

Background Tutorial · April 18, 2026

Medicine and AI

A Self-Study Guide for Attendees

SECTION 1

About This Guide & the Helix Format

VERIFICATION NOTICE

Facts, links, and quantitative claims last verified: **April 11, 2026**. Regulatory counts, citation totals, and organizational metrics may change over time.

This guide is designed for college-educated attendees (MS, PhD, MD, or equivalent) who may be deeply familiar with *either* AI/machine learning *or* medicine/clinical practice — but not both. It provides concise background on the concepts, vocabulary, and research questions you'll need to follow the roundtable discussion productively.

THE HELIX CENTER FORMAT

Helix Center roundtables are **moderated conversations**, not lecture panels. Five experts with different (often opposing) perspectives sit in a semicircle and engage in real-time dialogue for 90 minutes. No PowerPoints, no prepared talks. The audience observes the conversation and joins during a Q&A period. The format is designed to surface *productive disagreement* — and this guide will help you anticipate where those disagreements lie.

How to Use This Guide

- **If you come from medicine:** Start with Section 2 (AI Foundations) to build vocabulary, then jump to Section 5 (Speakers) and Section 6 (Fault Lines).
- **If you come from AI/ML:** Start with Section 3 (Medical Foundations) for clinical context, then read Section 5 and 6.
- **If you're a generalist:** Read straight through — the guide is designed as a narrative arc from foundations to debate.
- **Short on time?** Read Section 5 (speakers) and Section 6 (fault lines) only — ~15 minutes.

SECTION 2

AI/ML Foundations for Medical Professionals

If you're a clinician, researcher, or medical professional, this section provides the AI/ML vocabulary and concepts you'll need to follow the roundtable's technical discussions.

What “AI” Means in This Context

When we say “AI in medicine,” we almost always mean **machine learning (ML)**: algorithms that learn patterns from data rather than following hand-coded rules. The key categories relevant to this roundtable:

- **Supervised learning:** The model learns from labeled examples (e.g., “this tissue slide shows Alzheimer’s pathology”). Most clinical AI today uses supervised learning.
- **Deep learning:** ML using neural networks with many layers. Enables analysis of complex data like images, text, and genomic sequences.
- **Foundation models / Large Language Models (LLMs):** Very large models trained on broad data, then adapted to specific tasks. GPT (text) and ESM (proteins) are generative examples. NYUTron is a BERT-like clinical language model trained on EHR text for predictive tasks — it is not a general-purpose generative chatbot, though it is sometimes loosely called an “LLM.”
- **Computer vision:** AI that analyzes images — in medicine, this means radiology, pathology, dermatology, and retinal imaging.

How Clinical AI Models Work (Simplified)

THE BASIC PIPELINE

1. **Data collection** → EHR records, imaging, genomic data, wearable sensors
2. **Preprocessing** → Cleaning, de-identification, formatting
3. **Training** → The model learns patterns from thousands/millions of examples
4. **Validation** → Testing on held-out data the model hasn’t seen
5. **Deployment** → Integration into clinical workflow (the hardest part — ask Aphinyanaphongs)
6. **Monitoring** → Ongoing surveillance for model drift and bias

Key Concepts You’ll Hear

- **Training data:** The dataset the model learns from. Its composition determines what the model can and cannot do. If the training data underrepresents a population, the model will perform worse for that population (a central theme of Chunara’s work).

- **Model validation:** Testing model performance on new data. Internal validation uses data from the same institution; external validation uses data from different institutions. External validation is the gold standard but rarely done.
- **Sensitivity and specificity:** Sensitivity = how well the model catches true positives. Specificity = how well it avoids false positives. These trade off against each other, and the right balance depends on clinical context.
- **Black box vs. interpretability:** Many deep learning models can't explain *why* they made a prediction. This is a major barrier to clinical trust and regulatory approval.
- **Model drift:** Real-world data changes over time (new diseases, new documentation practices, population shifts). Models that worked well at deployment may degrade. Aphinyanaphongs' team monitors more than 30 models for this.

SECTION 3

Medical & Clinical Foundations for AI Researchers

If you're an AI researcher, engineer, or technologist, this section provides the medical context you'll need to understand the clinical stakes of the roundtable discussion.

How Clinical Decisions Are Made

Clinical medicine is fundamentally about **decision-making under uncertainty**. A clinician integrates:

- Patient history and symptoms (subjective)
- Physical examination (objective but variable)
- Lab results and imaging (quantitative but noisy)
- Published evidence (population-level, may not apply to this patient)
- Clinical experience and intuition (powerful but biased)
- Patient preferences and values (often unspoken)

AI enters this process as an *additional input*, not a replacement for it. The question is: where in this decision chain does AI add value, and where does it introduce risk?

The Electronic Health Record (EHR)

The EHR is the central nervous system of modern hospitals. It stores clinical notes, lab results, medications, imaging, billing codes, and more. For AI researchers, two things to know:

- **EHR data is messy:** It's generated for billing and legal documentation, not for research. Notes contain abbreviations, copy-pasted text, and errors. Lab values have different reference ranges across hospitals.
- **EHR data is biased:** What gets recorded reflects who seeks care, who has insurance, and what clinicians choose to document. Absence of a diagnosis doesn't mean absence of disease.

