

# NCCHCA Annual Primary Care Conference: Moving Health Forward

## Operationalizing Responsible, Ethical, and Equitable AI at a Large, Multi-State Health System

June 5, 2026 • Washington Duke Inn & Golf Club (Durham, NC)



**Justin Kramer, PhD, MAT**

Assistant Professor  
Wake Forest University School of Medicine

Core Faculty, Maya Angelou Research Center for Health Communities

Evaluation Consultant, Advocate Health AI Studio

# Operationalizing Responsible AI in Healthcare



## **Need for Operational AI Governance**

Responsible AI in healthcare requires frameworks that translate safety, equity, and accountability into daily decisions.

## **Risk Assessment and Decision Making**

Distinguishing between low-risk automation and high-risk clinical decision support guides review and deployment processes.

## **Patient and Clinician Trust**

Maintaining trust requires integrating patient perspectives and transparent AI communication in healthcare delivery.

## **Pilot Testing and Emerging Challenges**

Rigorous pilot studies and preparation for agentic AI ensure safe, scalable adoption in complex healthcare environments.



# Presentation Overview and Structure



## **Practical AI Governance for Healthcare**

Focus on operationalizing responsible, ethical, and equitable AI in complex healthcare organizations.

## **Importance of Patient Engagement**

Patient feedback is integrated into governance and communication to build trust and equity in AI use.

## **Real-World AI Application**

Case study on a voice-enabled AI tool for hypertension management demonstrating governance-informed evaluation.

## **Future AI Challenges**

Addressing emerging autonomous AI systems and evolving governance to manage agentic AI risks.



# **Developing an Enterprise AI Evaluation and Governance Framework:**

Framework for the Appropriate Implementation and Review of Artificial Intelligence (FAIR-AI)



# Rationale for Creating FAIR-AI

## **Addressing Governance Gaps**

FAIR-AI was created to fill gaps in enterprise AI governance with consistent, operational review mechanisms tailored for AI complexities.

## **Operational Framework Design**

Designed as an operational guide, FAIR-AI supports scalable, flexible AI decisions integrated into existing risk management.

## **Patient Safety and Trust**

FAIR-AI emphasizes alignment with patient safety, ethical standards, equity, privacy, and trust in AI healthcare applications.

## **Standardized Review Process**

The framework provides a transparent and consistent AI review process, fostering trust among clinicians, patients, and teammates.



# Stakeholder Informed Development Process



## **Stakeholder Engagement Phases**

Two waves of interviews involved diverse stakeholders to identify priorities (pre-development) and evaluate the AI governance framework (post-acceptance).

## **Multidisciplinary Design Session**

Experts in ethics, compliance, legal, cybersecurity, clinical care, and data science collaborated to design the FAIR-AI.

## **Iterative Framework Development**

The governance framework is treated as a living system that evolves through continuous feedback and adaptation.



# Pre-Development Stakeholder Priorities

## **Balancing Risk and Benefit**

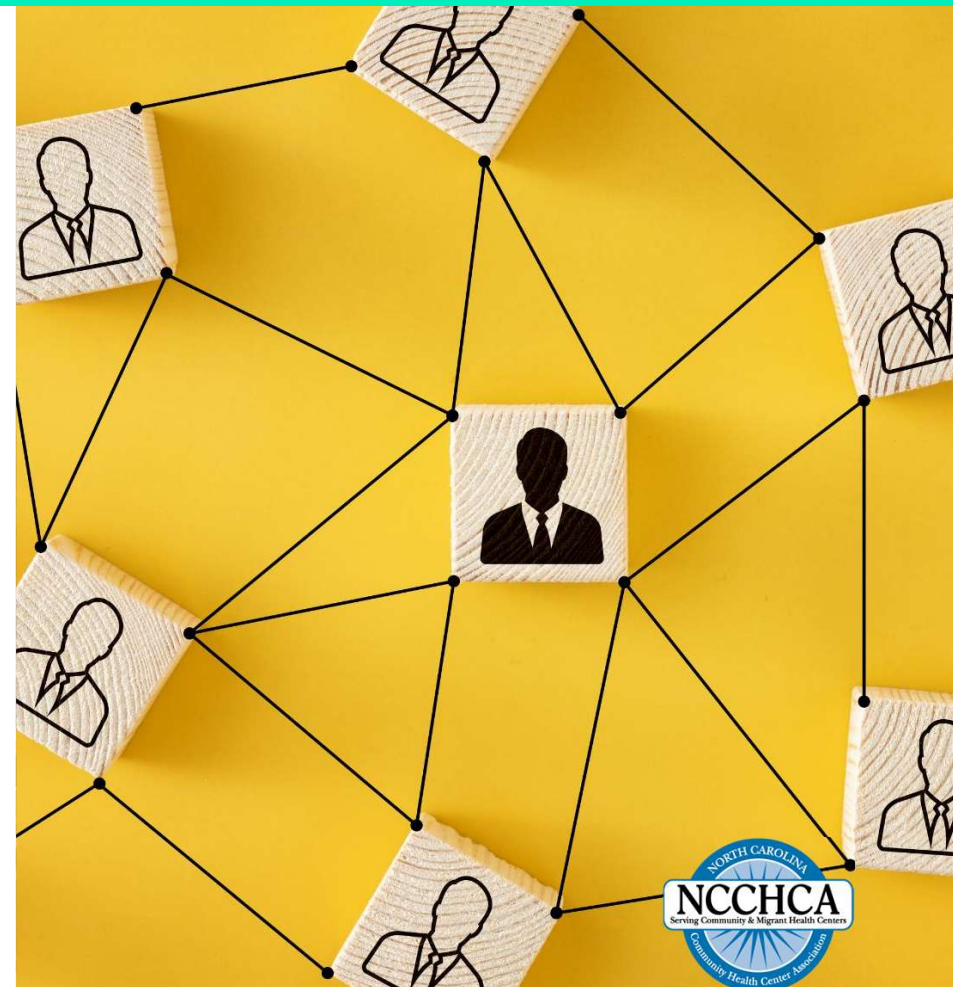
Stakeholders stressed balancing risk tolerance with potential benefits by using a risk-stratified AI governance approach (e.g., low, medium, high, intolerable risk)

## **Human Oversight Necessity**

Direct human oversight (e.g., human-in-the-loop) is essential to support clinical AI and prevent automation bias and accountability loss.

