



COGNITIVE VIEW
AGETECH & AGED CARE • AI ASSURANCE

AgeTech AI Developer & Vendor Readiness Guide

A practical, evidence-based guide for procurement, pilot, deployment and ongoing assurance

CognitiveView | Version 1.0 | September 2026 | www.cognitiveview.com

What this guide is for

AgeTech AI can affect care delivery, safety, independence, workflow and access to sensitive information. This guide helps developers and vendors explain—not simply claim—why an AI system is appropriate for its intended older-adult population and care environment. It also helps provider organizations structure questions across clinical, technology, privacy, risk, AI governance and procurement functions.

Important

This is an educational and practical readiness resource. Completing the guide does not constitute regulatory approval, clinical validation, certification or procurement authorization. Each organization remains responsible for its own governance, risk acceptance, clinical/care review, procurement and deployment decisions.

Contents & How to Use This Guide

Part	Purpose
1. Readiness at a Glance	A concise view of what procurement and deployment readiness means.
2. Organization & Product Profile	Establish what the vendor is, what the AI system does and how it is deployed.
3. Intended Use & Care Context	Define users, population, environment, workflow and boundaries.
4. AI Risk Classification	Determine how much assurance is proportionate to potential harm and influence.
5. Data & Older-Adult Population Representation	Explain data provenance, quality, representation and subgroup considerations.
6. Responsible AI Testing & Evaluation	Describe how usefulness, performance, fairness, safety, usability and robustness are evaluated.
7. Safety, Reliability & Failure Management	Explain what happens when the system is wrong, uncertain, unavailable or degraded.
8. Human Oversight & Accountability	Define who reviews, interprets, escalates and can override AI outputs.
9. Transparency & Documentation	Provide Model/System Card information, limitations and traceable evidence.
10. Real-World Validation	Show whether the system works in representative care environments.
11. Privacy, Security & Compliance	Document data handling, safeguards, security and applicable requirements.
12. Health-System Procurement & Decision-Making	Prepare for the functions that may evaluate and approve an AI vendor.
13. Assurance Frameworks & Standards	Map the system to relevant voluntary frameworks and applicable requirements.
14. Governance, Change Management & Monitoring	Maintain assurance throughout the lifecycle.
15. Vendor Evidence Package & Due-Diligence Questions	Turn the guide into a practical procurement evidence room.
16. Readiness Decision & Reassessment	Record the decision, conditions, residual risk and reassessment triggers.

Recommended approach

Read this document as a narrative first, then use the evidence prompts and readiness questions at the end of each section. The objective is not to maximize the number of boxes checked. The objective is to create a coherent, traceable assurance story from intended use → risk → evidence → human oversight → deployment → monitoring.

1. Readiness at a Glance

Provider organizations need to make several different judgments before an AgeTech AI system is procured or deployed. A strong vendor package connects the product's intended purpose to its risks, evaluation evidence, governance arrangements and operating controls.

Readiness dimension	What a provider should be able to understand
Purpose	What problem does the AI solve, for whom, and in what care or living environment?
Risk	What could go wrong, how serious could the harm be, and what safeguards are required?
Population	Was the system developed and tested with populations relevant to intended older-adult users?
Performance	What evidence demonstrates useful, reliable performance for the intended use?
Safety	How are false alarms, missed events, unsafe outputs and system failures handled?
Human oversight	Who is accountable for interpreting outputs, intervening and escalating concerns?
Transparency	Can users and buyers understand the system, limitations, evidence and version?
Real-world fit	Has the system been validated in an environment representative of deployment?
Privacy & security	How is sensitive information collected, used, retained, protected and shared?
Lifecycle assurance	How are incidents, drift, changes and revalidation handled after deployment?

A useful procurement principle

An assurance-ready vendor does not simply state that its AI is safe or responsible. It provides dated, traceable evidence showing what was tested, with whom, under what conditions, what the limitations are, and how risk is managed in operation.

2. Organization & Product Profile

Before evaluating the AI itself, a provider needs a clear picture of the supplier, the product and the technical boundary being assessed. This prevents ambiguity about which components are covered by the assurance evidence and which are supplied by third parties.

What should be described

- Legal entity, product owner and accountable executive.
- Primary technical, clinical/care, privacy/security and governance contacts.
- AI system/product name, version and release status.
- AI application type: Generative, Predictive or Hybrid.
- Primary functions, material dependencies and third-party components.
- Deployment model, integrations and operating environment.
- Customer support, incident response and escalation contacts.

Evidence to prepare

Use these artifacts to make the discussion concrete and reduce repeated evidence requests.