Key Medical Domains in This Roundtable

FIVE DOMAINS, FIVE SPEAKERS

Digital mental health (Choudhury): Using smartphones and wearables to detect stress, depression, and behavioral changes in real time.

Genomics & protein science (Brandes): Predicting how genetic mutations affect protein function and disease risk.

Neuropathology (Crary): Diagnosing brain diseases by examining tissue under a microscope — now with AI.

Health equity (Chunara): Ensuring AI doesn't worsen existing disparities in who gets what care.

Clinical operations (Aphinyanaphongs): Actually deploying AI tools in a working hospital system.

Why “Deploying AI” in Medicine Is Hard

- **Regulation:** The FDA regulates AI as a medical device (SaMD). Approval pathways are evolving and inconsistent.
- **Liability:** If an AI recommends a wrong treatment, who is responsible? The developer? The hospital? The clinician who followed the recommendation?
- **Workflow integration:** Clinicians are already overwhelmed. Adding another alert or tool can cause “alert fatigue” and be ignored.
- **Trust:** Clinicians won't use a tool they don't trust, and they won't trust a tool they can't understand.
- **Equity:** A model that works well at an academic medical center may fail at a community hospital with different patient populations.

SECTION 4

The Intersection: AI in Medicine Today

This section maps the current landscape — what’s working, what’s not, and where the roundtable speakers sit within it.

Where AI Is Already in Clinical Use

- **Radiology:** As of the FDA’s March 2026 update (covering authorizations through the end of 2025), radiology accounts for roughly 1,104 of about 1,451 total authorized AI/ML-enabled medical devices (~76%). Most assist with detection (e.g., flagging suspicious findings on mammograms or CT scans). *Counts are periodically updated; verify against the FDA list before reuse.*
- **Pathology:** Digital pathology is accelerating, with AI tools for tumor detection, grading, and — as Crary demonstrates — brain aging.
- **Clinical decision support:** EHR-integrated alerts for sepsis, deterioration, and mortality risk (Aphinyanaphongs’ domain).
- **Drug discovery:** AI-assisted molecular design has entered clinical trials, though no AI-discovered drug has yet received full FDA approval.
- **Mental health:** Digital therapeutics and monitoring tools (Choudhury’s domain), including FDA-cleared apps for substance use disorders.

Where AI Is Struggling

- **Generalizability:** Models trained at one institution often fail at another. The “lab-to-bedside” gap remains wide.
- **Equity:** Multiple studies show AI models perform worse for underrepresented populations. Chunara’s work at NYC safety-net hospitals makes this concrete.
- **Regulation:** The FDA’s framework is playing catch-up with continuously learning models. Current rules assume static software.
- **Trust:** Clinician adoption remains low for many tools. The “last mile” from model to clinical impact is the hardest part.
- **Evidence standards:** The field lacks consensus on what constitutes sufficient evidence for clinical AI deployment. RCTs? Prospective studies? Retrospective validation?

CLINICAL AI EVIDENCE & REPORTING STANDARDS

Several reporting frameworks now set the benchmark for evaluating clinical AI evidence:

- **CONSORT-AI / SPIRIT-AI** — for clinical trials involving AI interventions
- **DECIDE-AI** — for early-stage clinical deployment evaluation
- **STARD-AI** — for diagnostic accuracy studies
- **TRIPOD-AI / PROBAST-AI** — for prediction models and risk-of-bias assessment

These standards are curated by the EQUATOR Network and referenced in FDA guidance on AI/ML-based software as a medical device.

Privacy, Consent & Data Governance

GOVERNANCE ESSENTIALS

Clinical AI development raises critical governance questions that go beyond model performance:

- **HIPAA & GDPR:** U.S. and EU frameworks impose strict requirements on how patient data is collected, stored, and shared for AI training and inference.
- **De-identification:** Methods such as Safe Harbor and expert determination aim to protect patient privacy, but re-identification risks grow as datasets are combined.
- **IRB oversight:** Research involving patient data typically requires Institutional Review Board review or documented waivers, even for retrospective studies.
- **Secure enclaves:** Controlled-access computing environments allow AI training on sensitive data without exposing raw records to researchers.
- **Foundation model training:** Training large models on clinical data raises distinct governance and consent questions, particularly when patients did not originally consent to AI use of their records.

The Regulatory Landscape

KEY REGULATORY FACTS

As of the FDA’s March 2026 update (covering authorizations through the end of 2025), the FDA had authorized about 1,451 AI/ML-enabled medical devices, with radiology accounting for roughly 1,104 (~76%). The agency’s 2021 “Action Plan” for AI/ML-based SaMD proposed a framework for continuously learning systems, but implementation is still evolving. The EU AI Act entered into force in 2024, with obligations phasing in over subsequent years; many medical-AI systems are likely to fall into high-risk categories, creating additional compliance obligations alongside existing medical-device rules. No single global standard exists.

