



AI-ON-Lab

Arsela Prelaj, MD, PhD,
Fondazione IRCCS Istituto Nazionale Tumori di Milano, Italy
arsela.prelaj@istitutotumori.mi.it

WHERE WE ARE WITH AI?

Different way to build our AI biomarkers



PREDICTIVE
AI TOOLS



FOUNDATION
MODELS



LLM and
MULTIMODAL
FOUNDATION
MODELS



AGENTIC AI
(All in ONE
MODELS in a
Dynamic
Fashion)

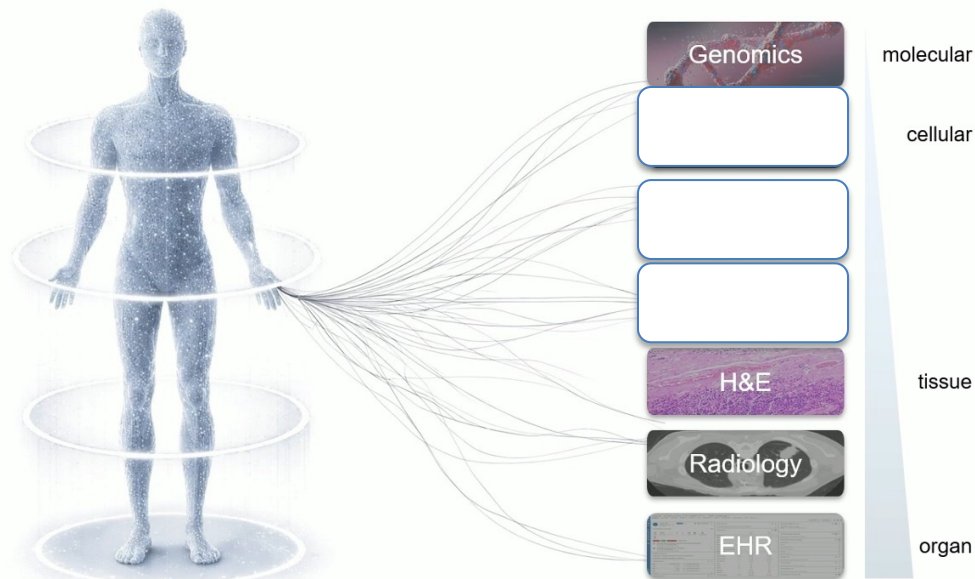
STATIC MODELS



DYNAMIC MODELS

INTEGRATED CLINICAL BIOMARKERS

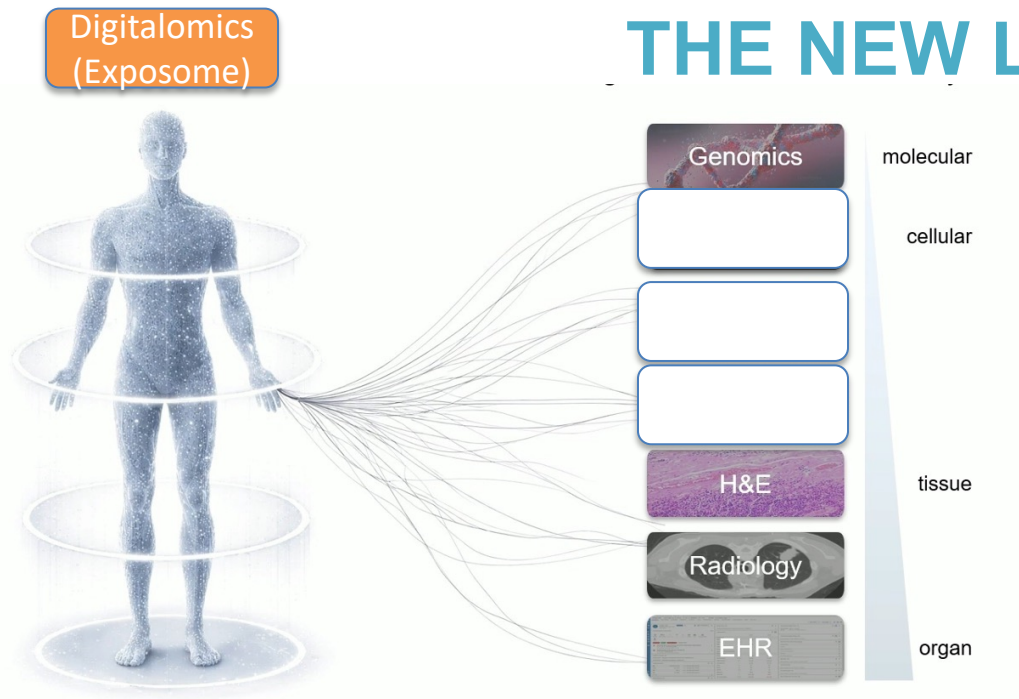
THE VISIBLE LAYER



