



FAIR LABS

FAIR ARTIFICIAL INTELLIGENCE RESEARCH LABS

POLICY BRIEF ♦ AI & HEALTH ♦ 2025

Responsible AI in Telehealth

Guidance for health systems adopting AI-assisted telehealth – safety, oversight, consent, and equity, from procurement to the point of care.

PROGRAM

AI & Health

DOCUMENT TYPE

Policy Brief

AUDIENCE

Health Systems & Regulators

PUBLISHED

2025

FAIR Labs · A 501(c)(3) nonprofit research institute · Washington, DC

CONTENTS

Table of Contents

—	Executive Summary	3
I	Telehealth's AI Moment	5
II	Use Cases & Their Risks	7
III	Clinical Validation & Safety	10
IV	Clinician Oversight & the Human Role	13
V	Consent, Privacy & Data Governance	16
VI	Equity & Access	19
VII	Reimbursement, Liability & Regulation	22
VIII	Recommendations for Adoption	24
—	Notes & References	27

EXECUTIVE SUMMARY

The bottom line for health-system leaders

IN ONE PARAGRAPH

Health systems do not need to become AI developers to adopt AI-assisted telehealth responsibly. They need a disciplined way to sort telehealth AI by clinical risk, an evidence standard that demands validation on their own patients before deployment, an oversight design that keeps a clinician meaningfully in charge rather than nominally in the loop, a consent and data-governance posture that treats patients as parties rather than data sources, and an explicit commitment that every deployment be examined for its effect on the populations least served by remote care. This brief supplies each of these as an instrument that can be lifted into local policy.

4

use-case families that structure risk: triage, diagnostic support, documentation, and remote monitoring

3

clinical-risk tiers that set the depth of validation and oversight for any tool

5

stages of the validation lifecycle, from intended-use definition through post-deployment surveillance

What we recommend

1**Classify by clinical risk before you buy.**

Sort every candidate tool by how much it steers a clinical decision and how reversible its errors are. The tier—not the vendor's enthusiasm—sets the evidentiary and oversight bar for everything that follows.

2**Validate on your own patients.**

A model's published accuracy describes the population it was tested on. Require local validation on a representative sample of your patients, disaggregated by the groups you serve, before a tool touches a live encounter.

3**Keep a clinician genuinely in charge.**

Design against automation bias deliberately. A clinician who lacks the time, information, or standing to overrule the system is a liability shield, not an overseer.

4**Treat consent, privacy, and equity as deployment gates.**

Patients should know when AI is involved and be able to decline it without losing care; data should be governed to a defined purpose; and no tool should launch without an equity review of who it may leave behind.

The sections that follow develop each recommendation in operational detail, with tables and checklists a health system can adapt directly into procurement and clinical policy.

SECTION I

Telehealth's AI Moment

Telehealth crossed from the margins of medicine to its mainstream in the space of a few years, driven by a public-health emergency, relaxed reimbursement, and a patient population that discovered it preferred not to sit in waiting rooms. What began as a stopgap—video visits to keep chronic care from lapsing—has settled into a permanent modality. Into that newly normalized channel, artificial intelligence is now arriving on a similarly compressed schedule: symptom checkers at the front door, decision support inside the visit, ambient systems that write the note, and monitoring platforms that watch the patient between appointments. The convergence is not coincidental. Remote care generates the structured, digital, continuous data that machine learning feeds on, and machine learning promises to make remote care scale without a proportional increase in clinicians.

That promise is real, and so are the stakes. Telehealth AI operates where three difficult conditions coincide. The clinician is at a distance, deprived of the physical examination and the tacit signals of an in-person encounter. The patient is often on consumer hardware of uneven quality, in an uncontrolled setting, sometimes navigating a second language or an unfamiliar interface. And the system is making or shaping decisions—who is urgent, what is likely, what should be documented—that were until recently the exclusive province of trained human judgment. AI can extend scarce clinical expertise across that distance. It can also transmit error across it at scale, faster than any single clinician could.

Why telehealth is a distinctive setting for AI

It is tempting to treat telehealth AI as ordinary clinical AI that happens to run over video. That framing understates three features that change the governance problem. First, the **loss of the physical exam** removes information a clinician would normally use to sanity-check an algorithm's suggestion; the human's own error-correction is weaker precisely where the machine is relied upon more. Second, the **heterogeneity of the patient's environment**—lighting, camera, connectivity, background noise, digital fluency—injects variance into the model's inputs that no laboratory validation captures. Third, the **speed and volume** of remote encounters invite exactly the throughput pressures under which humans defer to automation. Each feature is individually familiar; their combination is what makes telehealth a demanding proving ground.

A WORKING DEFINITION

In this brief, **telehealth AI** means any capability whose behavior is materially shaped by a model learned from data and that operates within a remote-care pathway—symptom intake, triage, diagnostic or decision support, clinical documentation, or remote monitoring—together with the devices, data pipelines, clinicians, and patient interactions that surround it. The unit we govern is the *care pathway*, not the algorithm.