Post-Market Surveillance & Drift Monitoring

Clinical AI requires ongoing real-world performance monitoring, drift detection, and update governance. For adaptive systems, regulators have emphasized **predetermined change control plans** and post-market surveillance rather than one-time validation alone. The FDA’s AI/ML SaMD Action Plan outlines the evolving expectations for how deployed models should be monitored and updated over their lifecycle.

SECTION 5

Meet the Five Speakers

Each speaker brings a distinct lens to the question of AI in medicine. Understanding their research and perspective will help you follow the roundtable's dynamics.

THE SENSOR SCIENTIST — DIGITAL MENTAL HEALTH & UBIQUITOUS COMPUTING

Tanzeem Choudhury, PhD

Professor of Information Science, Jacobs Technion-Cornell Institute, Cornell Tech & Cornell University; Co-founder, HealthRhythms

Choudhury is a pioneer in using smartphone and wearable sensor data to monitor mental health. Her lab developed systems that detect stress, depression, and behavioral changes through passive sensing (accelerometer, microphone, GPS, typing patterns) without requiring active user input. She co-founded HealthRhythms to translate this research into clinical tools, and her work on the StudentLife dataset is one of the most cited studies in mobile health.

Questions she'll raise: Can continuous sensing detect mental health crises before they happen? Who has access to this data, and who should? What happens when AI knows you're depressed before you do?

PREPARATION READING

TALK Tanzeem Choudhury — Tracking mental health through mobile sensing (talk)

WEBSITE Cornell Tech faculty page

THE MOLECULAR DECODER — GENOMIC AI & BIOLOGICAL DATA GOVERNANCE

Nadav Brandes, PhD

Assistant Professor, Center for Human Genetics & Genomics, NYU Grossman School of Medicine

Brandes develops protein language models to predict how genetic mutations affect protein function. His 2023 *Nature Genetics* paper predicted effects of all ~450 million possible missense variants in the human genome. He created ProteinBERT (*Bioinformatics*, 2022) and PWAS (*Genome Biology*, 2020), and his collaboration with Jennifer Doudna and Carl June on mitigation of CRISPR-Cas9-induced chromosome loss in clinical T cells was published in *Cell*. He is also an active voice on AI safety and biological data governance, contributing to LessWrong and the AI Alignment Forum and co-authoring a 2026 *Science* article on data governance.

Questions he'll raise: Who controls biological data when AI can decode it? What are the long-term safety implications of increasingly capable AI? How should we govern genomic data in an age of foundation models?

PREPARATION READING

PAPER Genome-wide prediction of disease variants (*Nature Genetics*, 2023)

PAPER ProteinBERT (*Bioinformatics*, 2022)

WEBSITE Incrementally — Brandes' blog on AI, biology, and safety

THE DIAGNOSTIC ARCHITECT — AI NEUROPATHOLOGY & BRAIN AGING

John F. Crary, MD, PhD

Professor of Pathology, Molecular and Cell-Based Medicine, Neuroscience, and AI & Human Health; Vice Chair for AI, Dept. of Pathology, Icahn School of Medicine at Mount Sinai

Crary is a board-certified neuropathologist who created HistoAge (the first AI to predict biological brain age from histopathology slides) and co-developed an AI system detecting alpha-synuclein pathology in peripheral biopsy tissue with reported sensitivity and specificity near 99%. His 2014 definition of PART (Primary Age-Related Tauopathy) is widely cited (over 1,000 citations, Google Scholar, accessed April 2026) and changed how the field classifies brain aging. He directs a major Alzheimer's brain bank.

Questions he'll raise: When AI can read a slide with near-perfect accuracy, what is the role of the human pathologist? How do we validate AI diagnostic tools for tissues we've never seen before? What does "trust" mean when the AI sees something the human can't?

PREPARATION READING

PAPER PART consensus definition (Acta Neuropathologica, 2014)

GUIDE Transforming Diagnostics: Advances in Digital Pathology (review)

THE EQUITY ARCHITECT — ALGORITHMIC FAIRNESS & POPULATION HEALTH

Rumi Chunara, PhD

Associate Professor of Biostatistics and Computer Science and Engineering, NYU; Director, Center for Health Data Science

Chunara pioneered social media-based disease surveillance, developed influential frameworks for equitable AI in health (*Nature Machine Intelligence*, 2021), and has directly measured algorithmic bias at NYC Health + Hospitals. Her 2024 *PNAS* paper demonstrated that utilizing big data without domain knowledge impacts public health decision-making. She is an ACM Distinguished Member, NSF CAREER awardee, and MIT TR35 honoree and is widely cited (Google Scholar, accessed April 2026).

Questions she'll raise: Whose data trains the algorithm, and who is missing? Can fairness be “added on” after a model is built, or must it be designed in? When AI scales a biased system, who bears the harm?

