–53}

Fairness metrics

Designed to analytically assess equality and equity issues in order to give insight into the nature of the model's performance.

- 1) Group Unawareness: Group unawareness means that a model does not use a sensitive variable during prediction⁵⁴. For example, in a simple linear regression, the coefficient of the sensitive variable would be set to 0. Group Unawareness can be assessed using metrics such as SHAP, or by looking at the statistical significance of a univariable Ordinary Least Squares Regression that has been fit to predict the outcome of interest can be predicted using a single feature. Caution when using Group Unawareness is needed in scenarios where the sensitive variable may be highly correlated with a proxy variable.
- 2) Statistical Parity/Demographic Parity: measurement of whether a model predicts positive outcome at equal rates for each segment of a subgroup.⁵⁵
- 3) Equal Opportunity: measures whether individuals from each segment of a protected class, who are eligible/qualified, have the same probability of receiving a positive outcome (e.g. being offered participation in a clinical trial). However, assessment of eligibility/qualification can be subjective, and Equal Opportunity does not address biases in the assessment process.^{52,56}
- 4) Equal Odds: this metric is used to quantify whether equal true positive rates and false positive rates exist between different groups. It is more restrictive than equal opportunity making it more appropriate when there are existing biases in the data. However, this is a very restrictive metric and may reduce the model's performance.^{57,58}

- 5) Positive Predicted Value (PPV) - Parity: given a positive prediction, the precision is equal across different groups. For example, if a model positively predicts that a patient should receive “treatment X”, the probability of this treatment being successful is equal in all segments of a protected class.⁵⁹
- 6) False Positive Rate (FPR) - Parity: this metric is the opposite of PPV-Parity and wants to ensure that each segment of a protected class has the same false positive rate. For example, if the model positively predicts that a patient should receive “treatment X”, the probability of this treatment not being successful is equal in all segments of the protected class.⁶⁰

Open-Source fairness metric libraries

- 1) [FairLearn](#): An open-source python package designed for metric calculation and reduction of bias in algorithms.
- 2) [Fairness Indicators](#): A package from Google designed to work with TensorFlow. Can be used for evaluation and visualization of group disparities.
- 3) [AIF360](#): An open-source package used to detect and mitigate biases. It is known for its extensive documentation and tutorials.
- 4) [Themis-ML](#): An open-source package for easy integration with scikit-learn. It is used mainly for binary classes and evaluation.

3.4. Model Training and Testing

Acquisition of a third-party model

Assessment of an AI-model, or AI-Solution, acquired from a third party for deployment in a healthcare institution is imperative for the safety of patients. Comprehensive questioning is required across multiple domains. Some questions that should be asked include:

1. General Information
 - a. What is the intended purpose and scope of the model?
 - b. Has the model been used in clinical settings similar to ours?
 - c. What are the specific clinical problems it aims to solve?
 - d. Were patients and end-users consulted during the conception and development of this model?
2. Development and Validation
 - a. What data was used to train and validate the model? (e.g., size, sources, geographic diversity, demographic representation) Are we able to do an internal assessment of the data?
 - b. What validation processes were followed? (e.g., external validation, cross-validation)

- c. What are the key performance metrics, and how do they vary across subpopulations of interest?
- 3. Bias and Fairness Assessment
 - a. Was the training data representative of the populations the model will serve? This may be challenging if the training data cannot be accessed and compared to local data.
 - b. What methods, if any, were used to detect and mitigate biases in the development phase?
 - c. Are there known performance disparities across demographic groups (e.g., age, sex, race, ethnicity)?
 - d. What fairness frameworks or metrics were used to evaluate the model?
 - e. Does the model include safeguards to minimize inequities in its recommendations?
 - f. Are there transparency mechanisms to report when biases or unfair outcomes are detected post-deployment?
- 4. Safety and Risk Management
 - a. What safeguards are in place to ensure patient safety in case of errors?
 - b. Has the model undergone stress testing for edge cases or rare clinical scenarios?
 - c. What adverse outcomes or unintended consequences were identified during testing?
 - d. Is there a protocol for monitoring and reporting errors or adverse events post-deployment?
- 5. Compliance and Regulatory Adherence
 - a. Does the model comply with relevant regulations?
 - b. Are there documented audit trails for data usage and model outputs?
- 6. Operational Integration
 - a. What are the technical requirements for integration with our existing systems (EHR, PACS, etc.)?
 - b. What resources are needed for deployment and maintenance?
 - c. What user training and support are provided?
- 7. Post-Deployment Monitoring
 - a. What tools or processes are available for continuous monitoring of model performance, clinical impact, and operational and regulatory adherence?
 - b. How frequently should the model be retrained or updated?

- c. How is feedback from clinicians and patients incorporated into updates?
 - d. Are there mechanisms to adjust or recalibrate the model based on observed disparities?
- 8. Ownership and Intellectual Property
 - a. Who owns the model and any updates made during its use?
 - b. What are the terms of data sharing, if applicable?

Data Representativeness and Quality

When the retrospective dataset is split into training and testing cohorts (through e.g. randomized stratified splitting, bootstrapping, time wise split), investigators should take care to assure each cohort retains the same representation of known biases and inequities (i.e. the distributions should be the same when compared to the full retrospective dataset).

Data quality should also be assessed at this step. Data missingness and consistency in variable naming are two such tests that should be completed prior to commencing ML training. Additionally, if using longitudinal data, an understanding of whether there were any changes in clinical practice over time is required to avoid temporal bias (e.g. variable naming standards, standard treatment regimens, pandemics affecting patient presentation).

Mitigation methodologies

Similarly to the fairness and equity measures section presented in [Section 3.3](#) of this appendix, this section is used only to highlight a few different methods to improve model fairness. Additional methods can be found in some of the open-source fairness metric libraries mentioned [above](#). A literature review is recommended if none of the methods below are appropriate, or to determine if there are newer and more relevant methods that have been developed.