The drivers pushing adoption

Four forces are pulling AI into telehealth simultaneously, and understanding them helps a health system separate durable value from hype. There is a **capacity** driver: chronic workforce shortages and clinician burnout make any tool that promises to reduce cognitive or clerical load attractive, sometimes uncritically so. There is an **access** driver: the hope that automated triage and monitoring can extend care to populations and hours that human staffing cannot reach. There is a **data** driver: remote care produces the digital exhaust—structured intake forms, device streams, transcribed conversations—that makes model-building feasible. And there is a **commercial** driver: a vigorous vendor market that has correctly identified telehealth as fertile ground and that markets accordingly.

None of these drivers is illegitimate, but each can distort a procurement decision if left unexamined. The capacity driver, in particular, deserves caution: a tool adopted primarily to relieve an overworked staff will be used by an overworked staff, under precisely the conditions that make careful oversight least likely. A governance program that does not account for the reason a tool is wanted will be surprised by how it is actually used.

The stakes, in both directions

The failure that makes headlines is the dramatic one: a triage system that sends a heart attack home, a monitoring platform that misses a deterioration. Those matter, and the tiering that follows is designed around them. But the more common institutional harm is quieter—an ambient scribe that subtly distorts what a patient said, a symptom checker that works well for the digitally fluent and poorly for everyone else, a decision aide that clinicians learn to accept because contesting it costs time they do not have. These failures rarely produce a lawsuit or a news story. They erode, slowly, the trust on which remote care entirely depends.

Equally real is the cost of paralysis. A health system that cannot distinguish a low-risk documentation aide from a high-risk autonomous triage tool will tend to govern both as if they were the latter—smothering harmless efficiency in review committees while, paradoxically, still lacking the capacity to scrutinize the deployments that genuinely warrant it. Proportionality is not a convenience here; it is what makes rigorous oversight affordable where a patient's safety depends on it. The rest of this brief is an attempt to hold both truths at once: to move quickly where AI is low-stakes, precisely so that the system can move carefully where it is not.

SECTION II

Use Cases & Their Risks

The phrase "AI in telehealth" hides four quite different activities, each with its own benefit, its own characteristic failure, and its own appropriate level of scrutiny. Governing them as one category guarantees that the riskiest are under-examined and the safest are over-burdened.

We group telehealth AI into four use-case families. The families are not perfectly exclusive—an ambient scribe may also surface diagnostic suggestions—but they are stable enough to anchor a risk conversation, and they map cleanly onto the different kinds of evidence and oversight later sections require. What distinguishes them is not the underlying technology, which is often shared, but the clinical consequence of the system being wrong.

Triage and symptom intake

At the front door of remote care sit tools that ask a patient about their symptoms and route them: to self-care, to a nurse line, to an urgent visit, or to an emergency department. This is where AI most directly touches the decision that matters most in acute care—who is seen, and how fast. A triage tool that is too cautious floods clinicians and erodes its own credibility; a triage tool that is too permissive sends a serious presentation home. Because triage decisions are consequential and, when wrong in the dangerous direction, hard to reverse, this family sits at the top of the risk hierarchy whenever the tool's output is allowed to influence disposition without a clinician in between.

Diagnostic and clinical decision support

Inside the encounter, AI increasingly offers suggestions: a differential diagnosis, a flag on an uploaded image, a prompt about a drug interaction, a risk score. Some of this is regulated as a medical device; much of it is marketed as "decision support" that leaves the clinician nominally in charge. The risk here is subtle. These tools rarely act autonomously, so their danger is mediated through the clinician's response to them—which means their real-world safety depends heavily on the oversight conditions discussed in Section IV. A diagnostic aide that is right most of the time is precisely the kind of tool a busy clinician learns to trust reflexively.

Documentation and ambient scribing

Ambient documentation systems listen to a clinical conversation and draft the note. They are among the fastest-spreading telehealth AI applications because they target the clerical burden most associated with burnout, and because their failures feel low-stakes. They are not as low-stakes as they feel. A generated note is a legal and clinical record that downstream clinicians, coders, and courts will treat as an account of what happened. A scribe that omits a stated symptom, invents a

detail through confabulation, or subtly reframes a patient's words introduces error into the record at its source—and the clinician who is meant to review it is the same clinician the tool was sold to relieve.

Remote patient monitoring

Beyond the visit, monitoring platforms ingest data from connected devices—weight, glucose, blood pressure, cardiac rhythm, activity—and use models to detect deterioration, predict events, or prioritize which patients a care team should contact. The promise is genuine: continuous signal where medicine once had episodic snapshots. The risks are alert fatigue that trains staff to ignore the system, false reassurance from a device that has drifted out of calibration, and inequitable performance across skin tones, body types, or home environments that the device was not designed for. Monitoring also raises the sharpest data-governance questions, because it collects continuously, in the home, often from the whole household's ambient environment.