## **Streamlined Review for Low-Risk AI**

Expedited evaluation pathways help ensure timely adoption of low-risk AI tools while maintaining thorough review of high-risk applications.



# Post-Acceptance Evaluation of FARI-AI

## Seven central elements were identified as being essential to successful FAIR-AI integration:

- (1) Transparent and consistent reviews
- (2) Timely evaluations
- (3) Ongoing solution monitoring
- (4) Iterative framework refinement
- (5) Alignment with institutional priorities and regulatory standards
- (6) Multi-modal teammate education
- (7) Diverse patient dissemination efforts

## Transparency Builds Trust

Clear and consistent review processes enhance trust and usability by explaining AI decision-making and approvals.

"The more transparent we can turn the FAIR-AI process and why you made decisions the way you did, the better we uplift everybody's understanding and knowledge."

- Executive Leader



# Implemented FAIR-AI

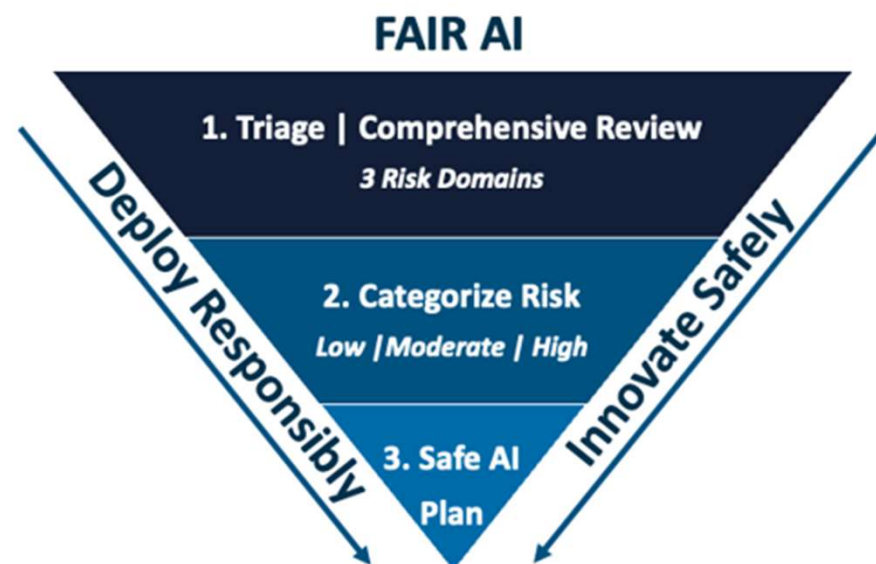
## Four Guiding Principles:

- (1) Equitable and Ethical
- (2) Valid and Reliable
- (3) Transparent
- (4) Impactful

### FAIR-AI Dissemination

A practical framework for appropriate implementation and review of artificial intelligence (FAIR-AI) in healthcare  
([npj Digital Medicine](#))

Leveraging stakeholder engagement to develop and evaluate a responsible artificial intelligence framework at a large, multi-state health system  
([Health Policy and Technology](#))



| Risk Category | Data Science Team                                             | AI Governance Committee                                                                                                      |
|---------------|---------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------|
|               | Evaluation                                                    | Escalation Determination                                                                                                     |
| Low           | All low risk criteria passed                                  | n/a                                                                                                                          |
| Medium        | No high risk criteria on in-depth review                      | n/a                                                                                                                          |
| High          | Any high risk criteria, and/or supporting evidence is lacking | Proceed under high risk conditions<br>Proceed as high risk - pilot or research required<br>Do not proceed - intolerable risk |

# Developing Patient-Facing AI Information, Disclosures, and Communications



# Why Patient Perspectives Matter for AI Governance



## **Patient Trust and Communication**

Patient trust depends on clear explanations and justifications of AI use in healthcare settings.

## **Patient Input Shapes Governance**

Patient feedback can meaningfully inform governance priorities like human oversight, clinician autonomy, and accountability.

## **Aligned Communication and Governance**

Integrating patient perspectives ensures communication effectively informs about AI review and monitoring processes.

## **Promoting Trust and Equity**

Engaging patients fosters trust, equity, and informed involvement in AI-supported healthcare decisions.



# Background on Patient-Facing AI Information Development

Expanding upon the FAIR-AI, and in collaboration with enterprise leaders – virtual health, compliance, patient experience, MaCX – we developed and tested patient-facing AI information.

- **Our approach to AI use** (general AI information)
- **Information on specific, highly-visible AI solutions**
  - Ambient Digital Scribe (ADS)
    - An ambient digital scribe adopted in the southeast region in June 2023; integrated into Epic in June 2024
  - In-patient Monitoring System (IMS)
    - Smart Hospital / “Room of the Future”
    - Uses cameras and microphones throughout the hospital, including in patient rooms
    - Integrated AI monitoring to identify risks and enhance monitoring of sick patients



# Project Goals and Methods

## Primary Goal:

To understand how to best inform patients about our responsible use of AI.