Data pre-processing

In some scenarios, it is appropriate to change or adjust the dataset to be fairer before training an ML model.

- 1) Relabeling and perturbation: These methods involve changing the end-point label or features in a dataset. Relabeling attempts to balance the dataset by changing the end-point label, while perturbation involves varying the features/variables to create a more balanced representation of the data. Two examples of these methods are disparate impact remover⁶¹ and “massaging”⁶². However, relabeling or perturbing data can introduce inaccuracies and distort the data's original distribution. This may lead to models learning incorrect or overly simplified

relationships, so these techniques must be thoroughly validated against the original data distributions.

- 2) Sampling: A dataset can be sampled up or down by adding or removing samples, respectively, to change the sample distributions to one that is more balanced. Up sampling the minority class can be done using duplication of existing samples or synthetic data, but caution is warranted since the model may start to overfit to duplicated or synthetic examples. Down sampling can involve removing majority group samples, but similarly, caution is warranted since data complexity can be reduced. Methods such as Synthetic Minority Over-sampling Technique (SMOTE)⁶³ attempt to balance these methods.

ML training methods

These methods are used to modify or alter ML training algorithms to improve model fairness.

- 1) Regularization and constraints: these methods alter the loss function of an algorithm. During regularization, an extra term is used to penalize discrimination, while constraints are used to limit the allowed bias level according to a certain loss function. Prejudice Remover⁶⁴, Exponentiated Gradient Reduction⁶⁵, Grid Search Reduction⁶⁵ and Meta Fair Classifier⁶⁶ are all examples of these techniques. A potential drawback of these methods is the difficulty in correctly defining the penalization terms or constraints. Overly strict constraints might lead to underfitting and reduced model performance, while poorly chosen terms can fail to adequately address bias. Additionally, these techniques may require significant computational resources and careful tuning, which can be challenging in practice.
- 2) Adversarial learning involves training two models that compete to improve their performance⁶⁷. Specifically, one model attempts to predict the true label of a dataset, while the other model attempts to exploit a known fairness issue using equality metrics. A drawback of adversarial learning is its complexity and the risk of instability during training, as the competing objectives of the models can lead to convergence issues. Adversarial models may inadvertently reduce overall predictive accuracy if fairness constraints conflict significantly with optimizing performance. Furthermore, adversarial models are not appropriate in scenarios where there are known differences between subgroups. As an example, when developing auto-segmentation models with a known performance difference between males and females adversarial learning should not be used since morphological differences are known to exist between sexes⁶⁸.

Prediction post-processing

Post-processing methods act on the model predictions and are used when access to training data or the model is limited. These methods would be more appropriately used in scenarios where the model is commissioned from an external group.

- 1) Classifier correction: A trained ML model is adapted to remove discrimination based on equalized odds and equality of opportunity constraints. Calibrated Equalized Odds⁶⁹ is an example of one of these methods. However, classifier correction depends on accurately identifying sources of bias and setting appropriate constraints. Poorly defined constraints can lead to either insufficient fairness improvements or a decrease in the model's predictive accuracy. Additionally, applying such corrections post-training may not address deeper issues of bias inherent in the data, limiting their effectiveness in mitigating unfairness.
- 2) Output correction: model outputs are modified to obtain fairer distribution of the data. Reject Option based Classification⁷⁰ is an example of this type of method that assigns more favourable outcomes to protected groups based upon low confidence regions of the classifier. While output correction can improve fairness metrics, it may distort predicted probabilities and reduce trust in model predictions. This approach addresses bias at the output level without resolving biases present in the training data or the model itself, potentially leading to superficial fairness improvements that fail to generalize to new datasets. Furthermore, aggressive corrections can harm overall model performance and usability for certain tasks.

4. Silent Deployment and Clinical Evaluation

4.1. Prospective Deployment Preparation

Approvals for prospective deployment

Similar to [Section 3.1.](#), access to and use of prospective clinical data requires institutional approval. Which type of institutional approval to seek depends on the specifics of the project, including the type of data that is required, the sensitivity of that data, the intended audience of the model outputs, and the nature of any potential interventions based on the model outputs.

Some important questions to consider:

1. Which clinical system(s) will provide prospective data?
2. Do you require data on demand (live data feed) or on a schedule?
3. Is there sufficient infrastructure to meet your data needs?
4. Are there limitations on how the prospective data can be used (e.g., Silent Mode data cannot be used to develop new models)?
5. What are the potential interventions that could be prompted by the deployed model, and could they have a direct impact on clinical decision-making?

Retrospective to prospective mapping

Healthcare data and EHR systems are constantly changing. In many cases, AI models are developed and tested using historical data, and/or from historical records systems. In these cases, every effort must be made to evaluate how this data corresponds to the data expected from defined sources of prospective healthcare data and clearly map the differences. It is also important to consider spectrum effects and any unintended disparities in data.

While navigating a change in EMR system between model development and silent deployment is a dramatic (albeit common) scenario, similar approaches are a necessary part of any prospective deployment. Even within any given EMR, important definitions such as procedure codes, names of fields, and other critical data may change over time without warning. After any initial data mapping is complete, prospective data needs to be compared to historical data to ensure consistency and accuracy. Questions you can ask yourself are:

1. Are all of the features included in your model being captured?
2. Are there any unexplained shifts in your data? (missingness, volumes, values, etc.)

This is an iterative process until the prospective data quality meets expectations but will continue as part of lifecycle monitoring.

Prospective Data Representativeness

In ideal scenarios the prospective data used during Silent Mode and Prospective Deployment will have demographic distributions and data collection methods that are similar. However, there are certain scenarios where retrospective and prospective datasets will differ, but the model still meets defined objectives and metrics. In this scenario, the model is robust to certain acceptable variable/data variations. These acceptable variations should be assessed and documented so that the model does not undergo unnecessary retraining or decommissioning.