TELEHEALTH AI USE CASES — BENEFIT, CHARACTERISTIC FAILURE, AND DEFAULT RISK TIER

Use-case family	Primary benefit	Characteristic failure	Default risk tier
Triage & symptom intake	Faster, consistent routing; extended after-hours access	Under-triage of a serious presentation; over-triage that floods capacity	Tier 3 — High (Tier 2 with clinician gate)
Diagnostic & decision support	Surfaces possibilities and risks a rushed clinician may miss	Confident wrong suggestion the clinician defers to	Tier 2–3, by autonomy and consequence
Documentation / ambient scribing	Reduced clerical burden; more attention to the patient	Omission, confabulation, or distortion in the record	Tier 2 — Elevated
Remote patient monitoring	Continuous signal; earlier detection of deterioration	Missed event, alert fatigue, inequitable sensor performance	Tier 2–3, by clinical dependence

THE TIERING RULE

Two questions set the tier for any telehealth tool. How much does its output *steer* a clinical decision—does a human meaningfully stand between it and the patient? And how *reversible* is its error—can a mistake be caught before it causes durable harm? A tool that steers heavily and errs irreversibly is Tier 3, whatever its use-case family and whatever the vendor calls it.

Autonomy is the hinge

The same underlying model can belong to different tiers depending on how much authority a deployment grants it. A symptom checker that produces advice a patient reads directly, with no clinician in the loop, is a higher-risk deployment than one whose output is visible only to a nurse who makes the disposition. A monitoring model that auto-escalates is riskier than one that populates a worklist a human triages. This is why the brief insists on tiering deployments rather than certifying tools: the identical algorithm can be safe in one configuration and unsafe in another, and the difference lies entirely in the human authority the health system chooses to retain.

The question is never whether the model is good. It is how much the patient's safety depends on the model being good, and who catches it when it is not.

— The governing premise of use-case tiering

SECTION III

Clinical Validation & Safety

A telehealth AI tool that performs well in a published study has demonstrated one thing: that it worked on the population, in the setting, and for the task the study measured. None of those may match the health system about to deploy it. The central discipline of safe adoption is refusing to accept generalization on faith—insisting that a tool prove itself on the patients it will actually serve, at the task it will actually perform, before it touches a live encounter. This is not bureaucratic caution. It is the difference between a validation and an advertisement.

Define the intended use before measuring anything

Regulators who oversee software as a medical device begin with a deceptively simple question: what is the tool for? The intended use—the clinical condition, the patient population, the role the output plays in a decision—fixes everything downstream. A model validated to *flag* images for a radiologist's review has not been validated to *replace* that review. A triage tool validated for adults has not been validated for children. Much of the danger in clinical AI comes not from models that fail at their intended use but from models used outside it, where no evidence exists at all. The first act of validation is therefore to write the intended use down, narrowly, and to treat any use beyond it as unvalidated by definition.

Evidence proportional to risk

Not every tool warrants a randomized trial, and demanding one for a documentation aide would simply block useful tools while consuming the scrutiny that a triage system deserves. The depth of evidence should scale with the tier established in Section II. A Tier 2 tool may be adequately established by a well-designed local evaluation against clinician judgment; a Tier 3 tool that steers consequential, hard-to-reverse decisions warrants prospective evidence, independent evaluation, and where feasible a comparison against the standard of care it is meant to augment or replace.

PRE-DEPLOYMENT CLINICAL VALIDATION CHECKLIST

Validation question	What good evidence looks like	Minimum by tier
Intended use defined	Written statement of condition, population, task, and the role of the output in the decision	All tiers
Analytical performance	Sensitivity, specificity, and calibration reported for the stated task, not a proxy	All tiers
Representative test data	Evaluation on data resembling the deploying system's actual patients and devices	Tier 2–3
Subgroup performance	Results disaggregated by age, sex, race/ethnicity, language, and device where relevant	Tier 2–3
Local validation	Prospective or silent-mode evaluation on the deploying system's own patients	Tier 3 (recommended Tier 2)
Failure-mode analysis	Documented behavior on edge cases, poor inputs, and out-of-distribution presentations	Tier 2–3
Human-factors evaluation	Study of how clinicians actually use the output, including override behavior	Tier 3
Post-deployment plan	Monitoring, drift detection, and a defined re-validation trigger and cadence	All tiers

Validate the pathway, not just the algorithm

A model's laboratory accuracy is measured on clean inputs. In telehealth the inputs are anything but clean: a photograph taken on a cracked phone camera in poor light, audio degraded by a weak connection, an intake form completed by a patient guessing at medical terminology. Analytical performance on curated data therefore systematically overstates real-world performance.

Meaningful validation exercises the whole pathway—the patient's device, the network, the human on both ends—because that is where telehealth AI actually fails. A tool that is accurate in the lab and brittle on a consumer phone is not a safe tool; it is a tool that has not been tested where it lives.

NON-NEGOTIABLE FOR TIER 3

For high-risk tools, the party that validates the system should be independent of the party that wants to deploy it—an internal quality or informatics unit with a separate reporting line, or an external evaluator. Self-graded homework is not evidence. A vendor's benchmark describes the vendor's benchmark; local validation on your own patients is what tells you whether the tool works here.