AI Biomarkers as a result of integration, leveraging data from Clinical Practice data

With this *clinical phenotype biomarker* we are able to predict around 80% of the patients

INTEGRATED CLINICAL + EXPOSOME BIOMARKERS

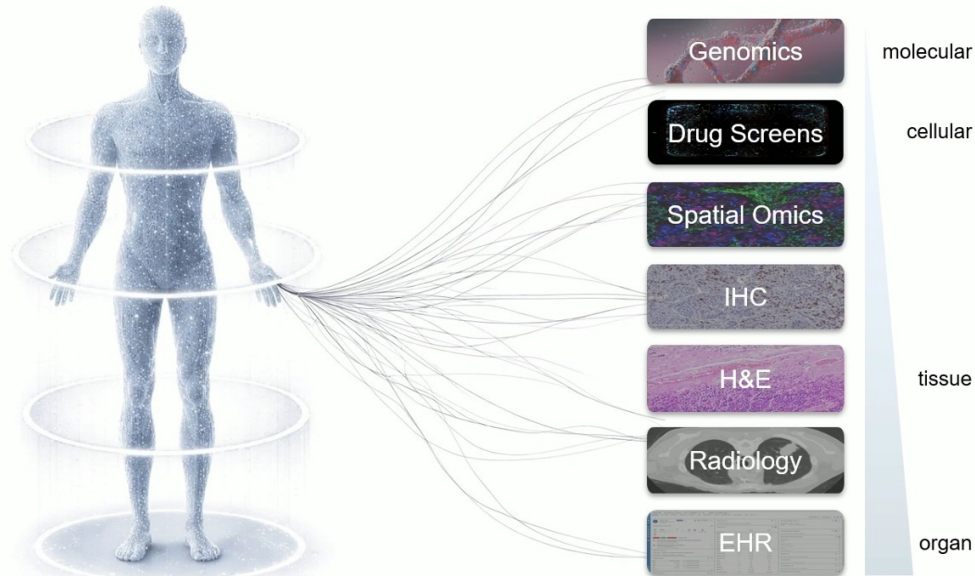


New Biomarkers collected through devices (via wearable or via app), generally ignored before using AI in healthcare

With this *expotype* we are able to further increase the prediction (no estimation is possible at the moment)