AI-Solution Failures and Adverse Events

Before any prospective deployment of an AI-solution occurs, the team should define the bounds for failure of the AI (a deviation from the target performance that is none-the-less within expectations for research and development), and the bounds for a critical incident or adverse event involving the AI (a deviation from target performance that is either far beyond expectation or that causes/could have caused harm).

For example, if a medical chatbot has been implemented to allow clinicians to ask for up-to-date medication recommendations for their patients, there may be failure if:

- when prompted by the clinician, the chatbot returns no answer at all
- when asked for a list of options, the chatbot returns a medication recommendation that is nonsense (hallucinated)

However, there may be an incident if:

- when prompted by the clinician, the chatbot returns offensive/derogatory language
- when asked about a certain medication, the chatbot returns the names and contact information for all patients enrolled in a clinical trial related to that medication (a privacy incident)

Some situations are clear cut as either failures or critical incidents (for example, any physical harms to patients would be critical incidents). However, some situations may be less clear. The entire cross-functional team should be engaged to deliberate on the bounds of expected and unexpected AI-solution behaviour, and in particular, ethicists and patient and caregiving partners' expertise should be drawn on. There are situations in

which a clinician or researcher may perceive deviations in the AI solution to be a minor issue, but a patient on the receiving end may feel harmed.

4.2. Prospective Model Assessment (Silent Mode)

Silent Deployment Considerations

Many approaches to silent mode

There are many different potential approaches and phases to silent mode, and what is best will completely depend on the specific solution being evaluated. As a general guide, silent mode involves: (1) developing required workflow integration components that consider human factors and user experience; (2) validation of the model's performance; (3) reliability and integration within the clinical workflow under real-world settings; and (4) identification of potential biases, usability issues and any unintended consequences.

Threshold selection

Some AI-Solutions will require the selection of a statistical threshold, above and below which different actions are taken. This threshold determines the binarization point of a ML prediction. This threshold should be selected in collaboration with clinical-end-users and the project's cross-functional team. It should also be chosen to balance false positive rates and false negative rates, to minimize unnecessary operational burden and interventions, and delayed diagnosis or intervention, respectively⁷¹. These thresholds should be selected with consideration of subgroups.

Assessing Silent Mode Performance

Silent mode performance should broadly assess how the model performs in a real-world setting, as well as how the solution is integrated into current clinical workflows. The cross-functional team should be consulted if the model fails to meet any of the defined statistical measures of performance, equity objectives or fairness metrics to determine if the model should be adapted, mitigation strategies should be implemented, or if the AI-Solution is not meeting the intent of the project and should be discontinued.

Define auditing methods

In collaboration with the cross-functional team, an auditing plan should be generated prior to solution rollout. The frequency of audits will be determined by the team. Audits of the solution are comprehensive and are not meant to replace regular assessment of model performance. An audit of the solution should consider the following:

1. Ethics, security and privacy check of solution and data being used.
2. Changes in standard practice, structural changes, or other sources of data drift, that may affect the data being used.
3. Model performance across patient subgroups.
4. Whether or not equity objectives and fairness metrics are still appropriate and being met.
5. Feedback from end-users regarding experience, questions, adoption and any additional recommendations about the solution.
6. Failures that have occurred since the last audit (including if they were properly communicated and addressed)

4.3. Prospective Clinical Evaluation

Education materials

Developing education materials to provide during the Clinical Evaluation of an AI-Solution in healthcare is critical to ensure all stakeholders understand the technology, its goals, and their roles. This section highlights a sampling of materials a cross-functional team may need to develop for their Clinical Evaluation, depending on the stakeholders involved.

1. End-users:
 - a. User Manuals and Quick Start Guides:
 - Include examples of AI outputs and how to interpret them.
 - Provide mechanisms for users to report issues or provide feedback.
 - b. Decision Support Documentation:
 - Highlight evidence supporting the AI model's recommendations.
 - Address limitations, biases, and areas of uncertainty.
 - Describe situations where the model should not be used or where caution is needed.
 - c. Compliance and Safety Guidelines:
 - Clarify the trial's regulatory compliance and safeguards to prevent patient harm.
2. For Patients (and Caregivers, if applicable):
 - a. Simplified Information Sheets or Videos:
 - Describe the AI-Solution, its role in their care, and expected benefits.
 - b. Trust-Building FAQs:
 - Address concerns about AI (e.g., "Will AI replace my doctor?").
 - Reassure participants about clinician oversight and privacy safeguards as relevant.
 - c. Patient Feedback Channels:

- Provide materials explaining how participants can offer feedback or raise concerns.

Recruit end-users for future clinical integration

When assessing the AI-Solution during a Clinical Evaluation, obtaining feedback from end-users and patients involved in the trial is essential to evaluate whether clinical integration is effective. Various methods can be used to obtain this feedback including surveys, interviews, focus groups, or direct observation. Here are some examples of feedback to consider:

Feedback from end-users:

1. Is the AI-Solution intuitive and user-friendly?
2. Are the AI insights/actionable recommendations accurate and clinically meaningful?
3. Do the AI insights/actionable recommendations align with evidence-based practices or clinical guidelines?
4. Does the AI tool save time or reduce manual workload? Does it increase confidence in developing a care plan? Does it improve equitable care?
5. Is it reducing cognitive load or decision fatigue?
6. Is the rationale behind AI outputs clear and explainable?
7. *If the end-user is a clinician:* Do you feel confident relying on the tool?
8. *If the end-user is a clinician:* Are you noticing measurable improvements in patient care quality or outcomes?
9. Was adequate training provided on how to use this tool?
10. Is ongoing technical support accessible and helpful?

Feedback from Patients:

1. Are you aware that AI is part of your care process?
2. Do you understand how the AI contributes to your treatment?
3. Do you feel the tool is enhancing your care (e.g., faster results, personalized treatments, fewer errors or delays)?
4. Are there any concerns about depersonalization of care?
5. Do you feel your data is secure and used responsibly?
6. Are you comfortable with AI being used in decision-making?
7. Do you trust the collaboration between clinicians and AI?