PREPARATION READING

PAPER ML and algorithmic fairness in public and population health (*Nature Machine Intelligence*, 2021)

PAPER Big data without domain knowledge impacts public health decisions (*PNAS*, 2024)

PAPER Algorithmic bias in the safety net (*npj Digital Medicine*, 2025)

THE SYSTEMS BUILDER — CLINICAL AI DEPLOYMENT AT SCALE

Yindalon (Yin) Aphinyanaphongs, MD, PhD

Research Professor, Dept. of Population Health & Medicine; Director, Division of Applied AI Technologies (DAAIT), NYU Langone Health

Aphinyanaphongs co-led the development of NYUTron, a BERT-like clinical language model trained on 7.25 million clinical notes, published in *Nature* (2023). His mortality prediction system, deployed as interruptive EHR alerts, has been associated with a substantial increase in goals-of-care conversations and a reduction in the Vizion Mortality Index. He has deployed more than 30 operational AI models at NYU Langone, led the rollout of GenAI Studio (a HIPAA-compliant generative AI platform with broad internal adoption), and oversees the Ultraviolet HPC cluster, one of the largest academic GPU clusters in the U.S.

Questions he'll raise: What does it take to move an AI model from publication to production? When is the evidence “enough” to deploy? What organizational infrastructure sustains clinical AI at scale? Can AI improve not just efficiency but the *quality* of care conversations?

PREPARATION READING

PAPER NYUTron: Health system-scale language models (*Nature*, 2023)

PAPER Mortality risk predictions for advance care planning (*NEJM Catalyst*)

SECTION 6

The Five Fault Lines

The roundtable is organized around five interlocking tensions. For each, we provide the core question, which speakers are most directly engaged, and what to think about beforehand.

1. EVIDENCE & DEPLOYMENT

Core question: What counts as sufficient evidence to deploy AI in clinical care?

Aphinyanaphongs has deployed more than 30 models; Chunara has shown that similar models can encode bias in safety-net hospitals. The threshold for “enough” evidence looks very different depending on whether you prioritize speed of impact or equity of outcomes.

Think about: How many patients need to benefit before deployment is justified? Who decides, and how do you account for patients who might be harmed?

2. THE BODY AS DATA

Core question: What is gained and what is lost when the patient becomes a data substrate?

From Choudhury’s smartphone stress sensors to Crary’s brain tissue slides to Brandes’ protein sequences, AI increasingly “knows” the patient through computational representations rather than clinical encounters.

Think about: Does an AI model that reads your genome “know” you? Does it know you differently than a model that reads your clinical notes? What about one that listens to your voice for signs of stress?

3. EQUITY & SCALE

Core question: How do we prevent AI from scaling existing health disparities?

Chunara’s research shows bias is not hypothetical — it is measurable in real systems. Choudhury studies whether digital mental health tools work equitably across populations. Aphinyanaphongs deploys at scale across a major urban system.

Think about: If a model improves outcomes for 80% of patients but worsens outcomes for a vulnerable 20%, should it be deployed? Who makes that call?

Practical transparency tools: Model Cards, Datasheets for Datasets, and fairness-audit toolkits such as AIF360 and Fairlearn offer concrete approaches to bias documentation and mitigation.

4. TRUST ACROSS SCALES

Core question: Can a single framework for “trust in medical AI” span from molecules to health systems?

Trust means something different when a clinician trusts a protein variant prediction (Brandes), an AI pathology reading (Crary), a smartphone behavioral assessment (Choudhury), an EHR mortality alert (Aphinyanaphongs), or a fairness audit (Chunara).

Think about: Is “trust” even the right word? When you “trust” a molecular prediction, is that the same kind of trust as when you trust a doctor?

5. GOVERNANCE FROM THE INSIDE

Core question: What happens when the people building AI are also the ones governing it?

Without a dedicated regulator on the panel, governance questions fall to the builders. Brandes writes on data governance, Chunara on responsible AI, Aphinyanaphongs on institutional AI policy, Choudhury on navigating FDA/IRB processes as an entrepreneur.

Think about: Can builders effectively regulate themselves? What does “governance” look like when it comes from practitioners rather than external regulators? Is this better or worse?

SECTION 7

Key Tensions to Watch

The most generative conversations at Helix roundtables emerge from productive disagreements between speakers. Here are the pairings to watch:

Aphinyanaphongs

The deployer: “We have 31 models live. The evidence supports action.”

Chunara

vs. The auditor: “Similar models encode bias in safety-net hospitals. We must audit first.”

Choudhury

Patient-facing AI: wearables, apps, direct-to-patient sensing

Aphinyanaphongs

vs. Clinician-facing AI: EHR alerts, predictive models, workflow tools

Brandes

Foundational science: protein language models, pre-clinical tools

Aphinyanaphongs

vs. Operational deployment: the “last mile” from model to bedside

Brandes

Molecular scale: AI reads protein sequences

Crary

vs. Tissue scale: AI reads brain slides

Choudhury

Sensing as empowerment: early detection, patient agency

Chunara

vs. Sensing as surveillance: who is monitored, who benefits?

Aphinyanaphongs

Near-term: maximize clinical impact now

Brandes

vs. Long-term: existential AI safety risks demand caution

THE CENTRAL DEBATE

The **Aphinyanaphongs–Chunara** tension is the panel’s defining axis. Both are at NYU. Both care about patient outcomes. They approach deployment from opposite directions — scale versus scrutiny, action versus audit. The question of whether to deploy first and audit later, or audit first and deploy later, is the roundtable’s most consequential operational disagreement. Listen for it early and throughout.