- What patients want to know about how our health system uses AI
- How patients want to receive that information
  - Where – e.g., website, flyer, or pamphlet
  - Method – e.g., videos, pictures, text
- How to deliver information for visible AI solutions to patients
  - IMS
  - ADS

## Methods:

- Conducted multiple rounds of patient interviews pertaining to AI:

Phase I (phone) – obtain feedback prior to FAIR-AI development

Phase II (phone) – identify needs; draft patient-facing AI materials

Phase III (in-person) – obtain feedback on developed AI content

- Website (mock-up)
- V1 IMS Flyer (for in-hospital use)

NCCHCA • Annual Primary Care Conference

|                             | Phase 1:<br>FAIR-AI<br>Development | Phase 2:<br>FAIR-AI Patient<br>Education | Phase 3:<br>Patient Education<br>Refinement | TOTALS              |
|-----------------------------|------------------------------------|------------------------------------------|---------------------------------------------|---------------------|
| <b>When</b>                 | Spring 2024                        | Summer 2025                              | Fall 2025                                   |                     |
| <b>Market</b>               |                                    |                                          |                                             |                     |
| -- AH Charlotte             | 5 (100%)                           | 1 (20%)                                  | 0 (0%)                                      | 6 (33%)             |
| -- AH Wake Forest           | 0 (0%)                             | 3 (60%)                                  | 4 (50%)                                     | 7 (39%)             |
| -- Midwest                  | 0 (0%)                             | 1 (20%)                                  | 4 (50%)                                     | 5 (28%)             |
| <b>Sex</b>                  |                                    |                                          |                                             |                     |
| -- Female                   | 1 (20%)                            | 3 (60%)                                  | 6 (75%)                                     | 4 (56%)             |
| -- Male                     | 4 (80%)                            | 2 (40%)                                  | 2 (25%)                                     | 6 (44%)             |
| <b>Race</b>                 |                                    |                                          |                                             |                     |
| -- Black or AA              | 1 (20%)                            | 2 (40%)                                  | 0 (0%)                                      | 3 (17%)             |
| -- White                    | 4 (80%)                            | 3 (60%)                                  | 6 (75%)                                     | 13 (72%)            |
| -- Other                    | 0 (0%)                             | 0 (0%)                                   | 2 (25%)                                     | 2 (11%)             |
| <b>Ethnicity</b>            |                                    |                                          |                                             |                     |
| -- Non-Hispanic             | 5 (100%)                           | 5 (100%)                                 | 7 (88%)                                     | 17 (94%)            |
| -- Hispanic                 | 0 (0%)                             | 0 (0%)                                   | 1 (13%)                                     | 1 (6%)              |
| <b>Mean Age<br/>(range)</b> | 51.0 years<br>23 - 82              | 54.4 years<br>19 - 81                    | 48.9 years<br>24 - 75                       | 51 years<br>19 - 82 |
| <b>Education</b>            |                                    |                                          |                                             |                     |
| -- HS Diploma               | 0 (0%)                             | 1 (20%)                                  | 0 (0%)                                      | 1 (6%)              |
| -- Some College             | 0 (0%)                             | 1 (20%)                                  | 1 (13%)                                     | 2 (11%)             |
| -- Assoc. or 2yr            | 0 (0%)                             | 0 (0%)                                   | 1 (13%)                                     | 1 (6%)              |
| -- Bachelor's+              | 5 (100%)                           | 3 (60%)                                  | 6 (75%)                                     | 14 (78%)            |

# Phase I and II Findings

## What: Patients want to know...

- (1) How AI is being used at Advocate
- (2) What rules govern the use of AI
- (3) Why specific solutions were adopted
- (4) What “credentials” does the AI have
- (5) Who oversees adopted solutions for safety
- (6) Where is the data obtained by AI stored

## When, Where, and How findings exist too...

“The first thing I'd want to know is who oversees the AI process from the hospital standpoint. Are they medical people, or the IT people, or are they other folks that come with the AI?”

– West NC Patient

“Is there some certification, or something that gives you confidence that it really is state of the art. [...] Has it been certified, who certified it, and where was it tested, what does that look like?”

– South NC Patient



## Phase II Findings

### IMS Concerns:

- (1) A/V monitoring in patients' rooms was the primary issue, not the AI
- (2) Audio created more concern for folks around privacy
- (3) Hacking or improper access to A/V feeds was a major worry
- (4) Patients often felt A/V use would reduce in-person attention by the care team

### ADS:

**Patients were overall very comfortable with ADS being used..... when accessed via health system recording devices.**

- (1) Phones can be hacked, lost, or left unsecure and have increased risk of data breaches, which might jeopardize data privacy

**Many more findings on What, When, How, and Where**

**IMS:** “Personal discussions in a hospital room often take place, and nobody in that hospital has any reason to want to hear that or need to hear that.”

– West NC Patient

**IMS:** “It sounds like you’re going to take the human part out of it. [...] I just want to make sure that I’m getting the quality of care that you normally get when you go into the hospital.”

– West NC Patient

**ADS:** “Who is going to be protecting their phones with that app on it? Would that be secure?”

– West NC Patient



## Phase III - Patient and Provider Experiences

"I'm still in the phase where AI just concerns me in general. I feel like I trust the provider a lot more than I trust the AI system. So, I'm like, oh gosh, how much is the provider relying on this AI?"

- West NC Patient

"[What's] the overarching philosophy on the integration of AI into how the hospital provides care, like vis-a-vis humans."

- West NC Patient

**Patients had multiple questions pertaining to how AI impacts both patients and providers**

Specifically, patients wanted the provided information...

- (1) ...to focus more on patient safety and experience, as opposed to institutional efficiencies and saving the care team time.
- (2) ...to give more specific details on how AI interacts with providers and what tasks it performs.
  - Most notably: what is AI's role is diagnosing patients and making treatment recommendations?
- (3) ...to clearly state Advocate Health's considerations when adopting new AI solutions.
- (4) ...to address existing concerns about how AI might impact patient – [human] provider relationships.