The special hazard of generative documentation

Ambient scribing and other generative tools introduce a failure mode the accuracy-and-calibration frame does not fully capture: **confabulation**, the fluent generation of plausible detail that was never said. A number reported to three decimals is easy to distrust; a well-written clinical note is not. Validation of generative documentation must therefore measure not only how often the note is right but how its errors are distributed—whether they are the kind a reviewing clinician will catch or the kind that reads correctly and slips through. Fluency is a safety hazard precisely because it disarms scrutiny.

The safety case that emerges from validation should be honest about what remains unknown. A tool is rarely proven safe in the abstract; it is proven safe enough for a defined use, given stated evidence, with residual risks named rather than hidden. A validation report that cannot state the conditions under which the tool should not be trusted is not finished.

SECTION IV

Clinician Oversight & the Human Role

"Human-in-the-loop" is the most invoked and least examined phrase in clinical AI. A clinician who cannot understand the system, lacks the time to question it, or fears the consequences of overruling it is not an overseer. She is a liability shield wearing the costume of oversight.

Health systems reach for the clinician-in-the-loop as reassurance, and often as a regulatory strategy: because a licensed professional signs off, the tool is characterized as decision support rather than autonomous decision-making, and a lighter regulatory regime applies. But this reassurance is only as real as the oversight beneath it. Meaningful human oversight has preconditions, and a deployment that does not engineer them will produce the appearance of human control and the reality of automated care.

Automation bias is the central threat

Decades of human-factors research across aviation, industrial control, and medicine converge on a robust finding: people defer to automated systems, especially confident ones, and especially when overriding them requires extra work. In a busy telehealth clinic, every one of those conditions holds. The system is confident by design; the clinician is under throughput pressure; and overriding the machine means documenting a justification the clinician does not have time to write. Under these conditions, "the clinician reviews the output" quietly becomes "the clinician accepts the output," and the human-in-the-loop dissolves into a rubber stamp that carries the legal weight of a decision no human really made.

THE AUTOMATION-BIAS TRAP

An oversight design that ignores clinicians' predictable tendency to defer will produce documented human review and undocumented automated decision-making. Countermeasures—selective verification, friction on high-stakes overrides in *both* directions, and routine tracking of acceptance and override rates—must be designed in, not hoped for. If a tool is overridden almost never, that is a finding to investigate, not a success to celebrate.

Preconditions for meaningful oversight

For clinician oversight to be real rather than nominal, four conditions must hold, and a deployment plan should be able to point to each. Clinicians must *understand* the tool well enough to know when to distrust it—its intended use, its known failure modes, and the situations where its confidence is misleading. They must have the *time* to exercise judgment rather than clear a queue; oversight

designed to be squeezed between visits is oversight designed to fail. They must face *incentives* that do not punish a careful override, so that questioning the machine is professionally safe. And they must have *usable access to the information* behind the output—not a bare score, but enough context to reason about whether to trust it in this case.

A Calibrate the interruption to the stakes.

Present low-tier outputs as quiet, ignorable assistance; reserve interruptive, attention-demanding prompts for the high-stakes moments where deference is most dangerous. Alerting everything trains clinicians to dismiss everything.

B Make the override cheap and the record honest.

Overriding the system should be at least as easy as accepting it, and the record should reflect who actually decided. Oversight that is more work than compliance will not survive a busy afternoon.

C Show the clinician what she needs to judge.

Surface the basis for a suggestion—the inputs, the uncertainty, the comparable cases—so the clinician can reason about it rather than accept or reject it blind.

D Measure the human, not just the model.

Track acceptance and override rates and their outcomes as a first-class monitoring signal. How the clinician uses the tool is as much a part of the deployment's safety as the model's accuracy.

The scribe's paradox

Ambient documentation illustrates the oversight problem in miniature. The tool is sold to relieve clinicians of clerical work, and its safety case rests on the assumption that the clinician will review the generated note before signing it. But the same clinician, freed from writing the note, faces the same time pressure that made the tool attractive—and reviewing a fluent, plausible draft demands more vigilance than reading one's own terse notes, not less. The tool's value proposition and its safety precondition are in direct tension. Resolving that tension requires deliberate design—flagging low-confidence passages, preserving the source audio for spot-checks, and setting a cultural expectation that the note is the clinician's, not the machine's.

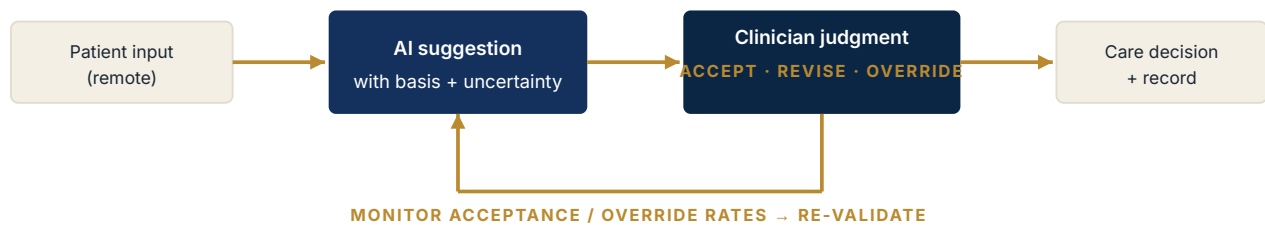


Figure 1. The oversight loop. The AI suggestion carries its basis and uncertainty forward so the clinician can reason rather than defer; acceptance and override behavior is fed back as a monitored safety signal, not treated as settled at deployment.