DEEP-OMIC BIOMARKERS

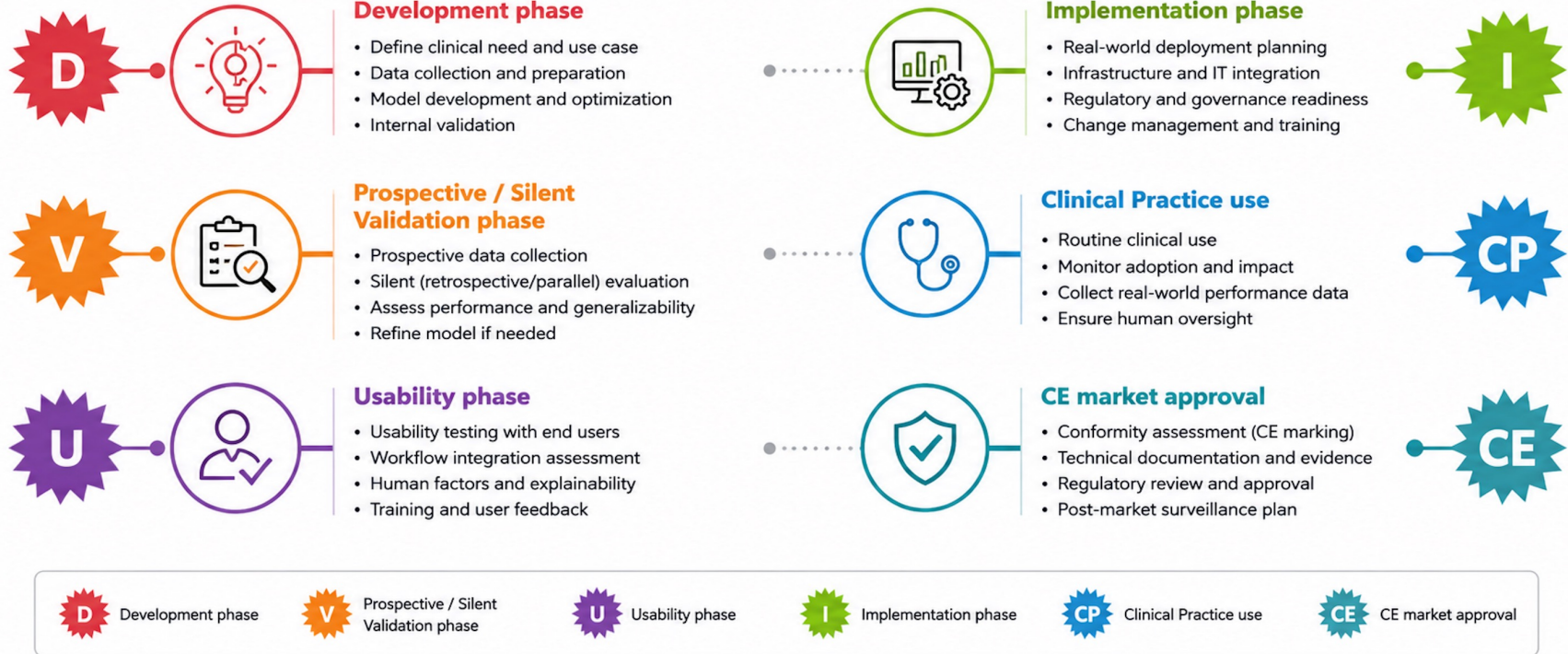
THE HIDDEN LAYER



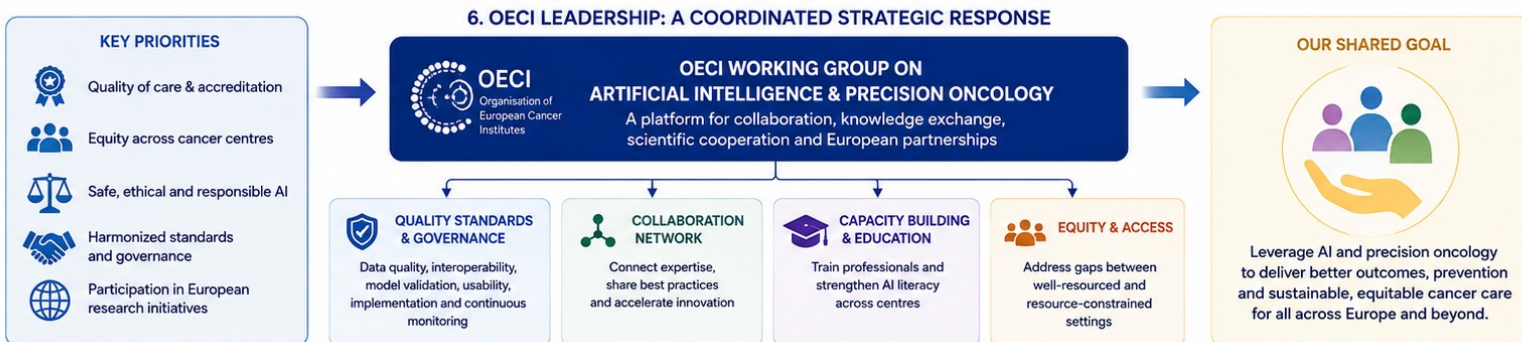
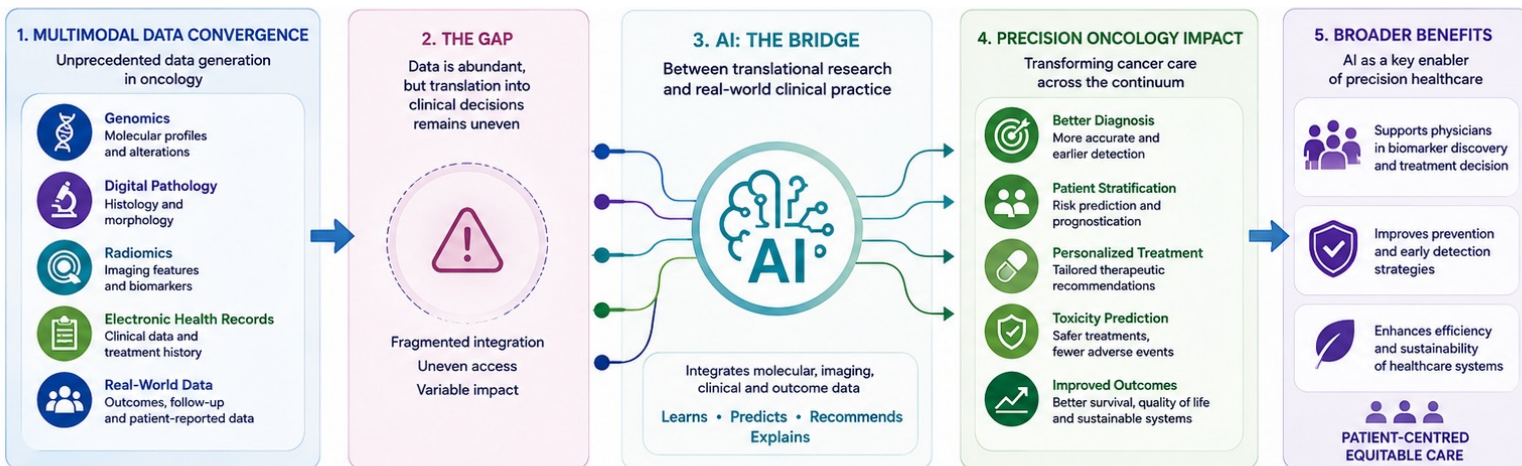
AI Biomarkers as a result of interconnection of different tissue, blood and ever-biosample omics (Not collected per clinical practise)

With this *genotype biomarker* we are able to further increase the prediction (no estimation is possible at the moment)

Life cycle of an AI biomarkers: what studies we need?