SECTION 8

Glossary of Key Terms

AI & Machine Learning

Large Language Model (LLM)

A neural network trained on massive text data to predict and generate language. Clinical language models like NYUTron are trained on medical records for predictive tasks; NYUTron specifically is a BERT-like encoder model, distinct from generative chatbot-style LLMs.

Protein Language Model

A neural network trained on amino acid sequences (instead of words) to learn the “grammar” of proteins, enabling prediction of how mutations affect function.

Foundation Model

A large AI model trained on broad data that can be adapted to many downstream tasks. ESM (proteins) and GPT (text) are examples.

Computer Vision

AI systems that analyze images. In medicine: reading radiology scans, pathology slides, retinal photos, or dermatology images.

Multiple Instance Learning (MIL)

An ML approach for analyzing very large images (like pathology slides) by dividing them into patches and learning patterns across the collection. Used by Crary’s HistoAge.

Algorithmic Bias

Systematic errors in AI predictions that disadvantage certain groups, typically arising from biased training data or biased problem framing.

Model Drift

Degradation in model performance over time as the real-world data distribution shifts away from the training data.

AI Alignment

The research field focused on ensuring advanced AI systems act in accordance with human values and intentions. Brandes contributes to this field alongside his genomics work.

Passive Sensing

Continuous data collection from devices (smartphones, wearables) without requiring active user input. Core to Choudhury's work.

Digital Biomarker

A measurable health indicator derived from digital data (e.g., voice patterns indicating stress, typing speed indicating cognitive state).

Medical & Clinical

EHR (Electronic Health Record)

Digital system storing patient data: clinical notes, lab results, medications, imaging, billing. Primary data source for clinical AI.

Clinical Decision Support (CDS)

Any system that provides clinicians with knowledge or patient-specific information to enhance decision-making. AI-powered CDS is what Aphinyanaphongs deploys.

Alert Fatigue

Clinician desensitization from excessive electronic alerts, causing important warnings to be dismissed. A major barrier to clinical AI adoption.

Missense Variant

A single-letter change in DNA that alters one amino acid in the resulting protein. Can be benign or disease-causing. Brandes' models predict which.

GWAS (Genome-Wide Association Study)

A study scanning the entire genome of many individuals to find genetic variants associated with a disease or trait.

Tauopathy / PART

Diseases involving abnormal tau protein in the brain. PART (Primary Age-Related Tauopathy), defined by Cray, is a common age-related form distinct from Alzheimer's.

Alpha-Synuclein

A protein whose abnormal aggregation is a hallmark of Parkinson's disease and Lewy body dementia. Cray's AI detects it in peripheral biopsies.

Histopathology

The study of tissue under a microscope to diagnose disease. "Digital pathology" converts this process to AI-analyzable images.

Social Determinants of Health (SDOH)

Non-medical factors (housing, education, income, neighborhood) that strongly influence health outcomes. Central to Chunara's equity work.

Safety-Net Hospital

A hospital that provides care regardless of ability to pay, serving predominantly low-income and uninsured populations. NYC Health + Hospitals is the largest in the U.S.

Advance Care Planning

Conversations between patients, families, and clinicians about goals of care, especially near end of life. Aphinyanaphongs' AI has been associated with a substantial increase in these conversations.

Vizient Mortality Index

A standardized metric comparing a hospital's actual mortality to expected mortality, adjusting for patient severity. Used to benchmark care quality.

SECTION 9

Further Reading & Courses

Books

BOOK Eric Topol — *Deep Medicine: How AI Can Make Healthcare Human Again*

BOOK Barocas, Hardt & Narayanan — *Fairness and Machine Learning* (free online)

Free Courses

COURSE MIT OCW — Machine Learning for Healthcare (full lecture videos + materials)

COURSE Coursera — AI for Medicine Specialization (deeplearning.ai)

COURSE Stanford Online — Artificial Intelligence in Healthcare

Institutional Resources

WEBSITE Stanford HAI — Healthcare & AI Research

WEBSITE WHO — Ethics and Governance of AI for Health (guidance document)

WEBSITE FDA — AI/ML-Based Software as a Medical Device

WEBSITE NEJM AI — The New England Journal of Medicine's AI journal

WEBSITE Nature — AI in Health Care collection

WEBSITE Google Health Research

Clinical Evidence & Governance

GUIDE FDA — AI/ML SaMD Action Plan & authorized device list

GUIDE EQUATOR Network — reporting guidelines for clinical AI studies (CONSORT-AI, SPIRIT-AI, TRIPOD-AI)

PAPER AI in health and medicine (Nature Medicine, 2021 — comprehensive review)

Bias, Fairness & Real-World Implementation

PAPER Algorithm fairness in AI for medicine and healthcare (Nature Biomedical Engineering, 2023)