SECTION V

Consent, Privacy & Data Governance

A patient who agrees to a telehealth visit is agreeing to be cared for at a distance. It does not follow that they have agreed to have their words transcribed by a model, their images processed by a third-party algorithm, or their home data streamed to a monitoring platform and reused to improve a commercial product. Consent that was designed for a video call does not automatically extend to the AI now layered on top of it, and treating it as though it does mistakes the absence of objection for the presence of agreement.

Disclosure: patients should know when AI is involved

The baseline obligation is candor. A patient has a legitimate interest in knowing when an algorithm is materially shaping their care—triaging their symptoms, drafting their record, watching their vitals—particularly when the AI's involvement is invisible in a way a human's would not be. Disclosure need not be alarming or technical; it needs to be truthful and timely. The relevant question for a health system is not whether disclosure is legally mandated in its jurisdiction, which varies, but whether a reasonable patient would want to know—and, for consequential uses, they plainly would.

Consent proportional to stakes and reuse

Not every use of AI requires an explicit, separate consent, and pretending otherwise would bury patients in consent fatigue that protects no one. The appropriate consent scales with two factors: how much the AI affects the patient's care, and how far their data travels beyond it. A model that operates entirely within the treatment relationship, on data the patient already provided for that care, sits at one end. A system that streams home-monitoring data to a vendor who reuses it to train and improve commercial products sits at the other, and the patient's agreement to the first should never be silently stretched to cover the second.

CONSENT AND DATA-GOVERNANCE POSTURE BY USE AND DATA FLOW

Scenario	Disclosure	Consent	Data reuse
AI within the encounter, no external reuse	Notice at point of care	Implied in consent to treatment	Confined to the patient's care
Ambient scribing of the visit	Explicit notice; audio handling explained	Opportunity to decline the recording	Note retained; audio disposition defined
Consequential triage or decision support	Explicit notice of AI's role	Right to a human alternative	Confined; logged for audit
Remote monitoring in the home	Explicit; scope and household effects	Affirmative, revocable	Defined purpose; no silent secondary use
Any secondary use (model training, product improvement)	Explicit and specific	Separate, opt-in; care not conditioned on it	Only under that specific consent

The right to a human alternative

For consequential uses, disclosure without an option is hollow. A patient who is told that an algorithm will triage them, but who has no way to reach a human instead, has been informed of a decision, not given a choice. The principle we recommend is straightforward: for Tier 3 telehealth AI, a patient may decline the AI pathway and receive care through a human route without penalty or meaningful delay. Where staffing genuinely cannot support a human alternative, that is a reason to reconsider the deployment, not to quietly remove the patient's choice.

A GOVERNANCE PRINCIPLE

Data collected for care is held in trust for care. Any use beyond the patient's treatment—training a model, improving a product, sharing with a partner—is a new purpose that requires its own justification and, for identifiable data, its own consent. "We may use your data to improve our services," buried in a terms-of-service agreement, is not consent to secondary use of health data; it is the absence of it dressed as agreement.

Governance obligations that do not depend on the patient

Consent is necessary but not sufficient; a patient cannot meaningfully consent to safeguards they cannot see. The health system therefore owes duties that operate regardless of any individual agreement. Data should be governed to a **defined purpose** and minimized to what that purpose requires. It should be **secured** in transit and at rest, with particular care for the consumer devices and home networks that telehealth necessarily involves. Access should be **logged and auditable**, so the system can answer who saw what and why. And responsibility for the data should be **contractually**

fixed with any vendor before deployment, including what happens to the data when the relationship ends. These are the terms on which a patient's trust is either honored or quietly spent.

Telehealth asked patients to let care into their homes. AI asks to let the data out. The second permission is not implied by the first.

— On the limits of inherited consent

SECTION VI

Equity & Access

Telehealth AI is deployed across exactly the fault lines—digital, linguistic, economic, and physiological—most likely to distort it. Equity is therefore not a downstream concern to be audited after launch. It is a property that must be designed in, or it will be designed out.

Every layer of a telehealth AI pathway can concentrate harm on the populations a health system is most obligated to serve. The patient needs a device, connectivity, and the digital fluency to use them; the model was trained on data that may under-represent them; the sensors may perform worse on their bodies; the interface may not speak their language. Each layer is a filter, and the filters compound. A tool that is equitable in the aggregate can still fail systematically for the tail of patients who most need care to work—and because those patients are also least able to complain, the failure is quiet.