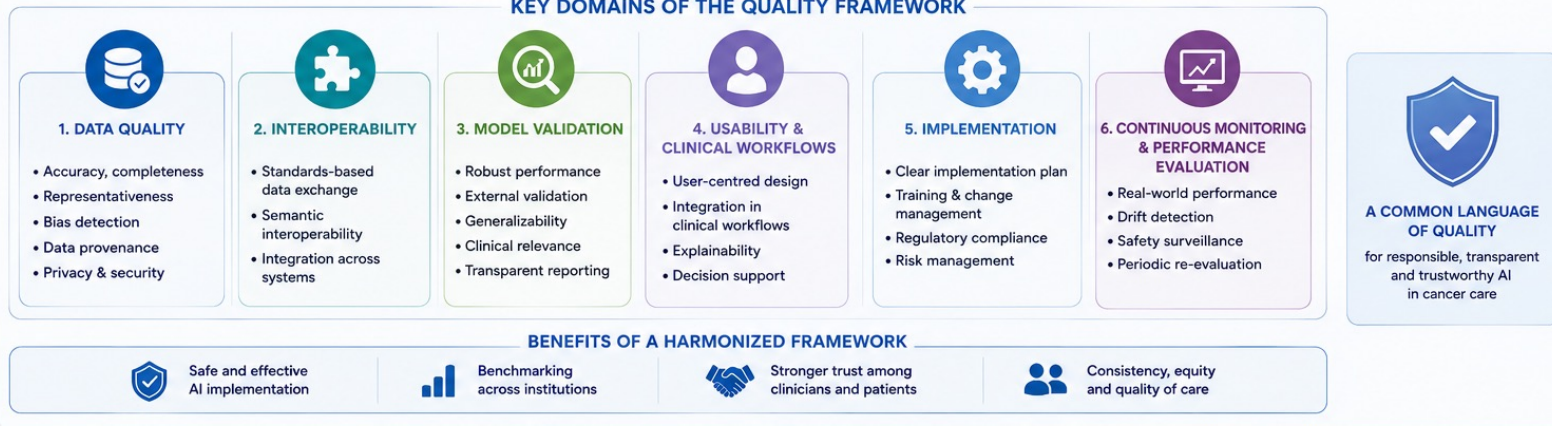


Where can OECl stand within the AI area?



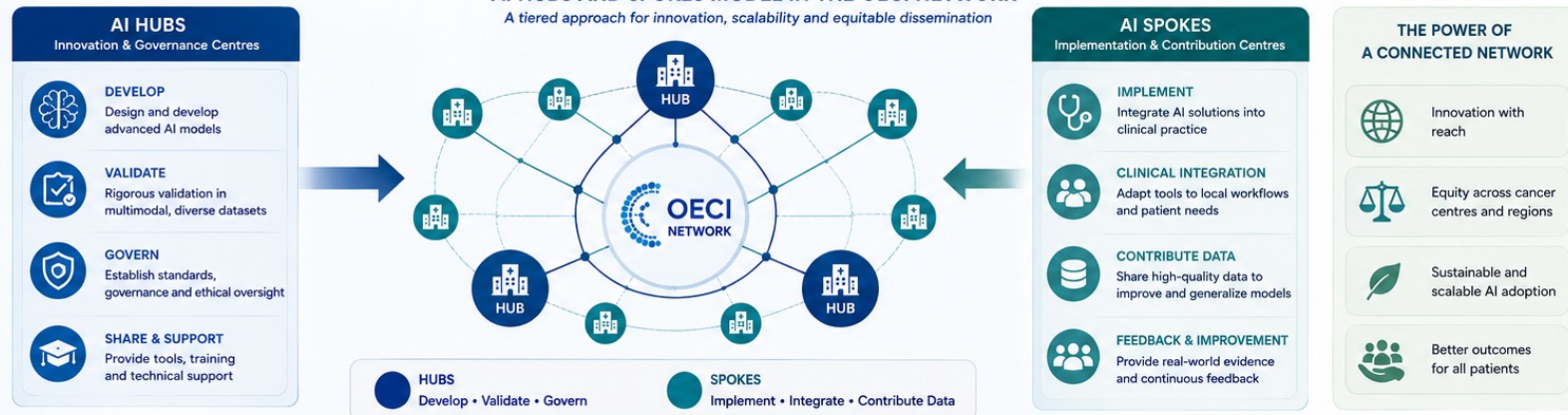
Harmonize Quality standards for AI usage in OECD affiliated centers