- Product/system profile.
- Architecture or high-level system description.
- Version and release history.
- Third-party model/component and subprocessor inventory.
- Named accountable contacts.

3. Intended Use & Care Context

Intended use is the anchor for the entire assurance process. A system cannot be meaningfully evaluated without knowing who will use it, what decision or action it supports, what older-adult population is involved, and what happens when the output is wrong or unavailable.

Question	What good looks like
What does it do?	A clear statement of the AI function and the problem it is intended to address.
Who uses it?	Older adults, caregivers, clinicians, care coordinators, administrators or other users are explicitly identified.
Where is it used?	Home, residential care, clinic, hospital, community or other operating context is defined.
What does it output?	Inputs, outputs, alerts, recommendations or generated content are described.
What action follows?	The workflow and responsible human role are documented.
What is outside scope?	Out-of-scope and prohibited uses are explicit.
What happens when it is uncertain?	Low-confidence, unavailable or degraded states have defined handling.

AgeTech-specific consideration

Older adults may interact with AI through different physical, cognitive, sensory, environmental and support contexts. The assurance case should therefore address accessibility, usability and care-team workflow—not only technical model performance.

- Q What would happen if the AI were unavailable for a day?
- Q Could a user reasonably mistake an AI recommendation for a clinical or authoritative decision?
- Q Which actions remain human decisions?
- Q What uses are explicitly not supported?

4. AI Risk Classification

Risk classification should reflect the potential impact of the AI in its intended context. The existing AgeTech materials use risk tiers to determine the depth of testing, monitoring, documentation and human oversight.

Tier	Typical interpretation	Illustrative examples	Assurance expectation
Low	Informational or lower-impact use without direct influence on health/safety decisions.	Education, wellness information, low-impact assistance.	Document intended use, limitations and baseline evaluation.
Moderate	Supports care or daily monitoring but does not directly trigger high-consequence intervention.	Medication reminders, caregiver support, daily activity monitoring.	More structured testing, user evaluation, subgroup analysis and ongoing monitoring.
High	Safety-relevant monitoring, detection or alerts that may influence intervention.	Fall detection, emergency alerts, deterioration detection.	Strong safety/reliability evidence, real-world validation, human oversight and continuous monitoring.

Risk should be revisited

Risk can change when the intended use, target population, data source, model, workflow, operating environment or level of automation changes. A previously acceptable risk assessment should not be treated as permanent.

Evidence to prepare

Use these artifacts to make the discussion concrete and reduce repeated evidence requests.

- Risk assessment and rationale.
- Risk tier and decision date.
- Potential harms and affected stakeholders.
- Required human oversight and escalation.
- Monitoring intensity and reassessment triggers.

5. Data & Older-Adult Population Representation

Population fit is central to AgeTech assurance. A system can perform well in aggregate while behaving differently for groups that are underrepresented in development or evaluation data. Vendors should therefore explain what data were used, how populations were represented, what quality limitations exist and how subgroup performance was considered.

Key areas

- Data sources and provenance are documented.
- Development, validation and testing populations are characterized.
- Relevant demographic and contextual characteristics are considered.
- Known under-representation and data gaps are disclosed.
- Data quality, bias and limitations are assessed.
- Material differences between development populations and intended deployment populations are evaluated.
- Privacy, consent and data-rights considerations are documented as applicable.

Evidence area	Example evidence
Population representation	Dataset summary and relevant demographic/contextual breakdowns.
Data quality	Missingness, labeling quality, known limitations and quality controls.
Subgroup performance	Performance or safety analysis across relevant groups.
Bias/fairness	Fairness methodology, results and mitigation actions where relevant.
Data governance	Data-use, retention, access and protection documentation.

6. Responsible AI Testing & Evaluation

The AgeTech Testing & Evaluation approach should be treated as a lifecycle rather than a single benchmark. Evaluation should cover the dimensions that matter for the intended use and risk: usefulness, performance, usability, fairness, safety, reliability, transparency, human oversight, deployment robustness, privacy/security and monitoring.

Evaluation dimension	What should be demonstrated
Usefulness	The AI addresses the intended problem and provides meaningful value.
Performance / efficacy	Defined metrics demonstrate performance appropriate to the intended task.
Usability & accessibility	Older adults and/or care teams can use the system appropriately and safely.
Fairness	Relevant subgroup differences are assessed and material concerns are addressed.
Safety	Critical failure modes, missed events, false alarms and unsafe outputs are assessed where relevant.
Reliability / robustness	Performance remains dependable under intended operating conditions.
Transparency	Limitations, uncertainty and system behavior can be understood by relevant users.
Human oversight	Human review, intervention and escalation work as intended.
Deployment robustness	The system is evaluated in representative environments and operating conditions.
Privacy & security	Relevant privacy and security risks are assessed.
Continuous monitoring	Post-deployment signals and thresholds are defined.