## Phase III - Concerns About Data Security and Privacy

Does the use of AI change the protections of your data in any way? And if so, we need to detail exactly how.”

- West NC Patient

“The cameras and microphones are always on—that’s terrifying. I don’t like this at all. [...] ‘Always watching over you’ feels Big Brother.”

- West NC Patient

“The inability to opt out—some people are probably going to have big problems with that.”

- Midwest (WI) Patient

**Patients wanted more clarity about how adopted AI solutions impact data security and raised multiple concerns pertaining to privacy, consent, and agency.**

Specifically, patients...

- (1) ...desired more information on how / if AI impacts established data safety and security protections.
- (2) ...raised strong privacy concerns about the use of always-on cameras and microphones.
- (3) ...expressed discomfort with the inability to opt out of AI solutions, comparing it to a loss of control.
- (4) ...noted confusion about multiple critical parts of the AI solutions, particularly about IMS’ Privacy Mode



## Phase III - Ongoing Monitoring and Human-in-the-Loop

**Patients found statements on direct human oversight to be reassuring and wanted high-levels of transparency pertaining to AI use.**

Specifically, patients wanted...

- (1) ...concise language and specifics on how tools are tested and monitored.
- (2) ...concrete examples of where AI is used across the care continuum and why those decisions were made.
- (3) ...transparency as to when AI is being used in their care and what benefits it provides.

"I want a really specific, detailed understanding of what exactly is being done."

- West NC Patient

"I feel like it makes it clear that these are the ways the team uses it, but the AI is not making decisions on its own."

- Midwest (WI) Patient



## Phase III - Accessibility and Design for Different Populations

**Consider different needs and concerns pertaining to AI for various patient groups, such as elderly or non-English-speaking patients.**

Specifically, patients wanted...

- (1) ...content and information design to meet the needs of older and/or internet non-users
- (2) ... information to be accessible via digital and non-digital modes at various touchpoints
- (3) ...translations of the information in Spanish and other languages

“As long as Aurora makes sure that elderly patients have what they need to make that transition—because I’ve seen people terrified by this [AI].”

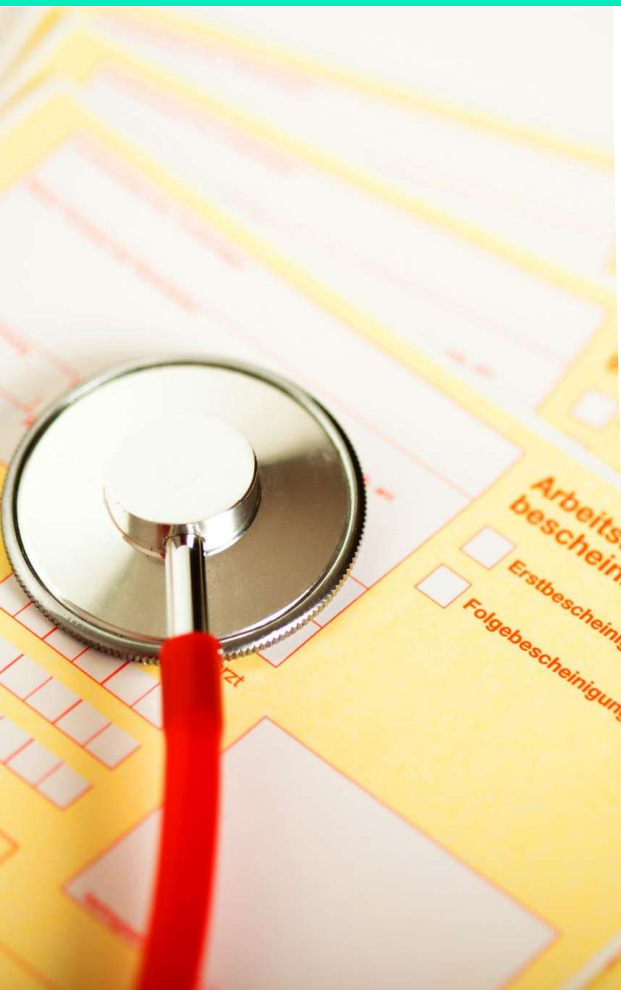
- Midwest (WI) Patient

“This, it has to be in different languages, right? [...] So, especially Spanish and the languages for the area.”

- Midwest (WI) Patient



# Design Principles for Patient-Facing AI Materials



## **Use of Plain Language**

Plain language with minimal jargon reduces anxiety and builds trust in AI communication materials.

## **Layered Information Approach**

Providing brief summaries with optional detailed paths balances clarity and completeness for diverse patient needs.

## **Accessibility and Equity**

Materials must be accessible in multiple languages and formats to serve diverse patient populations effectively.

## **Opportunities for Patient Questions**

Enabling patients to ask questions through clinicians or support roles fosters transparency and trust.