KEY DOMAINS OF THE QUALITY FRAMEWORK



AI HUBS AND SPOKES MODEL IN THE OECD NETWORK

A tiered approach for innovation, scalability and equitable dissemination

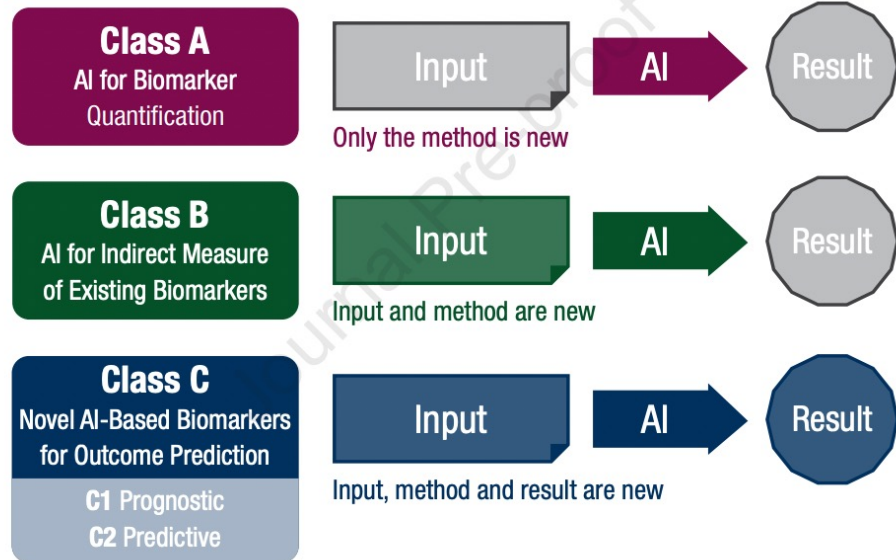


ESMO Basic Requirements for AI-based Biomarkers In Oncology (EBAI)

[M. Aldea](#)^{1,2,3} · [M. Salto-Tellez](#)^{4,5} · [A. Marra](#)^{6,7} · [R. Umeton](#)^{8,9,10,11} · [A. Stenzinger](#)¹² · [M. Koopman](#)¹³ · [A. Prelaj](#)^{14,15} · [K.L. Kehl](#)³ · [S. Gilbert](#)¹⁶ · [M-E Leßmann](#)¹⁶ · [J. Lipkova](#)^{17,18,19} · [L. Provenzano](#)^{20,21} · [F. Meric-Bernstam](#)²² · [S. Halabi](#)²³ · [J. Wu](#)²⁴ · [A. Pellat](#)²⁵ · [K.P.M. Suijkerbuijk](#)²⁶ · [B. Besse](#)^{1,2} · [B. Ryll](#)²⁷ · [C. Marchiò](#)^{28,29} · [M. Crispin-Ortuzar](#)³⁰ · [R. Fehrmann](#)³¹ · [J. Vibert](#)³² · [D. Ferber](#)³³ · [C. Pauli](#)³⁴ · [A. Valachis](#)³⁵ · [Federica Corso](#)¹⁴ · [T.J. Brinker](#)³⁶ · [J. Mateo](#)³⁷ · [N. Harbeck](#)³⁸ · [E.C. Winkler](#)³⁹ · [F. Lopez-Rios](#)⁴⁰ · [R. Perez-Lopez](#)⁴¹ · [G. Pentheroudakis](#)⁴² · [S. Delaloge](#)¹ · [C. Benedikt Westphalen](#)^{43,44} ✉ · [J.N. Kather](#)^{16,33,45} ✉ Show less