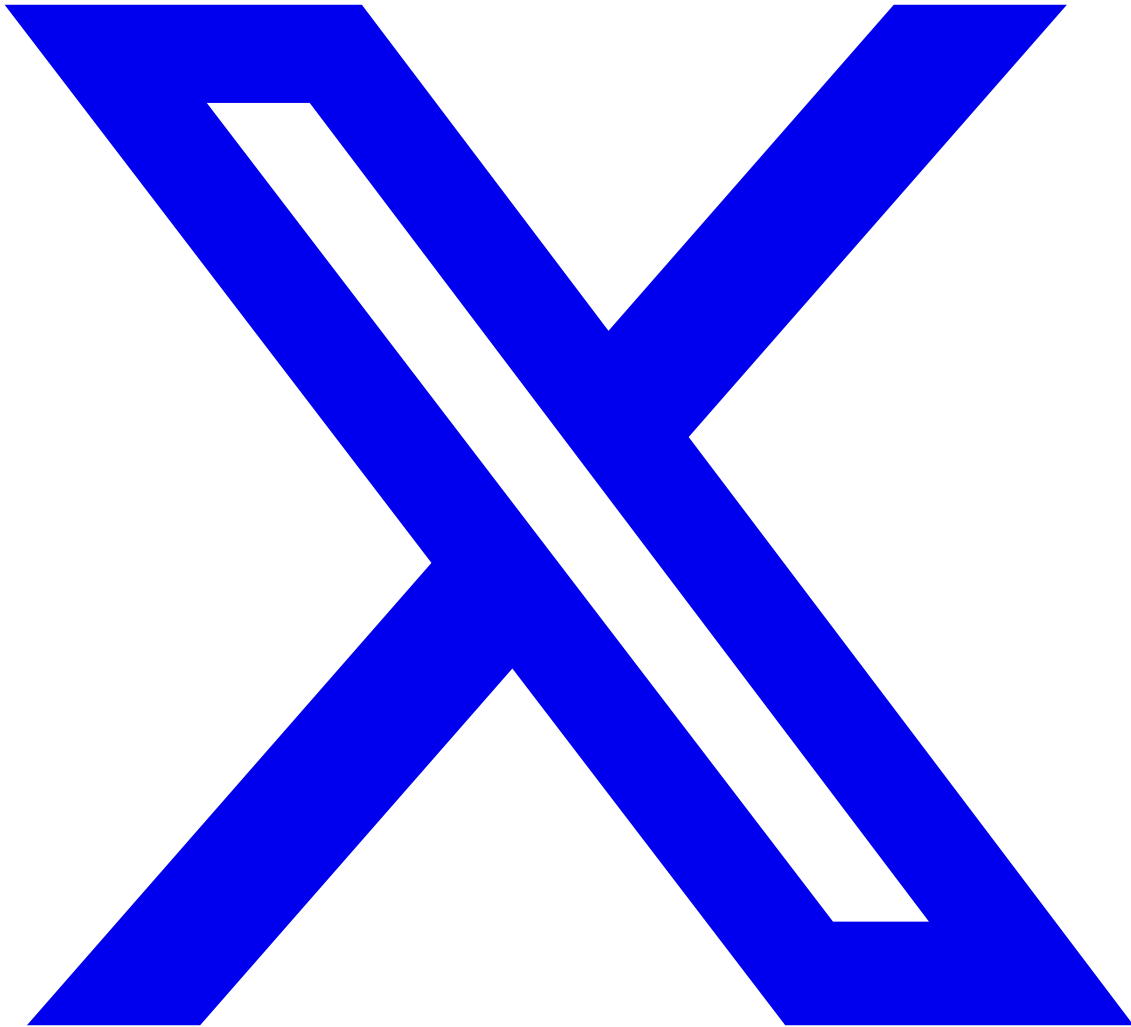
PAPER Clinical decision support with AI: overview (npj Digital Medicine, 2020)