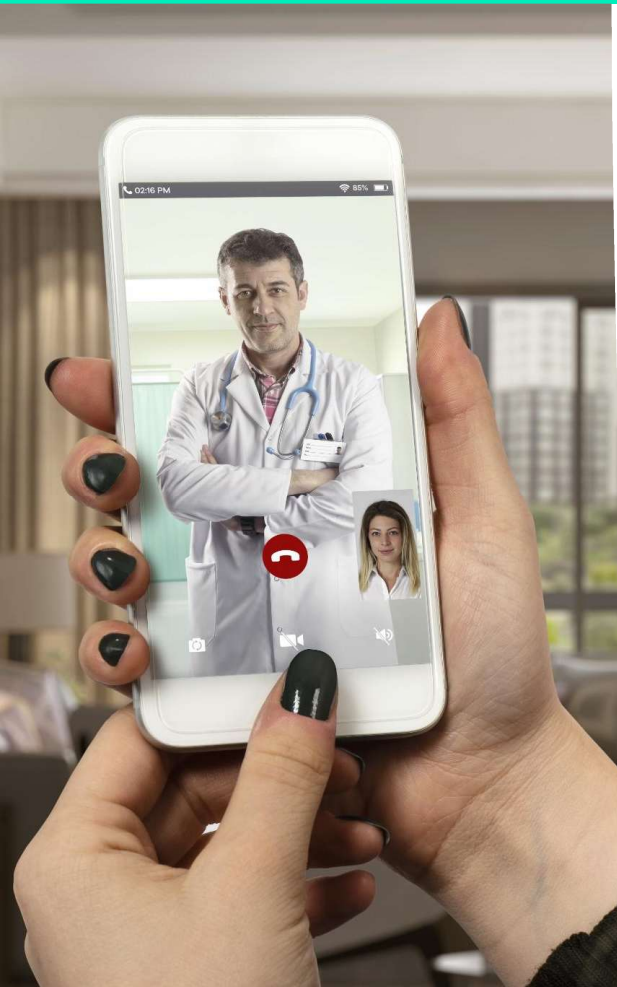


# **Evaluating High-Risk AI:**

## Hypertension Management Case Example



# Overview of the Hypertension AI Pilot



## **AI-Powered Hypertension Management**

A voice-enabled AI agent (LLM-HTN) provided education, blood pressure monitoring, and medication support through scheduled calls over four weeks.

## **High-Risk Use Case Evaluation**

The pilot test addressed safety and trust concerns by directly engaging consented patients for LLM-led interactions around chronic disease management.

## **Feasibility and Study Design**

Conducted as a proof-of-concept study focusing on feasibility, acceptability, appropriateness, and safety beyond clinical outcomes.

## **Integration with Care Workflows**

Embedded AI pilot within existing care teams and governance to test real-world AI oversight and collaboration.



# Overview of the Hypertension AI Pilot

## Study Goals:

- **(1)** Design a multi-week intervention using an LLM-HTN agent to assist patients with outpatient hypertension management
- **(2)** Conduct a small feasibility test / **proof-of-concept** assessment around using an AI agent for outpatient hypertension management

## AI Agent Development:

- **AI agent was developed via collaboration between the study team, enterprise leaders, and the LLM-HTN development team**
  - Development began late Aug 2025; agent was ready for testing early Dec 2025
  - Iterating clinician red team testing process guided revisions to call agenda and agent function
  - Formal User Acceptance Testing (UAT) for each of the four calls prior to study launch
- **4 weekly AI agent calls, with BP and med reconciliation / adherence assessed at each**
  - **Week 1:** BP measurement technique, general HTN barriers, symptoms, why BP control is important
  - **Week 2:** Diet and nutrition, alcohol and substances, goal setting, and motivational interviewing (MI)
  - **Week 3:** Review goals and barriers, budget friendly foods, recipes, and resources
  - **Week 4:** Exercise and lifestyle, goal setting, MI, and suggestions
- **Escalation protocols were established to handle emergent patient needs during AI agent calls**
  - Developed and monitored by enterprise operational leaders and care teams, not the study team
  - **High-level** received a warm hand-off to a care team
  - **Medium-level** were reviewed within 4 hours by centralized nurses
  - **Low-level** were noted in the summary note, but no follow-up was initiated by the AI agent



# Overview of the Hypertension AI Pilot

## 15 Patients were enrolled in this study.

- N=10 AHWFB Family Medicine
- N=5 SHVI HTN Boot Camp

## 13 / 15 patients completed entire 4-call series

## 2 patients did not complete a call

- 53 yrs, Black/AA, Male, AHWFB  
Attempted one call, but quit, due to the AI not understanding them during med rec.
- 64 yrs, Black/AA, Female, AHWFB  
Never engaged the AI, unreachable for the PI

|                                            | AHWFB     | SHVI     | Total     |
|--------------------------------------------|-----------|----------|-----------|
| <b>Total Participants</b>                  | 10 (100%) | 5 (100%) | 15 (100%) |
| <b>Mean Age</b>                            | 55 years  | 69 years |           |
| <b>(range)</b>                             | 35 - 71   | 51 - 82  |           |
| <b>Sex</b>                                 |           |          |           |
| -- Male                                    | 7 (70%)   | 3 (60%)  | 10 (67%)  |
| -- Female                                  | 3 (30%)   | 2 (40%)  | 5 (33%)   |
| <b>Race</b>                                |           |          |           |
| -- White                                   | 4 (40%)   | 5 (100%) | 9 (60%)   |
| -- Black or AA                             | 6 (60%)   | 0%       | 6 (40%)   |
| <b>Ethnicity</b>                           |           |          |           |
| -- Non-Hispanic                            | 10 (100%) | 5 (100%) | 15 (100%) |
| <b>Access to a Phone for Calls / Texts</b> | 10 (100%) | 5 (100%) | 15 (100%) |

|                                            | AHWFB   | SHVI    | Total    |
|--------------------------------------------|---------|---------|----------|
| <b>Using Technology for Health Reasons</b> |         |         |          |
| -- Daily                                   | 4 (40%) | 2 (40%) | 6 (40%)  |
| -- Weekly                                  | 2 (20%) | 1 (20%) | 3 (20%)  |
| -- Rarely                                  | 1 (10%) | 2 (40%) | 3 (20%)  |
| -- Never                                   | 3 (30%) | 0 (0%)  | 3 (20%)  |
| <b>Prior AI Use for Health Reason</b>      |         |         |          |
| -- Yes                                     | 2 (20%) | 2 (40%) | 4 (27%)  |
| -- No                                      | 8 (80%) | 3 (60%) | 11 (73%) |
| <b>Confidence Learning New Technology</b>  |         |         |          |
| -- Very Confident                          | 5 (50%) | 2 (40%) | 7 (47%)  |
| -- Somewhat Confident                      | 2 (20%) | 3 (60%) | 5 (33%)  |
| -- Neutral                                 | 1 (10%) | 0 (0%)  | 1 (7%)   |
| -- Somewhat Unconfident                    | 1 (10%) | 0 (0%)  | 1 (7%)   |
| -- Very Unconfident                        | 1 (10%) | 0 (0%)  | 1 (7%)   |