5. Operationalization and Lifecycle Monitoring

5.1 Preparation and Documentation

Comprehensive documentation

Comprehensive documentation bridges the gap between development, deployment, and real-world use. Documentation should be generated both for internal usage (to ensure safe, ethical, and effective implementation within the local care system) and external publishing (to foster transparency, reproducibility, and collaboration).

Some best practices to keep in mind for both internal and external documentation include:

- (1) generating plain language summaries for accessibility to non-technical stakeholders;
- (2) ensuring all documentation aligns with FAIR principles (Findable, Accessible, Interoperable, and Reusable)⁷² when publishing; and (3) including a revision history of any updates and improvements to the AI-Solution.

Below are some recommendations for documentation for internal usage and external publishing:

Internal Documentation:

1. Model Development and Technical Specifications

a. Model Overview:

- Objectives and intended use cases.
- Description of the problem the AI addresses.

b. Model Architecture:

- Detailed explanation of algorithms, architecture, and training pipeline.
- Hyperparameters and optimization techniques.

c. Data Sources and Preprocessing:

- Description of datasets used for training, validation, and testing.
- Details on data cleaning, augmentation, and handling missing values.
- Consider use of Data Labels⁷³

d. Fairness, Equity and Performance Metrics:

- Intended equity objective(s) and utilized fairness metric(s)
- Performance evaluation metrics (e.g., accuracy, precision, recall, AUC-ROC).
- Comparison with baseline methods.

e. Ethics and Bias:

- Description of steps taken to minimize bias and promote fairness.
- Description of known biases and inequities in AI-Solution.

- f. Versioning:
 - Model version history and updates.
- 2. Deployment Documentation
 - a. System Integration:
 - How the AI integrates with existing healthcare infrastructure (including e.g. what data sources are accessed, who the primary end-users are, etc).
 - b. Operational Guidelines:
 - Deployment process, runtime environment, and maintenance requirements.
 - c. Interfaces:
 - User interface guides for clinicians or administrators.
- 3. Regulatory and Compliance Documentation
 - a. Data Privacy and Security:
 - Compliance with regulations like GDPR, HIPAA, or equivalent.
 - Explanation of data storage, access controls, and encryption methods.
 - b. Risk Management:
 - Identified risks and mitigation strategies.
- 4. Monitoring and Post-Deployment Validation
 - a. Performance Monitoring Plans:
 - Procedures for tracking model performance in real-world settings.
 - b. Update and Retraining Guidelines:
 - When and how the model should be updated or retrained.
 - c. Audit Logs:
 - Documentation of decision-making processes for accountability.

External Documentation (For Publishing to Journals or Open Access Repositories): The TRIPOD+AI⁷⁴ statement is an excellent resource that should be followed for external publishing of an AI-Solution.

- 1. Research and Model Development
 - a. Introduction and Background:
 - Problem statement, clinical relevance, and literature review.
 - b. Methods:
 - Detailed description of the AI development process, including:
 - 1. Model architecture and algorithms.
 - 2. Training pipeline and dataset descriptions.
 - 3. Statistical methods used for validation.

4. Explanation of fairness assessments and steps to mitigate bias
- c. Results:
 - Intended equity objective and utilized fairness metric
 - Evaluation metrics (e.g., accuracy, precision, recall, AUC-ROC), statistical significance tests, and comparative analysis..
 - Comparison with baseline methods or alternatives.
 - Visualization of results (e.g., confusion matrices, ROC curves).
- d. Discussion:
 - Interpretation of findings, limitations, and implications for clinical practice.
 - Intended users
 - Known limitations
2. Dataset Documentation (if sharing data)
 - a. Data Description:
 - Features, data sources, and collection methods.
 - b. Data Preprocessing Steps:
 - Cleaning, normalization, and other transformations.
 - c. Data Dictionary:
 - Definitions of variables and labels.
 - d. Ethical and Legal Considerations:
 - How patient privacy was protected.
 - Any restrictions on dataset usage.
3. Open Access Repositories
 - a. Code Repository (e.g., GitHub, GitLab):
 - Source code with comprehensive comments and clear organization.
 - Instructions for reproducing results (e.g., environment setup, dependencies).
 - b. Model Repository (e.g., Hugging Face, Zenodo):
 - Pre-trained models with detailed usage instructions.
 - Associated metadata for easy reference.
4. FAQs:
 - a. Common questions from reviewers, researchers, or practitioners

5.2. Communication and Education

Modifying trial educational materials

Well-designed educational materials help build confidence, trust and engagement across all stakeholders, enabling a smooth clinical integration of the AI-Solution. During this step

the education materials made during the Clinical Trial should be modified for a broad audience. Some recommended additions to the above section on education materials ([Section 4.3](#)) include:

1. End-users:
 - a. User Manuals and Quick Start Guides:
 - Expand to explain system functionalities, workflows and common troubleshooting steps.
 - b. Clinical Workflow Integration Training:
 - Provide interactive sessions or videos on how to use the AI-Solution.
 - c. FAQ and Troubleshooting Tips:
 - Address common concerns and technical issues.
2. For Organizational Stakeholders (e.g. administrators, IT):
 - a. Implementation Guides:
 - Highlight infrastructure, security, and resource requirements.
 - b. Compliance and Governance Briefs:
 - Outline regulatory, ethical, and operational responsibilities.

Communication plan

The communication plan should assure that the team has the following information readily available:

- a. Equity objectives, as defined at the start of the project
- b. Fairness metrics, as defined at the start of the project
- c. Model performance (general and in relation to fairness metrics) across defined subgroups
- d. Information about non-technical components of AI-Solution
- e. Contact information of solution management team
- f. Educational material to avoid automation bias
- g. Educational material on system usage

It is important to ensure that the information is complete and, to the extent possible, free from reporting bias⁷⁵. Reporting bias could manifest in a variety of ways. For example, regular performance reports could under-report unexpected or undesirable results or overemphasize expected or desirable results, or educational materials could obfuscate the intended use of AI-Solution in relation to how it is actually being deployed. In many cases, these reporting errors are not made maliciously, which makes them all the more important to stay alert for.