Evidence should be traceable

For each material assurance claim, a buyer should be able to identify the evaluation method, evidence date/version, population or test conditions, result, limitation and owner of the evidence.

Evidence to prepare

Use these artifacts to make the discussion concrete and reduce repeated evidence requests.

- Evaluation plan and methodology.
- Results for performance, safety, fairness and usability as applicable.
- Test population description.
- Known limitations and unresolved issues.
- Real-world or pilot evaluation results.
- Evidence version/date and system version.

7. Safety, Reliability & Failure Management

Trustworthiness is tested most clearly when things go wrong. Vendors should explain how the system behaves when it produces an incorrect output, low-confidence result, false alarm, missed event, unavailable service or degraded signal.

Failure mode	Questions to answer
False positive / false alarm	How is the alert handled, and what burden or harm can repeated false alarms create?
False negative / missed event	What safeguards exist for events the AI fails to detect?
Low confidence	Does the system communicate uncertainty and trigger human review?
Unavailable system	What fallback or manual process is used?
Poor-quality input	How does the system detect or respond to missing/noisy/degraded inputs?
Material incident	Who is notified, what is investigated, and how is corrective action tracked?
Performance degradation	What monitoring detects deterioration and what happens next?

Safety principle

The assurance case should make clear which risks are accepted, which are mitigated through technical controls, and which require human intervention or operational safeguards.

8. Human Oversight & Accountability

Human oversight is not simply a statement that a human is 'in the loop'. The provider should understand exactly what the human sees, what decisions remain human, when review is required, how escalation works and whether the person has the authority and information needed to override or intervene.

Oversight element	What to define
Responsible role	Who is accountable for interpreting and acting on outputs?
Review threshold	When is human review mandatory?
Escalation	What triggers escalation and to whom?
Override	Can the output be overridden, suppressed or corrected?
User understanding	Are limitations and uncertainty communicated clearly?
Auditability	Can important decisions, interventions and incidents be reconstructed?

9. Transparency & Documentation

Transparency should enable a provider, care team and other relevant users to understand what the system is designed to do, how it was evaluated, what its limitations are and which version is being assessed.

- AI/System Model Card or equivalent documentation is available.
- Intended use, users, population and environment are documented.
- Out-of-scope uses and known limitations are disclosed.
- Evaluation methodology and key results are documented.
- Material changes are versioned and traceable.
- Current evidence can be distinguished from historical or superseded evidence.

Model/System Card

The Applied Model Card concept in the AgeTech materials provides a practical way to bring intended use, population, evaluation, limitations and governance information together in a consistent format.

10. Real-World Validation

Pre-deployment testing is necessary but may not demonstrate that an AI system works in the conditions where it will actually be used. Real-world validation should examine representative users, environments and workflows and should capture operational issues that laboratory testing can miss.

Validation area	What to examine
Representative population	Are intended older-adult users represented?
Care-team interaction	Can caregivers/clinicians interpret and act on outputs appropriately?
Workflow fit	Does the system work within the real process rather than creating unsafe workarounds?
Environment	Do home, residential, clinic or hospital conditions affect performance?
Operational reliability	Does the system remain available and dependable in practice?
Human outcomes	Are false alarms, missed events, usability and intervention consequences understood?
Pilot learning	Were findings fed back into the readiness/deployment decision?

11. Privacy, Security & Compliance

AgeTech products may process sensitive information, including health information and information about daily activities or living environments. Vendors should provide a clear account of what data are collected, why they are needed, how they are protected, how long they are retained, who receives them and how third-party dependencies are governed.

Area	Evidence / explanation expected
Data flows	What enters the system, where it goes and what is returned?
Purpose limitation	Why is each material data element collected or used?
Retention	How long is information retained and under what policy?
Sharing	Who can access or receive the information, including third parties?
Security	What access controls, encryption, monitoring and security practices apply?
Incident response	How are privacy/security incidents identified, escalated and communicated?
Applicable requirements	Which laws, contractual requirements and sector expectations apply?

Frameworks are not automatically requirements

The existing materials appropriately distinguish voluntary frameworks/guidance from mandatory legal or regulatory obligations. Applicability depends on the product, intended use, jurisdiction, customer and deployment context.