The digital divide is the first filter

Before any algorithm runs, telehealth requires a device, a connection, and a private space. These are unevenly distributed along lines of income, age, geography, and disability. AI can widen this gap or narrow it. A system designed for the newest smartphone and reliable broadband quietly excludes the patient on an old device and a metered connection—often the same patient with the greatest clinical need. A system designed to degrade gracefully to a phone call, to work on modest hardware, and to tolerate poor connectivity can instead extend reach. The design choice is not neutral, and it is made early, usually by default, unless someone insists otherwise.

EQUITY CONSIDERATIONS ACROSS THE TELEHEALTH AI PATHWAY

Layer	How inequity enters	Design and governance response
Access & devices	Excludes patients without current hardware, broadband, or private space	Support low-bandwidth and voice-only modes; provide human and phone alternatives
Digital fluency	Interfaces assume literacy and comfort many patients lack	Plain-language design; usability testing with the hardest-to-reach users
Training data	Under-representation yields worse performance for some groups	Demand subgroup performance; validate locally on the served population
Sensor / physiology	Devices calibrated on narrow populations mis-measure others	Test across skin tones, body types, and conditions; document known limits
Language access	Non-dominant languages are served worse or not at all	Validate performance per language; do not deploy where quality is unproven
Interpretation & trust	Outputs assume context patients may not share	Culturally competent design; maintain human relationships, not just automation

Algorithmic disparity: the average conceals the tail

A model reported to be accurate "overall" may be markedly less accurate for a subgroup that is a small fraction of the training data and of the aggregate metric. This is the best-documented failure mode in clinical AI, and it is not a rare accident; it is the default behavior of a system optimized for average performance on an unrepresentative sample. The remedy is unglamorous and specific: demand performance disaggregated by the groups the health system serves, treat a large subgroup gap as a deployment blocker rather than a footnote, and re-check those gaps after launch, because they can emerge or widen as the population shifts. A single aggregate accuracy number is not evidence of equity; it is the statistic most likely to hide its absence.

EQUITY AS A DEPLOYMENT GATE

No Tier 2 or Tier 3 telehealth AI should launch without an equity review that answers three questions in writing. Whom does this tool's access requirement exclude, and what is their alternative? How does its performance vary across the groups we serve? And what would we monitor to know if it is widening a disparity? A tool that cannot answer these is not ready, however strong its aggregate performance.

Language access is a safety issue, not a courtesy

Language-access failures in telehealth AI are patient-safety failures. A symptom checker or an ambient scribe that performs well in the dominant language and poorly in others does not merely

inconvenience limited-English-proficiency patients; it introduces clinical error precisely where a language barrier already elevates risk. The governance response is to treat each supported language as a distinct deployment requiring its own validation, and to decline to deploy in a language where performance is unproven rather than offering a degraded version to the patients least able to catch its mistakes. Extending a tool to a new language is a new clinical claim, and it deserves new evidence.

Designing for the margins improves the whole

There is a practical consolation in this work. Systems built to serve the hardest cases—the older patient on an old phone, the patient in a second language, the patient on a slow connection—tend to be more robust for everyone, because they cannot rely on the fragile assumptions that break first at the margins. Equity and quality are not in tension here. A telehealth AI program that designs for its most vulnerable patients is building a more reliable system for all of them, and one whose failures, when they come, are less likely to fall along the lines a health system is most accountable for.

SECTION VII

Reimbursement, Liability & Regulation

The clinical and ethical case for careful telehealth AI adoption meets a practical reality: the rules governing payment, liability, and regulatory clearance are unsettled, vary by jurisdiction, and were largely written for a world without learning systems in the care pathway. A health system cannot wait for that landscape to resolve before acting, but it can adopt in a way that is defensible under the rules as they stand and adaptable as they change.

The regulatory frame: software as a medical device

The most developed regulatory concept for clinical AI treats qualifying software as a medical device, subject to review proportional to its risk. The intuition is sound and portable across jurisdictions: a tool that informs a low-stakes decision warrants lighter oversight than one that drives a high-stakes one, and the intended use determines the category. Two features of learning systems, however, strain the traditional device paradigm. Models can **update** after clearance, so a system evaluated once is not the system in use a year later; regulators have responded with the idea of a predetermined change-control plan that authorizes bounded modification. And much telehealth AI is marketed as **clinical decision support** that leaves a clinician in charge, which in several regimes exempts it from device regulation entirely—an exemption that is only as protective as the oversight beneath it, as Section IV argued.

REGULATORY POSTURE

Do not let a tool's regulatory status substitute for your own validation. A device clearance attests that a tool met a regulator's bar for a stated intended use; it does not attest that the tool works on your patients, in your workflow, for your task. And a decision-support exemption does not lower the clinical risk—it relocates the responsibility for managing it onto the health system and the clinician. Either way, the burden of local assurance remains yours.