Communication plans should encompass dissemination of regular reports and audits, and patient and end-user educational materials. Plans should also be generated in advance

for events such as adverse event reporting (see [Communicating About Adverse Events](#)), AI-Solution update, AI-Solution failure, and AI-Solution decommissioning.

Communicating to Patients and Care Partners

For explainers, reports, and educational materials that will be provided to patients and care partners, consideration should be given to making the language and figures lay-accessible. There are many resources on the web that can be used to help convert medical jargon into more lay-friendly synonyms, such as the CDC's "[Everyday Words for Public Health Communication](#)" search tool. While many organizations encourage developing materials to be understandable to those with a Grade 6 reading level, there are no readability formulas that reliably evaluate the interpretability of complex information⁷⁶. Instead of focusing on the *readability* of patient and care partner materials, it is recommended that teams focus on *usability*, by using plain language wherever possible. The International Plain Language Federation considers a document to be in plain language if the communication's "wording, structure, and design are so clear that the intended readers can easily find what they need, understand what they find, and use that information."⁷⁷ Plain language can be facilitated with intentional word choice, well placed headings and subtitles, and graphic design.

In addition to lay-friendly, teams should strive to make patient and care partner materials accessible, for example for those with low vision/vision loss. There are simple steps teams can take to make material accessible for individuals with low vision, as laid out by the Canadian National Institute for the Blind's [Clear Print Guidelines](#). These include printing in black and white, keeping font size to a minimum of 12 to 18 pt font, using a medium heaviness font throughout most of the printed material, and printing on matte or non-glossy papers to minimize glare. Teams may note that some of the principles of plain language communication synergize well with principles of accessible design.

5.3. AI-Solution Rollout and Monitoring

Regular monitoring and reporting

The metrics that should be monitored and reported on will depend on the final AI-Solution. However, it is recommended that the AI-Solution monitoring focuses on performance, safety, compliance and operational impact. As examples: (1) model performance and data quality could be monitored by looking at error rates (false positives and false negatives), drift detection (in both data and model performance), and statistical metrics; (2) Clinical outcomes could be monitored by looking at workflow integration and adverse events, such as patient safety issues; (3) Operational metrics, security and privacy can be monitored by looking at uptime and response times, as well as usage metrics, data breaches and access logs; (4) finally, periodic revalidation of the AI-Solution may be beneficial using updated data or in the case of a system update that impacts the software and hardware used by the AI-Solution.⁷⁸

Utilizing the [developed communication plan](#), regularly report on the AI-Solution's impact and performance, as determined by audits, feedback from end-users and patients, and regular monitoring metrics. These regular reports should be comprehensive and include disparities in AI-Solution performance and current clinical adoption and impact.

Gathering Feedback for Continuous Quality Improvement

Very little literature exists on how to implement patient and end-user feedback in clinically integrated AI-solutions specifically. However, there is existing literature on integrating such feedback into other developed clinical workflows⁷⁹, which can be used as a scaffold.

It should be expected that if feedback is gathered from an appropriately diverse group of end users, some feedback will conflict – for example, for an AI-solution meant to flag patients at risk for pain crisis, some nurses may want to be alerted at frequently to an unlikely chance of crisis, while others may only want rare alerts for highly likely crises. Opportunity to improve the quality of the AI-solution overall lies in resolving these previously unaddressed tensions⁷⁹. Taking the previous example, the cross functional team may decide to implement two versions of the model depending on the time of day, so that nurses are alerted frequently during the day to control patients' pain, but less frequently during the night, when pain control measures might only serve to disturb the patients' sleep. Alternatively, the team might give only one nurse per shift access to the less sensitive model, and ask that nurse to decide, based on how well resourced the floor is, whether or not to follow up on low-certainty alerts.

Time must be taken to gather feedback specifically from patients affected by the AI-solution deployment. Patient feedback should be used to provide face validity for the impact of the AI-solution – if patients are noticing a primarily negative change, the team needs to evaluate if the clinical integration of the AI-solution is going as planned. Additionally, by continuously evaluating feedback from patients, change in the clinical environment and subsequent model drift may be detected earlier than by using numeric indicators of performance alone.

5.4. Updating or Decommissioning of AI-Solution

If at any point an AI-Solution is failing to meet the set objectives and requirements (i.e. fairness and performance metrics are worsening, institutional requirements are no longer being met, or end-users and/or patients are no longer seeing benefit from the solution), the cross-functional team should be consulted. During this consultation there are three possible next steps: Pausing the AI-Solution, Updating the AI-Solution, or Decommissioning the AI-Solution.

Each of these steps will look different depending on the institution and AI-Solution. Broadly speaking, it should be determined what each of these steps would look like from a clinical and regulatory perspective and who should be informed. Additionally, depending on the decided step, the results and/or changes will need to be communicated.

6. Community Tools

There are many free tools and resources to help in identifying, mitigating and evaluating bias and fairness. Below is a table containing some resources.