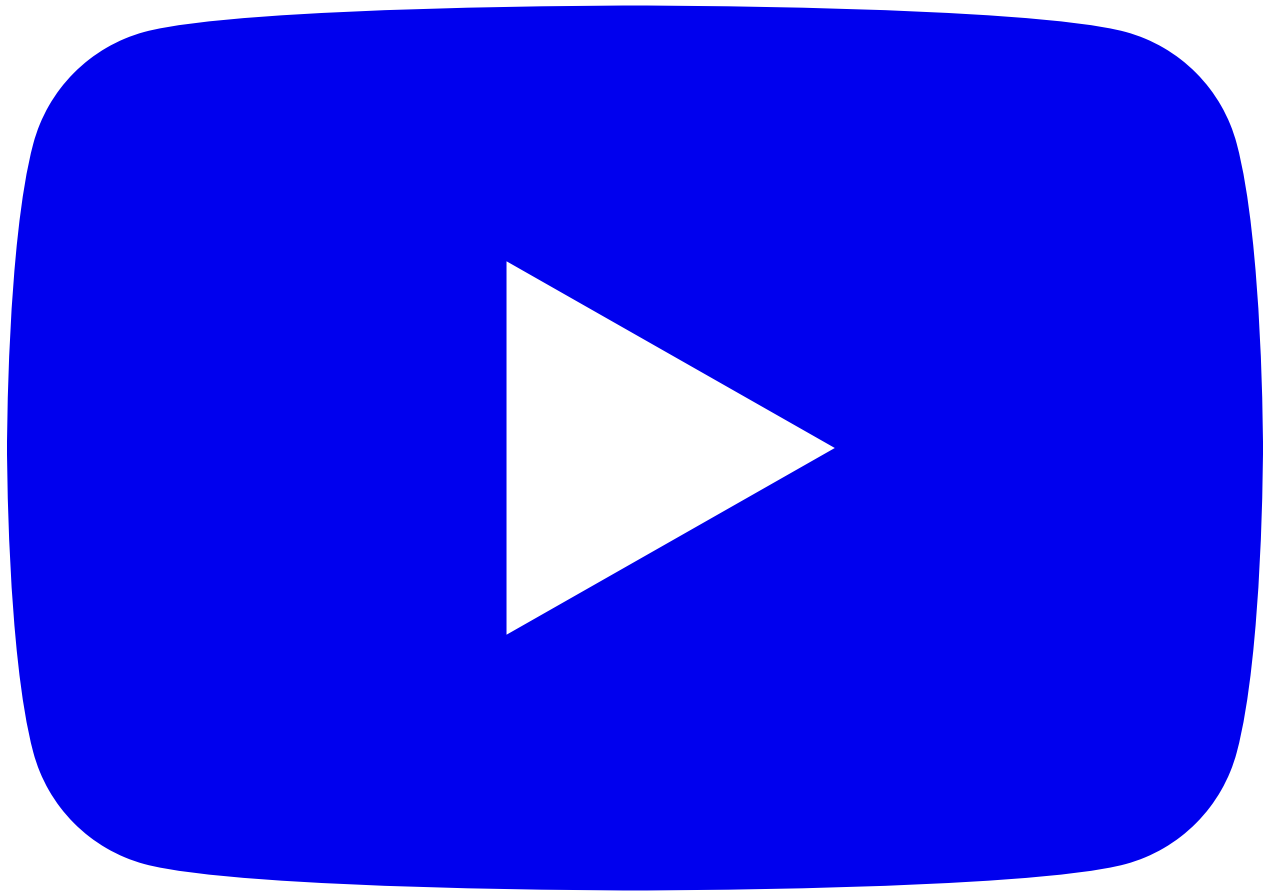
PAPER ML and algorithmic fairness in public and population health (Nature Machine Intelligence, 2021)

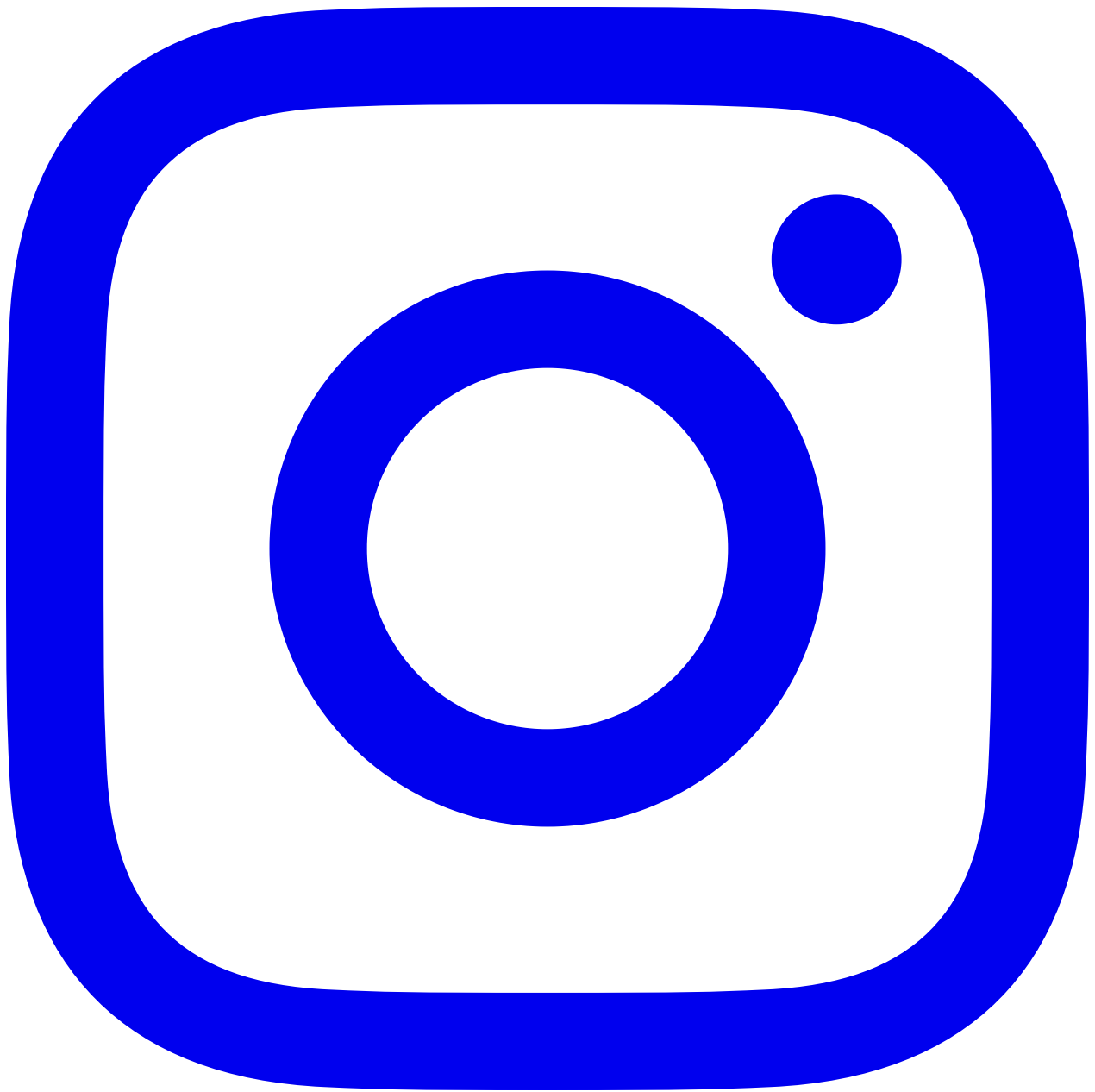
Prepared for the Helix Center Roundtable: AI + Medicine · April 18, 2026 · New York City