# Data Collection and Evaluation

## Clinician Reviews

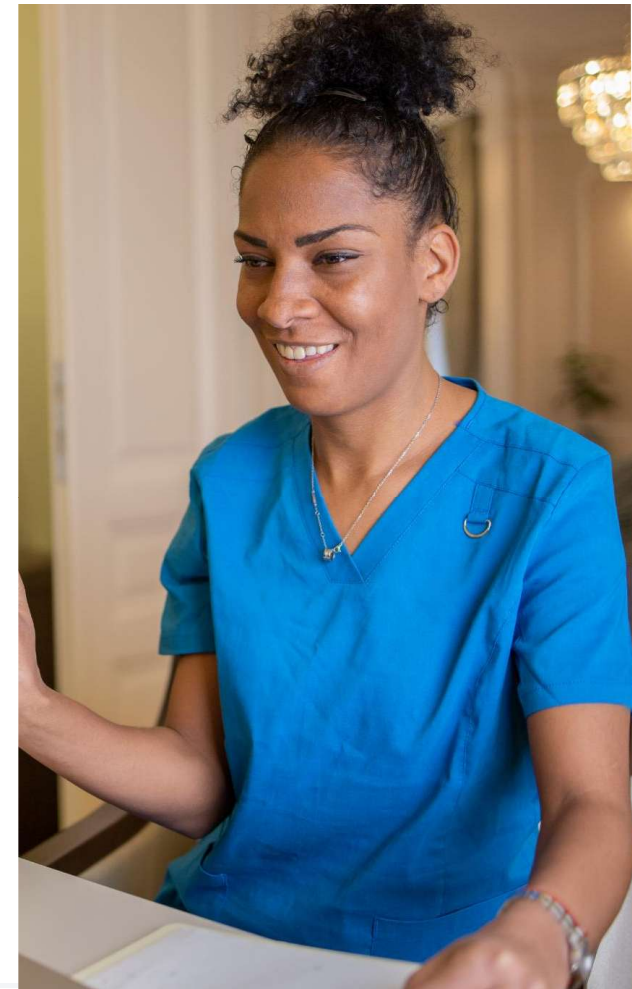
1. Care team evaluations in REDCap for all AI agent – patient calls (expected N=120; actual N=102)

## Data from LLM-HTN System

2. Patient call evaluations and Net Promoter Score (expected N=60; actual N=48)
  - Collected via the AI agent at the end of the call
3. LLM discrete data collection (expected N=60; actual N=51)
  - BP telemonitoring, medication adherence
4. AI agent call data (N = 186)
  - Includes patient call backs too
  - Data elements: when, duration, missed calls, PT call backs, etc.
5. AI Agent – patient call transcripts (expected N=60; actual N=51)

## Research Specific Activities

6. Pre- and post-trial surveys with patients (expected N=15; actual N=15 / N=13)
7. Semi-structured interviews with patients (N=10) and care team members (N=3)
  - If patients did not complete a call, the AI agent would call and text back 2 days later



# Patient Calls with LLM-HTN

**(1)** At the end of each call, the AI agent asked the patient to...

a) rate the call

b) if they would recommend the call (NPS)

**(2)** Calls 1 -3 received very high marks, while Call 4 observed a drop off

|        | Completed Calls | Call Evaluation* |                     | Net Promoter Score* |                |                   |               |                 |
|--------|-----------------|------------------|---------------------|---------------------|----------------|-------------------|---------------|-----------------|
|        |                 | Scores Collected | Mean Score (1 - 10) | Mean NPS            | Calculated NPS | Promoter (9 - 10) | Passive (7-8) | Detractor (1-6) |
| Call 1 | 13              | 13               | 8.9                 | 9.5                 | 76.9           | 11                | 1             | 1               |
| Call 2 | 12**            | 10*              | 9.6                 | 9.9                 | 100            | 10                | 0             | 0               |
| Call 3 | 13              | 13               | 9.2                 | 9.5                 | 76.9           | 10                | 3             | 0               |
| Call 4 | 13              | 12*              | 7.3                 | 8.7                 | 75.0           | 9                 | 3             | 0               |

\*Call eval and NPS were collected by the AI agent. When patients requested a shorter call, they may have been skipped.

\*\*AI agent erroneously skipped a patient's Week 2 call.

**(3)** Improved BP measurements

a) all patients had a last EHR BP  $\geq$  140/90 at time of enrollment

**(4)** Decreased medium-level+ escalations as the 4-call series stretched on

|        | Mean Blood Pressure Measurement |    | Identified Care Escalations |        |      |
|--------|---------------------------------|----|-----------------------------|--------|------|
|        |                                 |    | Low                         | Medium | High |
| Call 1 | 138                             | 83 | 2                           | 3      | 0    |
| Call 2 | 127                             | 72 | 1                           | 2      | 0    |
| Call 3 | 131                             | 79 | 1                           | 0      | 0    |
| Call 4 | 132                             | 81 | 3                           | 0      | 0    |

# Care Team Transcript Reviews

## (1) 51 AI agent – patient calls were completed

- AI agent skipped 1 patients' week 2 call

## (2) 102 total transcript reviews (2 per call)

- RN and PCP reviewer for Family Medicine
- Pharmacist and Cardiologist for Cardiology