Resource	Summary	Source
Data Nutrition Labels	The Data Nutrition Project aims to create a standard label for interrogating datasets. Website provides tools to evaluate data quality, fairness and bias in data.	https://datanutrition.org/
JarvAIs	An auto-machine learning tool for data quality assurance, training, and explaining data models. Includes input data bias analysis and model bias analysis.	https://github.com/pmcidi/jarvais
Med-ImageNet	DICOM image data processing tool. Unifies data format and indexes TCIA for easy search and download.	https://github.com/bhklab/med-imagenet
FairLearn	An open-source python package designed for metric calculation and reduction of bias in algorithms.	https://fairlearn.org/
Fairness Indicators	A package from Google designed to work with TensorFlow. Can be used for evaluation and visualization of group disparities.	https://www.tensorflow.org/tfx/guide/fairness_indicators
AIF360	An open-source package used to detect and mitigate biases. It is known for its extensive documentation and tutorials.	https://aif360.res.ibm.com/
Themis-ML	An open-source package for easy integration with scikit-learn. It is used mainly for binary classes and evaluation.	https://www.themisai.io/

References

- 1 Webster CS, Taylor S, Thomas C, Weller JM. Social bias, discrimination and inequity in healthcare: mechanisms, implications and recommendations. *BJA Education* 2022; **22**: 131–7.
- 2 Fernández Pinto M. Methodological and Cognitive Biases in Science: Issues for Current Research and Ways to Counteract Them. *Perspectives on Science* 2023; **31**: 535–54.
- 3 Popovic A, Huecker MR. Study Bias. In: StatPearls. Treasure Island (FL): StatPearls Publishing, 2024. <http://www.ncbi.nlm.nih.gov/books/NBK574513/> (accessed Oct 7, 2024).
- 4 Baumann AA, Cabassa LJ. Reframing implementation science to address inequities in healthcare delivery. *BMC Health Serv Res* 2020; **20**: 190.
- 5 National Academies of Sciences E, Division H and M, Practice B on PH and PH, et al. The Root Causes of Health Inequity. In: Communities in Action: Pathways to Health Equity. National Academies Press (US), 2017. <https://www.ncbi.nlm.nih.gov/books/NBK425845/> (accessed Oct 25, 2024).
- 6 Varkey B. Principles of Clinical Ethics and Their Application to Practice. *Med Princ Pract* 2021; **30**: 17–28.
- 7 Coalition for Health AI. Responsible AI Guidance. <https://www.chai.org/workgroup/responsible-ai> (accessed July 24, 2025).
- 8 Coeckelbergh M. Artificial intelligence, the common good, and the democratic deficit in AI governance. *AI Ethics* 2025; **5**: 1491–7.
- 9 Cortes C, Mohri M, Riley M, Rostamizadeh A. Sample Selection Bias Correction Theory. In: Lecture Notes in Computer Science. Berlin, Heidelberg: Springer-Verlag, 2008: 38–53.
- 10 National Cancer Institute. Definition of selection bias. NCI Dictionary of Cancer Terms. 2011; published online Feb 2. <https://www.cancer.gov/publications/dictionaries/cancer-terms/def/selection-bias> (accessed July 30, 2025).
- 11 Nickerson RS. Confirmation Bias: A Ubiquitous Phenomenon in Many Guises. *Review of General Psychology* 1998; **2**: 175–220.
- 12 Pilat D, Krastev S. Confirmation Bias - Biases & Heuristics. The Decision Lab. 2021. <https://thedecisionlab.com/biases/confirmation-bias> (accessed July 30, 2025).
- 13 Kea B, Hall MK, Wang R. Recognising bias in studies of diagnostic tests part 2: interpreting and verifying the index test. *Emerg Med J* 2019; **36**: 501–5.

- 14 Elston DM. Survivorship bias. *Journal of the American Academy of Dermatology* 2021; **0**. DOI:10.1016/j.jaad.2021.06.845.
- 15 Pasqualetti F, Barberis A, Zanotti S, *et al*. The impact of survivorship bias in glioblastoma research. *Critical Reviews in Oncology/Hematology* 2023; **188**: 104065.
- 16 Hawkins DM. The Problem of Overfitting. *J Chem Inf Comput Sci* 2004; **44**: 1–12.
- 17 Park Y, Ho JC. Tackling Overfitting in Boosting for Noisy Healthcare Data. *IEEE Transactions on Knowledge and Data Engineering* 2021; **33**: 2995–3006.
- 18 Usher-Smith JA, Stephen J Sharp, Griffin SJ. The spectrum effect in tests for risk prediction, screening, and diagnosis. *BMJ* 2016; **353**: i3139.
- 19 Abdelwanis M, Alarafati HK, Tammam MMS, Simsekler MCE. Exploring the risks of automation bias in healthcare artificial intelligence applications: A Bowtie analysis. *Journal of Safety Science and Resilience* 2024; **5**: 460–9.
- 20 Sabin JA. Tackling Implicit Bias in Health Care. *New England Journal of Medicine* 2022; **387**: 105–7.
- 21 PEIRS. Clayon Hamilton. <https://www.clayonhamilton.com/peirs> (accessed Oct 1, 2025).
- 22 Hamilton CB, Hoens AM, McKinnon AM, *et al*. Shortening and validation of the Patient Engagement In Research Scale (PEIRS) for measuring meaningful patient and family caregiver engagement. *Health Expectations* 2021; **24**: 863–79.
- 23 Goodman MS, Ackermann N, Haskell-Craig Z, Jackson S, Bowen DJ, Sanders Thompson VL. Construct validation of the Research Engagement Survey Tool (REST). *Res Involv Engagem* 2022; **8**: 26.
- 24 Do You Have ‘Shiny Object’ Syndrome? What It Is and How to Beat It. Entrepreneur. 2017; published online Feb 9. <https://www.entrepreneur.com/living/do-you-have-shiny-object-syndrome-what-it-is-and-how-to/288370> (accessed July 30, 2025).
- 25 Guo LN, Lee MS, Kassamali B, Mita C, Nambudiri VE. Bias in, bias out: Underreporting and underrepresentation of diverse skin types in machine learning research for skin cancer detection—A scoping review. *Journal of the American Academy of Dermatology* 2022; **87**: 157–9.
- 26 Seyyed-Kalantari L, Zhang H, McDermott MBA, Chen IY, Ghassemi M. Underdiagnosis bias of artificial intelligence algorithms applied to chest radiographs in under-served patient populations. *Nat Med* 2021; **27**: 2176–82.