Reimbursement shapes behavior, for better and worse

Payment rules are among the most powerful and least examined forces shaping telehealth AI. When a monitoring service is reimbursed per patient enrolled, the incentive is to enroll broadly, whether or not each patient benefits. When documentation tools are justified by throughput, the incentive is to see more patients in less time—the very condition that undermines the clinician oversight the tool's safety depends on. A health system should scrutinize what a reimbursement pathway is actually paying for and whether that aligns with patient benefit, because the payment model will shape clinical behavior more reliably than any policy statement. Reimbursement that rewards volume over vigilance will get volume over vigilance.

Liability: distributed authorship, concentrated exposure

When a telehealth AI tool contributes to a harm, responsibility is genuinely distributed—across the developer who built the model, the health system that deployed it, and the clinician who acted on its output—yet the law's instinct is to find a responsible human, and the clinician is the nearest one. This produces a precarious position: the clinician bears liability for decisions influenced by a system she did not build, cannot fully inspect, and may not have been given the time to properly question. The oversight preconditions in Section IV are, among other things, the health system's best protection here—for its clinicians and for itself. A clinician who is genuinely equipped to supervise a tool is defensible; one who was set up to rubber-stamp it is not, and neither is the system that arranged it.

ALLOCATING RESPONSIBILITY ACROSS THE TELEHEALTH AI LIFECYCLE

Responsibility	Developer / vendor	Health system	Clinician
Valid model for stated use	Primary	Verify locally	—
Fit for this population & workflow	Support	Primary	—
Oversight conditions in place	—	Primary	Exercise
The individual clinical decision	—	Enable	Primary
Post-deployment monitoring	Cooperate	Primary	Report
Disclosure & consent to patient	—	Primary	Deliver

Contract as governance instrument

Because much telehealth AI is bought rather than built, the contract is the point of maximum leverage—and the moment it is signed is when that leverage evaporates. Several protections can only be secured in the agreement, before deployment: the right to be notified of and to re-validate before a material model change; the vendor's obligation to cooperate with incident investigation on a defined timeline, including access the health system cannot generate itself; auditability sufficient to the tool's risk tier; clear allocation of data responsibility and its disposition at exit; and the ability to decommission cleanly without punitive cost. Lock-in is not merely a commercial hazard in this setting; it is a clinical one, because it raises the price of pausing a tool that patient safety may require pausing.

SECTION VIII

Recommendations for Adoption

No health system adopts responsible telehealth AI in a single decision. The realistic path is staged, and the aim is not to eliminate risk—which is not on offer—but to make the health system able to answer, for any tool it runs, what it is, why it is trusted, who is accountable, and how it would know if it were causing harm.

A staged adoption maturity ladder

TELEHEALTH AI GOVERNANCE MATURITY — FROM AD-HOC TO INSTITUTIONALIZED

Level	Character	What it looks like in practice
1 · Ad-hoc	Unmanaged	AI enters through individual departments and vendor pilots; no inventory, no shared standard, no one accountable across tools.
2 · Aware	Reactive	The system recognizes telehealth AI as a distinct risk; a policy exists on paper; application is inconsistent across service lines.
3 · Defined	Proportionate	Risk tiering, a validation checklist, and consent and equity gates are documented and applied to new tools.
4 · Managed	Assured	Monitoring, oversight metrics, and incident response operate for consequential tools; accountability is real and resourced.
5 · Institutionalized	Self-improving	Governance is routine and funded; lessons flow back into procurement and standards; the program improves itself.

A pragmatic first year

1

Inventory and tier what you already run (months 1–3).

You cannot govern what you cannot see. Build a register of telehealth AI already in use—including quiet vendor pilots and departmental tools—and assign each a provisional risk tier. This step alone typically surfaces surprises, including tools no one realized were making consequential decisions.

2**Adopt the gates for new tools (months 2–6).**

Apply the validation checklist, the consent posture, and the equity review to everything new, where your leverage is highest and the cost of doing so is lowest. Do not attempt to retrofit every existing tool at once; start at the front door.

3**Name accountable clinicians for Tier 3 tools (months 3–6).**

Assign a specific, senior clinician who owns each high-risk deployment's safety and holds the authority to pause it. Ownership is the control that makes every other control enforceable; a tool no one can stop is a tool the system does not control.

4**Stand up monitoring where failure is most consequential (months 6–12).**

Begin post-deployment surveillance—performance, subgroup gaps, drift, and clinician acceptance/override rates—on Tier 3 tools first. A modest monitoring capability on the riskiest tools beats an elaborate plan that never launches.

Standing commitments beyond the first year

A**Re-validate on a schedule and on triggers.**

Treat every authorization as conditional and expiring. Re-validate on a defined cadence and whenever the model updates, the population shifts, or an incident suggests the original assessment was optimistic.

B**Watch for scope creep.**

The most common way a safe tool becomes unsafe is quiet expansion of use—a documentation aide that starts driving decisions, a triage tool extended to a new population. Make "how is this actually being used?" a standing audit question, not a one-time assumption.

C**Keep the human relationship intact.**

Automation should extend clinicians, not replace the relationship at the center of care. Preserve accessible human routes, especially for the patients least served by remote and automated care.

D Close the loop from incidents to procurement.