12. Health-System Procurement & Decision-Making

When an AgeTech vendor contracts with a clinic, health system, ACO or other care organization, procurement may involve multiple functions. The exact structure varies by organization, so the developer should prepare for a cross-functional evidence review rather than assuming that one individual owns the entire decision.

Function	Typical focus	Evidence that helps
Care / business sponsor	Problem, value, workflow fit	Use case, usefulness, workflow and user evidence
Clinical / care leadership	Safety, care impact, appropriateness	Safety results, population evidence, human oversight
Digital / IT / data	Integration, architecture, operations	Architecture, dependencies, reliability, data flows
Privacy / security	Data use and protection	Privacy documentation, security assessment, controls
Risk / compliance / legal	Accountability and requirements	Risk classification, governance, contractual/legal evidence
AI governance / Responsible AI	AI-specific risk and assurance	Evaluation evidence, Model/System Card, control/evidence mapping
Procurement / vendor management	Supplier due diligence	Evidence package, policies, attestations, commercial/operational information
Executive / approval committee	Overall risk acceptance	Readiness decision, residual risk, conditions and decision record

What accelerates procurement

A procurement-ready evidence package can reduce repeated requests by giving different reviewers a common, traceable set of artifacts. It does not replace the customer's own governance or approval process.

- Q What assurance evidence will each reviewing function require?
- Q Which committee or accountable role has authority to approve deployment?
- Q What customer-specific controls or contractual requirements apply?
- Q What conditions must be met before pilot or production use?

13. Assurance Frameworks & Standards Alignment

There is no single universal AI assurance standard for every AgeTech system. A provider may reference different frameworks or guidance depending on risk, intended use and jurisdiction. The developer should therefore explain alignment in an evidence-based way rather than presenting a framework name as proof of safety.

Reference	How it may be useful
NIST AI Risk Management Framework	Structured approach to AI risk management and trustworthy AI characteristics.
ISO/IEC 42001	AI management-system governance and organizational controls.
ISO/IEC 23894	AI risk-management guidance.

Reference	How it may be useful
CHAI	Healthcare-focused AI assurance and responsible AI considerations.
WHO ethics & governance guidance	Ethical and governance considerations for AI in health.
FDA / ONC resources, where applicable	Relevant health/medical AI or health IT expectations depending on product and use.
Local / sector requirements	Customer, jurisdictional, contractual and regulatory requirements that apply to the specific system.

How to make an alignment claim credible

For every framework claim, identify the relevant requirement/control, the evidence that demonstrates alignment, the system/version covered, the assessment date and any exceptions or gaps.

14. Governance, Change Management & Continuous Monitoring

AI assurance continues after procurement. Material changes to models, data, sensors, integrations, workflows or intended use can change the risk profile. Providers and vendors should agree on what changes require notification, reassessment or revalidation.

Lifecycle activity	Expected practice
Baseline	Establish expected performance and operational behavior.
Monitoring	Track performance, safety, fairness and relevant operational signals.
Incident management	Record, investigate and remediate material incidents and near misses.
Change management	Version and review material model/system/data changes.
Revalidation	Define triggers for repeating evaluation after significant change.
Risk review	Reassess residual risk and conditions of use when circumstances change.
Governance record	Maintain decision, evidence and accountability records.

- Performance deterioration against baseline.
- Model or data drift where relevant.
- Changes to operating conditions, devices, sensors or connectivity.
- Material changes to model, architecture, data source or workflow.
- Changes to intended use or target population.
- Safety incidents, near misses or emerging harm patterns.
- Emerging fairness or subgroup performance concerns.
- Significant user or operational feedback indicating changed risk.
- Changes in applicable customer, jurisdictional or regulatory requirements.

15. Vendor Evidence Package & Due-Diligence Questions

The following evidence package converts the narrative guidance into a practical procurement room. Vendors should keep the evidence current, versioned and traceable to the AI system being assessed.

Evidence package	What to provide
System profile	Product, version, intended use, users, environment and dependencies.
Risk assessment	Risk tier, rationale, potential harms and controls.
Population evidence	Development, validation and testing population characteristics.
Evaluation results	Performance, usefulness, usability, fairness, safety and reliability results as applicable.
Model/System Card	Purpose, limitations, population, evaluation and known risks.
Human oversight	Roles, review thresholds, escalation and override mechanisms.

Evidence package	What to provide
Real-world validation	Pilot or deployment evidence representative of intended use.
Privacy & security	Data flows, privacy practices, security controls and assessments.
Monitoring plan	Metrics/signals, frequency, thresholds and escalation.
Governance	Change management, incident response, revalidation and accountability.
Evidence status	Version, date, owner, limitations and open gaps.