- 27 Khor S, Haupt EC, Hahn EE, Lyons LJL, Shankaran V, Bansal A. Racial and Ethnic Bias in Risk Prediction Models for Colorectal Cancer Recurrence When Race and Ethnicity Are Omitted as Predictors. *JAMA Netw Open* 2023; **6**: e2318495.
- 28 Gichoya JW, Banerjee I, Bhimireddy AR, *et al.* AI recognition of patient race in medical imaging: a modelling study. *Lancet Digit Health* 2022; **4**: e406–14.
- 29 Reeder-Hayes KE, Jackson BE, Kuo T-M, *et al.* Structural Racism and Treatment Delay Among Black and White Patients With Breast Cancer. *J Clin Oncol* 2024; **42**: 3858–66.
- 30 Ring A, Harder H, Langridge C, Ballinger RS, Fallowfield LJ. Adjuvant chemotherapy in elderly women with breast cancer (AChEW): an observational study identifying MDT perceptions and barriers to decision making. *Ann Oncol* 2013; **24**: 1211–9.
- 31 Protière C, Viens P, Rousseau F, Moatti JP. Prescribers' attitudes toward elderly breast cancer patients. Discrimination or empathy? *Crit Rev Oncol Hematol* 2010; **75**: 138–50.
- 32 Ring A. The influences of age and co-morbidities on treatment decisions for patients with HER2-positive early breast cancer. *Crit Rev Oncol Hematol* 2010; **76**: 127–32.
- 33 Haase KR, Sattar S, Pilleron S, *et al.* A scoping review of ageism towards older adults in cancer care. *J Geriatr Oncol* 2023; **14**: 101385.
- 34 Cirillo D, Catuara-Solarz S, Morey C, *et al.* Sex and gender differences and biases in artificial intelligence for biomedicine and healthcare. *NPJ Digit Med* 2020; **3**: 81.
- 35 Lee MS, Guo LN, Nambudiri VE. Towards gender equity in artificial intelligence and machine learning applications in dermatology. *J Am Med Inform Assoc* 2022; **29**: 400–3.
- 36 Larrazabal AJ, Nieto N, Peterson V, Milone DH, Ferrante E. Gender imbalance in medical imaging datasets produces biased classifiers for computer-aided diagnosis. *Proceedings of the National Academy of Sciences* 2020; **117**: 12592–4.
- 37 Turner M, Fielding S, Ong Y, *et al.* A cancer geography paradox? Poorer cancer outcomes with longer travelling times to healthcare facilities despite prompter diagnosis and treatment: a data-linkage study. *Br J Cancer* 2017; **117**: 439–49.
- 38 Ambroggi M, Biasini C, Del Giovane C, Fornari F, Cavanna L. Distance as a Barrier to Cancer Diagnosis and Treatment: Review of the Literature. *Oncologist* 2015; **20**: 1378–85.

- 39 Maly RC, Umezawa Y, Leake B, Silliman RA. Mental health outcomes in older women with breast cancer: impact of perceived family support and adjustment. *Psychooncology* 2005; **14**: 535–45.
- 40 Bevan JL, Pecchioni LL. Understanding the impact of family caregiver cancer literacy on patient health outcomes. *Patient Educ Couns* 2008; **71**: 356–64.
- 41 Mickel J. The Importance of Multi-Dimensional Intersectionality in Algorithmic Fairness and AI Model Development. 2023; published online May. <https://hdl.handle.net/2152/122644> (accessed Dec 3, 2024).
- 42 Hajian-Tilaki K. Receiver Operating Characteristic (ROC) Curve Analysis for Medical Diagnostic Test Evaluation. *Caspian J Intern Med* 2013; **4**: 627–35.
- 43 Çorbacioğlu ŞK, Aksel G. Receiver operating characteristic curve analysis in diagnostic accuracy studies: A guide to interpreting the area under the curve value. *Turk J Emerg Med* 2023; **23**: 195–8.
- 44 Czakon J. F1 Score vs ROC AUC vs Accuracy vs PR AUC: Which Evaluation Metric Should You Choose? neptune.ai. 2022; published online July 21. <https://neptune.ai/blog/f1-score-accuracy-roc-auc-pr-auc> (accessed July 30, 2025).
- 45 Richardson E, Trevizani R, Greenbaum JA, Carter H, Nielsen M, Peters B. The receiver operating characteristic curve accurately assesses imbalanced datasets. *PATTER* 2024; **5**. DOI:10.1016/j.patter.2024.100994.
- 46 Bertels J, Robben D, Vandermeulen D, Suetens P. Theoretical analysis and experimental validation of volume bias of soft Dice optimized segmentation maps in the context of inherent uncertainty. *Medical Image Analysis* 2021; **67**: 101833.
- 47 Taha AA, Hanbury A. Metrics for evaluating 3D medical image segmentation: analysis, selection, and tool. *BMC Medical Imaging* 2015; **15**: 29.
- 48 Gorriz JM, Clemente RM, Segovia F, Ramirez J, Ortiz A, Suckling J. Is K-fold cross validation the best model selection method for Machine Learning? 2024; published online Nov 8. DOI:10.48550/arXiv.2401.16407.
- 49 Jung Y, Hu J. A K-fold Averaging Cross-validation Procedure. *J Nonparametr Stat* 2015; **27**: 167–79.
- 50 Carey S, Pang A, Kamps M de. Fairness in AI for healthcare. *Future Healthcare Journal* 2024; **11**: 100177.
- 51 Gao J, Chou B, McCaw ZR, et al. What is Fair? Defining Fairness in Machine Learning for Health. 2025; published online May 27. DOI:10.48550/arXiv.2406.09307.
- 52 Rajkomar A, Hardt M, Howell MD, Corrado G, Chin MH. Ensuring Fairness in Machine Learning to Advance Health Equity. *Ann Intern Med* 2018; **169**: 866–72.