## (3) Call audio was only reviewed by initial reviewer

- Transcript Accuracy was assessed once per call

\*Mean values presented; \*\*Scale: 1 poor – 10 excellent

|        | Transcript Accuracy<br>(RN/Pharm) | AI Summary Note Accuracy |             | Call Completed its Objectives |             | Confidence in AI to Perform Call Tasks |             |
|--------|-----------------------------------|--------------------------|-------------|-------------------------------|-------------|----------------------------------------|-------------|
|        |                                   | RN / Pharm               | PCP / Cards | RN / Pharm                    | PCP / Cards | RN / Pharm                             | PCP / Cards |
| Call 1 | 6.8                               | 7.8                      | 9.2         | 7.2                           | 8.6         | 5.5                                    | 8.5         |
| Call 2 | 8.2                               | 7.9                      | 9.4         | 8.2                           | 8.9         | 6.2                                    | 8.4         |
| Call 3 | 7.7                               | 8.4                      | 9.7         | 8.1                           | 9.4         | 6.7                                    | 9.6         |
| Call 4 | 7.2                               | 8.2                      | 9.8         | 7.8                           | 9.5         | 7                                      | 9.9         |

### Key Takeaways:

#### (1) Confidence in AI performance climbed with each call

#### (2) Summary note accuracy and AI successfully completing call objectives also had an upward trend

### Safety Occurrences:

#### (1) All reviewers assessed safety. Across the 102 reviews...

- 8 times where the AI agent provided “inaccurate information”
- 4 times where the AI agent provided “unsafe information”

**Completed Call:** “To what extent did the AI agent complete all the call objectives and perform to a satisfactory level?”

**Confidence:** “How confident are you that the AI agent could be used in the place of a care team member to assist patients with the tasks of this call?”

#### (2) Error Severity

- All 12 instances were scored; **1 critical – 10 harmless**
- Mean Seriousness 7.5 || Median Seriousness: 8
- Only 2 reviews included an issue < 6
- From the same patient call, RN and PCP both gave a 4



# Patient Post-Trial Evaluations

\*N=13 patients\*

## AIM-4 and IAM-4 Scales (Likert):

Strongly Disagree: 1 – Strongly Agree: 5

### Acceptability of Intervention Measure (AIM-4)

|                                           |             |
|-------------------------------------------|-------------|
| -- Mean summed AIM scores (out of 20)     | 16.69       |
| -- <b>Mean average across the 4-items</b> | <b>4.17</b> |
| -- Range of summed scores                 | 12 - 20     |

**AIM-4:** Assesses how agreeable, satisfactory, or appealing an intervention is to its intended users.

### Net Promoter Score (NPS)

|                                       |             |
|---------------------------------------|-------------|
| -- <b>Calculated NPS (out of 100)</b> | <b>61.5</b> |
| -- Mean patient NPS                   | 8.92        |
| -- Range                              | 6 - 10      |

**NPS:** Captures users' likelihood of recommending a product or service as an indicator of overall satisfaction and loyalty.

### Intervention Appropriateness Measure (IAM-4)

|                                       |             |
|---------------------------------------|-------------|
| -- Mean summed IAM scores (out of 20) | 17.46       |
| -- <b>Mean average across 4-items</b> | <b>4.37</b> |
| -- Range of summed scores             | 14 - 20     |

**IAM-4:** Evaluates the perceived fit, relevance, or compatibility of an intervention within a given setting or role.

### System Usability Scale (SUS)

|                                |                     |
|--------------------------------|---------------------|
| -- Mean SUS score (out of 100) | 80.38               |
| -- <b>Median SUS score</b>     | <b>80 (A grade)</b> |
| -- Range                       | 57.50 - 97.50       |

**SUS:** Provides a global, standardized assessment of the usability of a system, tool, or technology.



# Post-Trial Interviews

## What Worked Well:

- (1) Human-like voice and tone, often described as pleasant, calm, and encouraging.
- (2) Increased access and call-back flexibility, enabling patients to engage on their own schedule.
- (3) BP-technique coaching and live BP capture during calls
- (4) Behavior support, such as nudges, tailored goal setting, and routine-building, were linked to increased patient activation and perceptions of being supported.
- (5) Increased PT comfort asking the AI questions, perceiving it to be judgement-free and never busy.

## What Didn't Work Well:

- (1) Multiple speech recognition and cadence issues, e.g., trouble understanding accents.
- (2) Many patients reported altering their normal speech cadence to accommodate the AI agent.
- (3) Patients found both the content and the AI agent's manner of engagement repetitious as the calls went on, reducing perceived naturalness and value. This was most prominent in calls 3 and 4.
- (4) Patients reported confusion about ambiguous BP goals, e.g., a BP < 120/80 called "a good start."
- (5) AI-generated call summaries sometimes included misheard (i.e., wrong) information and values.

"Anytime you can give a patient more information, that's good... the education felt substantial in the early calls."

– NC Family Medicine Patient

"It reminded me to stay on a schedule. It reinforced my routine of taking my blood pressure the right way."

– NC Family Medicine Patient

"I didn't feel like I was taking time away from my doctor or nurse... that was a big relief."

– NC Cardiology Patient

# Post-Trial Interviews

## Key **ACTIONABLE** Suggestions

- (1) Refine BP target messaging; state explicit goal ranges. If possible, tailor BP goals to patients.
- (2) Expand variability across calls to reduce redundancy, both for content and AI word choice.
- (3) Improve med list accuracy and med reconciliation processes; prompt patients for their medication lists, rather than asking them if they are taking “X, Y, and Z.”
- (4) Include longer rest times when soliciting BP readings; potentially offer a 5-minute wait option.
- (5) The AI agent should utilize prompts early in the call, e.g., “I’m having trouble hearing, please speak a bit slower/louder,” to help inform patients and avoid recurring issues.
- (6) Allow AI agent tone to be configurable to meet patient preferences and clinical norms.
- (7) Expand call rescheduling to use text messaging; ensure phone #s do not hit caller IDs as spam.
- (8) Provide patients content prior to HAI engagement that details AI agent capacity (e.g., questions that can be asked, content covered) and processes for human review and follow-up when needed.