Vendor Due-Diligence Questions

- Q What is the intended use, primary user and out-of-scope use?
- Q Which older-adult populations were represented during development and testing?
- Q How was performance evaluated across relevant demographic or contextual groups?
- Q What safety testing was performed, including false alarms and missed events where relevant?
- Q How does the system communicate uncertainty or low confidence?
- Q Who is responsible for interpreting and acting on AI outputs?
- Q What human override and escalation mechanisms are available?
- Q Can you provide an AI/System Model Card and document known limitations?
- Q What real-world or pilot validation evidence is available?
- Q How are performance, drift, incidents and material changes monitored?
- Q What privacy and security controls and assessments are available?
- Q What governance process applies to material changes and ongoing operation?

Vendor Trust: Red Flags

- No documented AI governance or assurance process.
- Claims of safety, fairness or accuracy without supporting evidence.
- No clear intended-use or out-of-scope statement.
- No documented risk classification or rationale.
- Limited or unclear representation of relevant older-adult populations.
- No meaningful subgroup/fairness evaluation where relevant.
- No evidence of safety/reliability testing appropriate to the use case.
- No defined handling of low-confidence or uncertain outputs.
- No clear human oversight, escalation or override mechanism where appropriate.
- No real-world validation evidence for the intended environment.
- Unclear data collection, use, retention or sharing practices.
- No clear privacy/security evidence.
- No post-deployment monitoring or change-management process.
- Material limitations, incidents or known risks are not transparently disclosed.

16. Readiness Decision & Reassessment

Readiness should be a documented decision based on the evidence available for the intended use and risk tier. A completed checklist or published documentation is not, by itself, an approval. The decision should state what is known, what remains unresolved and what conditions apply.

Decision	When it is appropriate
READY	Required evidence is available, material risks are addressed, and deployment controls are in place.
READY WITH CONDITIONS	Deployment may proceed only with documented limitations, controls, monitoring or remediation conditions.
NOT READY	Mandatory evidence is missing, material risks remain unresolved, or required safeguards are not in place.

Decision record field	Record
AI system / version	

Decision record field	Record
Intended use	
Risk tier	
Assessment date	
Reviewing functions / committee	
Decision	READY / READY WITH CONDITIONS / NOT READY
Conditions / restrictions	
Residual risks	
Evidence gaps	
Accountable decision-maker	
Next reassessment / trigger	

17. Practical Implementation: Build an Evidence Room

For a developer or vendor, the most useful implementation is a structured evidence room in which every material assurance claim can be traced to an artifact. This makes the checklist easier to use and allows the same package to support procurement, pilot review and ongoing governance.

Folder / evidence area	Suggested contents
01_Product	System profile, architecture, version history, intended use.
02_Risk	Risk assessment, risk tier, harms, controls and residual risk.
03_Data	Data sources, population representation, quality and governance.
04_Evaluation	Test plans, metrics, results, subgroup/fairness analysis.
05_Safety	Failure modes, safety testing, incident scenarios and mitigations.
06_Human_Oversight	Roles, escalation, override and operating procedures.
07_Validation	Pilot results, real-world validation and workflow evidence.
08_Privacy_Security	Data flows, privacy, security assessments and controls.
09_Governance	Policies, change management, incident response and approvals.
10_Monitoring	Baselines, metrics, thresholds, monitoring and reassessment records.
11_Framework_Mapping	NIST, ISO, CHAI and other applicable mappings.
12_Decisions	Readiness decisions, conditions, residual risk and review records.

Central resource

The AgeTech AI Assurance & Trustworthiness resource center can serve as the central destination for the detailed AgeTech Best Practice Guide, Testing & Evaluation Framework, Applied Model Card, procurement guidance, vendor evidence matrix, real-world validation guidance, continuous monitoring guidance and governance responsibilities.

Executive Summary

A strong AgeTech AI assurance case is a connected story: the system has a clearly defined purpose; the risk is understood; the intended older-adult population and environment are represented; testing addresses performance, fairness, usability, safety and reliability; humans remain accountable for appropriate decisions; evidence is transparent and versioned; real-world validation supports the intended workflow; privacy and security are addressed; and post-deployment monitoring and change management are in place.

The standard to aim for

Be able to answer, with evidence: What is this AI intended to do? Who is it for? What could go wrong? What did you test? With whom and under what conditions? What did you learn? Who remains accountable? What happens after deployment?