- 53 Chen RJ, Wang JJ, Williamson DFK, *et al.* Algorithmic fairness in artificial intelligence for medicine and healthcare. *Nat Biomed Eng* 2023; **7**: 719–42.
- 54 Fabris A, Esuli A, Moreo A, Sebastiani F. Measuring Fairness Under Unawareness of Sensitive Attributes: A Quantification-Based Approach. *Journal of Artificial Intelligence Research* 2023; **76**: 1117–80.
- 55 Bou V. Achieving Demographic Parity Across Multiple Artificial Intelligence Applications: A New Approach for Real-Time Bias Mitigation. 2024; published online Dec 5. DOI:10.20944/preprints202412.0468.v1.
- 56 Davoudi A, Chae S, Evans L, *et al.* Fairness gaps in Machine learning models for hospitalization and emergency department visit risk prediction in home healthcare patients with heart failure. *International Journal of Medical Informatics* 2024; **191**: 105534.
- 57 Roberts-Licklider K, Trafalis T. Machine Learning Techniques with Fairness for Prediction of Completion of Drug and Alcohol Rehabilitation. 2024; published online April 23. DOI:10.48550/arXiv.2404.15418.
- 58 Feng Q, Du M, Zou N, Hu X. Fair Machine Learning in Healthcare: A Review. 2024; published online Feb 1. DOI:10.48550/arXiv.2206.14397.
- 59 D’Souza G, Zhang HH, D’Souza WD, Meyer RR, Gillison ML. Moderate predictive value of demographic and behavioral characteristics for a diagnosis of HPV16-positive and HPV16-negative head and neck cancer. *Oral Oncol* 2010; **46**: 100.
- 60 Kim S-Y, Choi Y, Kim E-K, *et al.* Deep learning-based computer-aided diagnosis in screening breast ultrasound to reduce false-positive diagnoses. *Sci Rep* 2021; **11**: 395.
- 61 Feldman M, Friedler SA, Moeller J, Scheidegger C, Venkatasubramanian S. Certifying and Removing Disparate Impact. In: Proceedings of the 21th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining. New York, NY, USA: Association for Computing Machinery, 2015: 259–68.
- 62 Kamiran F, Calders T. Classifying without discriminating. In: Control and Communication 2009 2nd International Conference on Computer. 2009: 1–6.
- 63 Chawla NV, Bowyer KW, Hall LO, Kegelmeyer WP. SMOTE: synthetic minority over-sampling technique. *J Artif Int Res* 2002; **16**: 321–57.
- 64 Kamishima T, Akaho S, Asoh H, Sakuma J. Fairness-Aware Classifier with Prejudice Remover Regularizer. In: Flach PA, De Bie T, Cristianini N, eds. Machine Learning and Knowledge Discovery in Databases. Berlin, Heidelberg: Springer, 2012: 35–50.

- 65 Agarwal A, Beygelzimer A, Dudík M, Langford J, Wallach H. A Reductions Approach to Fair Classification. 2018; published online July 16. DOI:10.48550/arXiv.1803.02453.
- 66 Celis LE, Huang L, Keswani V, Vishnoi NK. Classification with Fairness Constraints: A Meta-Algorithm with Provable Guarantees. 2020; published online April 15. DOI:10.48550/arXiv.1806.06055.
- 67 Yang J, Soltan AAS, Eyre DW, Yang Y, Clifton DA. An adversarial training framework for mitigating algorithmic biases in clinical machine learning. *NPJ Digit Med* 2023; **6**: 55.
- 68 Fan Y, Penington A, Kilpatrick N, *et al.* Quantification of mandibular sexual dimorphism during adolescence. *J Anat* 2019; **234**: 709–17.
- 69 Pleiss G, Raghavan M, Wu F, Kleinberg J, Weinberger KQ. On Fairness and Calibration. 2017; published online Nov 3. DOI:10.48550/arXiv.1709.02012.
- 70 Franc V, Prusa D, Voracek V. Optimal Strategies for Reject Option Classifiers. *Journal of Machine Learning Research* 2023; **24**: 1–49.
- 71 Ruopp MD, Perkins NJ, Whitcomb BW, Schisterman EF. Youden Index and Optimal Cut-Point Estimated from Observations Affected by a Lower Limit of Detection. *Biometrical Journal* 2008; **50**: 419–30.
- 72 Wilkinson MD, Dumontier M, Aalbersberg IJ, *et al.* The FAIR Guiding Principles for scientific data management and stewardship. *Sci Data* 2016; **3**: 160018.
- 73 Chmielinski KS, Newman S, Taylor M, *et al.* The Dataset Nutrition Label (2nd Gen): Leveraging Context to Mitigate Harms in Artificial Intelligence. 2022; published online March 10. DOI:10.48550/arXiv.2201.03954.
- 74 Collins GS, Moons KGM, Dhiman P, *et al.* TRIPOD+AI statement: updated guidance for reporting clinical prediction models that use regression or machine learning methods. *BMJ* 2024; **385**: e078378.
- 75 Sinha MK, Montori VM. Reporting bias and other biases affecting systematic reviews and meta-analyses: a methodological commentary. *Expert Review of Pharmacoeconomics & Outcomes Research* 2006; **6**: 603–11.
- 76 Government of Canada PS and PC. Readability formulas, programs and tools: Do they work for plain language? – The Our Languages blog – Resources of the Language Portal of Canada – Canada.ca. 2025; published online Oct 1. <https://www.noslangues-ourlanguages.gc.ca/en/blogue-blog/readability-formulas-eng> (accessed Oct 1, 2025).

- 77 Plain Language Definitions - International Plain Language Federation. 2019; published online June 28. <https://www.iplfederation.org/plain-language-definitions/> (accessed Oct 1, 2025).
- 78 Feng J, Phillips RV, Malenica I, *et al.* Clinical artificial intelligence quality improvement: towards continual monitoring and updating of AI algorithms in healthcare. *NPJ Digit Med* 2022; **5**: 66.
- 79 Nordberg SS, McAleavey AA, Moltu C. Continuous quality improvement in measure development: Lessons from building a novel clinical feedback system. *Qual Life Res* 2021; **30**: 3085–96.