“I wasn’t sure how deep the AI could go. I didn’t realize I could ask anything. If I had known that earlier, it would’ve made things easier.”

– NC Family Medicine Patient

“If she asked me to get my cuff at the beginning and then we talked for a few minutes, that could be the rest time before taking it.”

– NC Family Medicine Patient

“Overall appropriate and safe, but sometimes the summaries or the agent misheard a medication or blood pressure number.”

– NC Family Medicine RN

## Summary of Main Findings

- Patients generally found the AI agent easy to engage, helpful, and would recommend its use to other patients
- Patients found AI agent-facilitated BP measurement and data sharing to be easy.
- Clinicians were increasingly confident in the ability of the AI agent to assist with hypertension management as the trial continued.
- Clinicians are optimistic about the potential for the AI agent to assist **both** patients and care teams with hypertension management.
  - **Patients:** Improved access to hypertension-focused education and support
  - **Care Teams:** Assistance with BP telemonitoring processes and provision of patient resources
- Improvements to the AI agent are needed, notably around conversation flow, understanding patients, and ensuring engagement does not become overly repetitious.



# **Looking Ahead:** Agentic AI and the Future of Governance



# Emergence of Agentic AI in Healthcare

## **Autonomous AI in Healthcare**

Agentic AI systems perform healthcare tasks with reduced human input, improving efficiency and decision-making.

## **Governance Challenges**

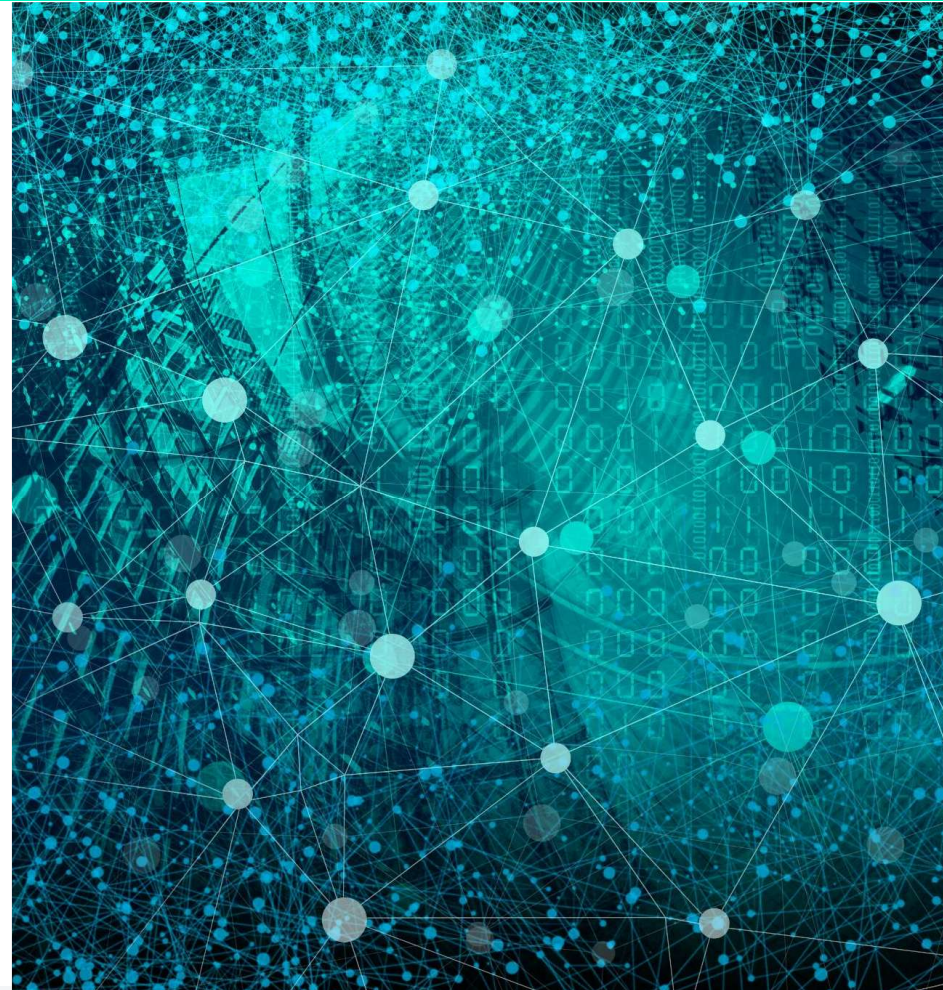
Agentic AI blurs the line between recommendations and actions, raising complex accountability and risk concerns.

## **Limitations of Existing AI Frameworks**

Agentic AI's autonomy challenges existing AI governance frameworks, highlighting the need for updated oversight methods.

## **Proactive Governance Planning**

Similar to FAIR-AI development (2024), our healthcare organization is planning a design session to address governance gaps before widespread agentic AI adoption.



# Future Directions and Key Takeaways



## **Operational AI Governance**

Effective AI governance guides decisions, integrates with workflows, and evolves based on feedback.

## **Patient-Centered Trust**

Building trust requires transparency, safety, and human oversight aligned with patient expectations.

## **Evaluation and Safeguards**

Pilot studies offer insights into AI performance, highlighting where safeguards are necessary.

## **Future-Focused Governance**

Adaptive frameworks are essential as AI evolves to manage new autonomy and risk challenges.



# Thank You!

## Contact Information:

Justin Kramer  
[Justin.Kramer@wfusm.edu](mailto:Justin.Kramer@wfusm.edu)