This guide is intended as background preparation. All linked resources are publicly accessible.



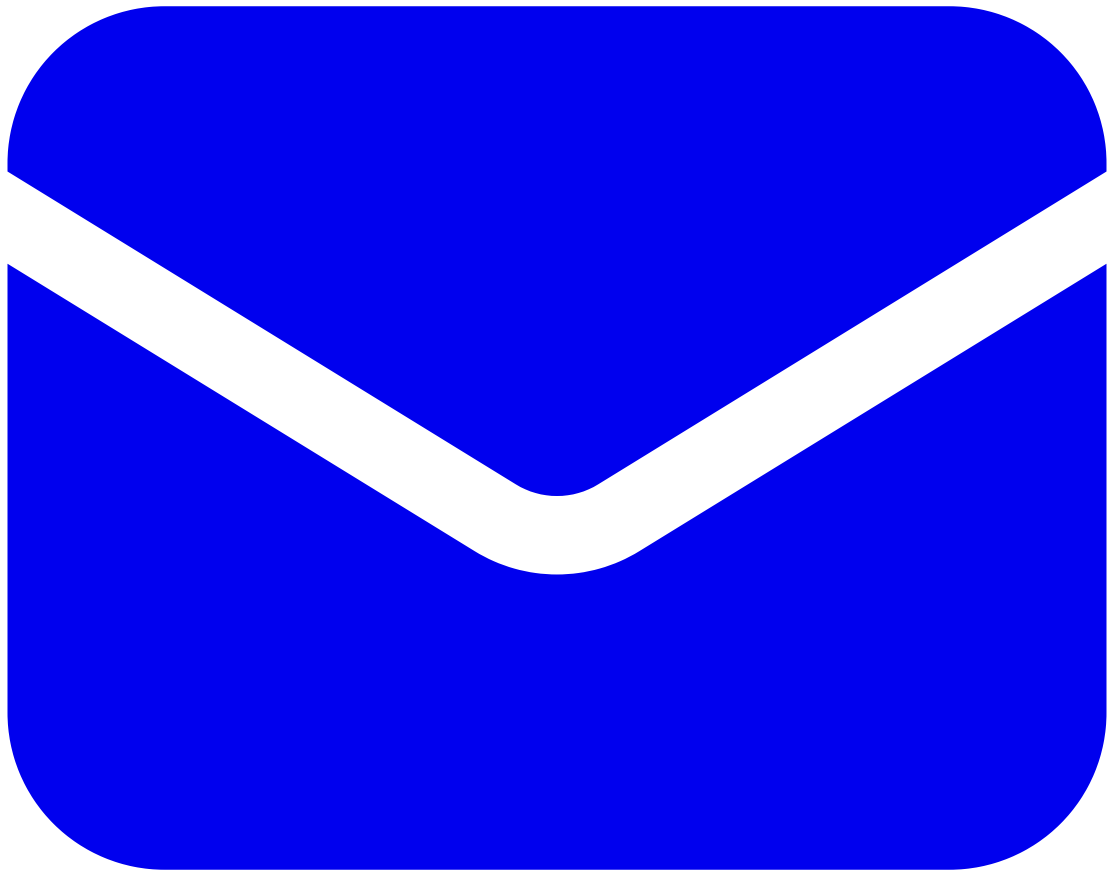












The Helix Center

The Helix Center for Interdisciplinary Investigation

The Marianne & Nicholas Young Auditorium

247 East 82nd Street

New York, NY 10028

[View Map →](#)

Explore

[About](#) [Speakers](#) [Roundtables](#) [Disciplines](#) [Leadership](#)

Engage

[Watch on YouTube](#) [Topics](#) [FAQ](#) [Contact](#)

Join Our Mailing List

First name	Last name
Email address	Sign Up

© 2026 The Helix Center for Interdisciplinary Investigation. All rights reserved. [About](#) [Contact](#)
[Privacy](#) [Terms](#) [RSS](#)