When a tool fails, feed the lesson back into how the next tool is validated and bought. A governance program that does not learn from its own incidents will keep buying the same failure.

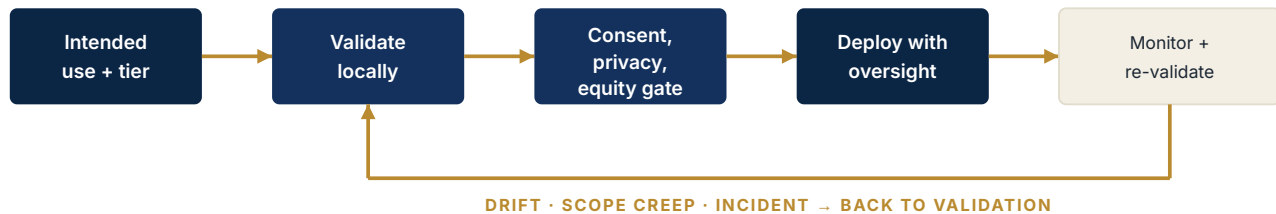


Figure 2. The adoption workflow as a loop. Intended use and tier set the depth of validation; consent, privacy, and equity are gates rather than afterthoughts; and deployment is the beginning of oversight, not the end—monitoring feeds drift, scope creep, and incidents back to re-validation.

What good looks like

A health system running this guidance well has not banished risk from remote care; it has made itself able to answer, for any telehealth AI tool it operates, a short and honest list of questions without scrambling: *What is it for, and what tier is it? What evidence, on our own patients, supports its use? Who is the clinician accountable for it, and can they pause it? Do patients know it is there, and can they decline it? Whom might it leave behind, and what are we watching to find out? And how quickly would we know, and act, if it began to cause harm?* An institution that can answer these has adopted AI in telehealth responsibly. The instruments in this brief exist to make those answers routine.

CLOSING NOTE

The aim of this guidance is not to slow telehealth down. It is to let health systems move quickly on the many genuinely low-stakes uses of AI in remote care—precisely because they have the discipline to move carefully on the few that decide who is seen, what is recorded, and who is watched. Proportionality is not the enemy of innovation in telehealth; it is what allows innovation to be trusted by the patients it is meant to serve.

APPENDIX

Notes & Further Reading

This brief synthesizes regulatory guidance, professional-society statements, and the peer-reviewed literature on clinical AI, telehealth, and health equity. The sources below are entry points into that literature; they are indicative rather than exhaustive.

1. World Health Organization. *Ethics and Governance of Artificial Intelligence for Health*. Guidance articulating principles including protecting autonomy, ensuring safety, fostering responsibility and accountability, and promoting equity.
 2. World Health Organization. *Regulatory Considerations on Artificial Intelligence for Health*. On lifecycle oversight, transparency, and validation for AI-enabled medical products.
 3. U.S. Food and Drug Administration. *Software as a Medical Device (SaMD): Clinical Evaluation and the Proposed Regulatory Framework for Modifications to AI/ML-Based SaMD (Predetermined Change Control Plans)*. On risk-based evaluation and governing model updates after clearance.
 4. International Medical Device Regulators Forum. *Software as a Medical Device: Possible Framework for Risk Categorization and Corresponding Considerations*. On tiering software risk by significance of information and state of the healthcare situation.
 5. National Institute of Standards and Technology. *AI Risk Management Framework (AI RMF 1.0)*. A voluntary framework organizing AI risk into govern, map, measure, and manage functions.
 6. Coalition for Health AI (CHAI). *Blueprint for Trustworthy AI Implementation Guidance and Assurance for Healthcare*. A consortium framework for assurance across the health-AI lifecycle.
 7. Sendak, M., et al. *Presenting Machine Learning Model Information to Clinical End Users ("Model Facts" labels)*. On concise, standardized disclosure of a model's intended use and performance to clinicians.
 8. Obermeyer, Z., et al. *Dissecting Racial Bias in an Algorithm Used to Manage the Health of Populations*. On how an aggregate-accurate model concealed systematic racial disparity.
 9. Liu, X., Rivera, S. C., et al. *SPIRIT-AI and CONSORT-AI* reporting guidelines. On standards for clinical trials of interventions involving AI.
 10. Collins, G. S., et al. *TRIPOD* (Transparent Reporting of a multivariable prediction model for Individual Prognosis Or Diagnosis) and its AI extension. On rigorous reporting of clinical prediction models.
 11. Parasuraman, R., & Riley, V. *Humans and Automation: Use, Misuse, Disuse, Abuse*. On automation bias and the conditions for meaningful human oversight.
 12. Mitchell, M., et al. *Model Cards for Model Reporting*; Gebru, T., et al. *Datasheets for Datasets*. On structured documentation of model performance across conditions and of data provenance and composition.
-

© 2025 FAIR Labs — Fair Artificial Intelligence Research Labs, a 501(c)(3) nonprofit, Washington, DC. This report may be reproduced and distributed for noncommercial, educational, and policy purposes with attribution. It constitutes research and policy guidance, not legal or clinical advice.