Affiliations & Notes Article Info

A



B

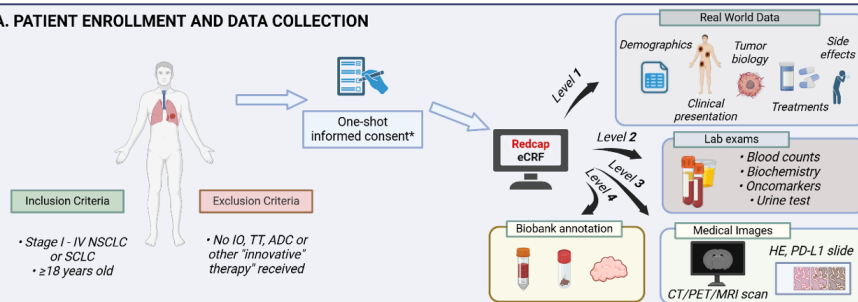
AI Biomarkers require clarity on



VALIDATION is key

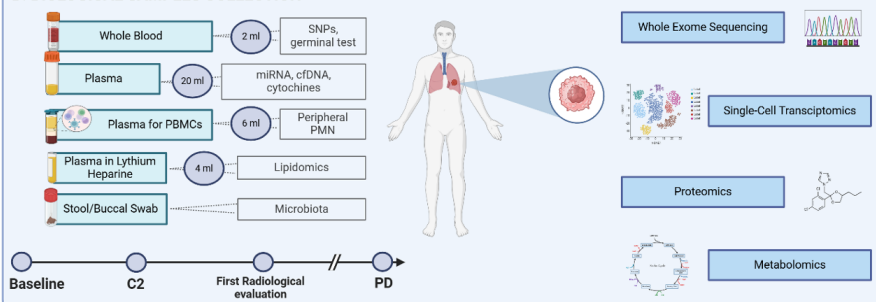
APOLLO11, BIO DATA DRIVEN MODEL

A. PATIENT ENROLLMENT AND DATA COLLECTION



*In the impossibility of requesting consent, General Authorisation to Process Personal Data for Scientific Research Purposes and General Authorisation to Process Genetic Data are applied

B. BIOLOGICAL SAMPLES COLLECTION



npj | precision oncology

Explore content ▾ About the journal ▾ Publish with us ▾

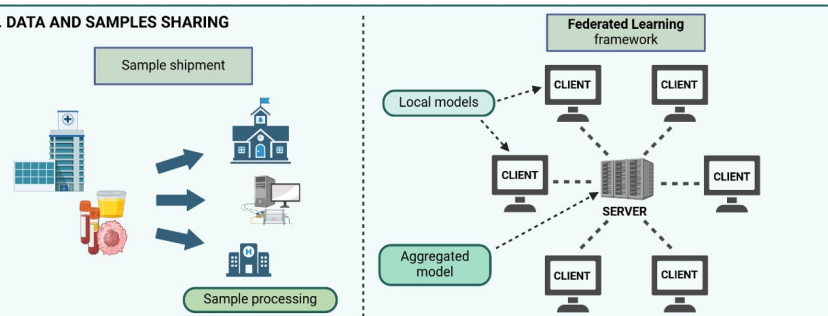
[nature](#) > [npj.precision oncology](#) > [articles](#) > [article](#)

Article | [Open access](#) | Published: 29 January 2026

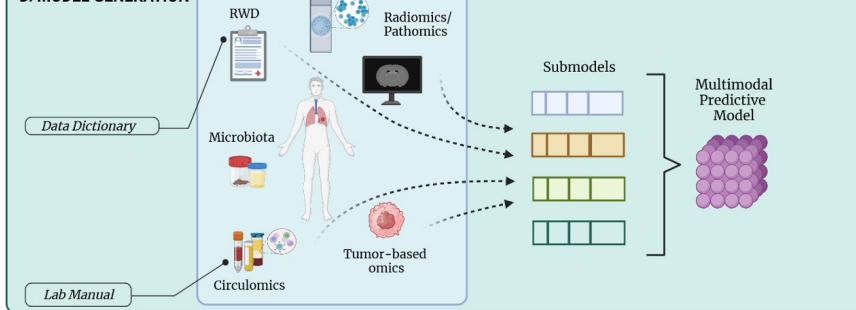
APOLLO11: a bio-data-driven model for clinical and translational research in lung cancer

[Arsela Prelaj](#) , [Leonardo Provenzano](#), [Vanja Miskovic](#), [Monica Ganzinelli](#), [Laura Mazzeo](#), [Maria Gemelli](#),

C. DATA AND SAMPLES SHARING



D. MODEL GENERATION



DATA remain key

Synthetic data for clinical trials

nature reviews cancer

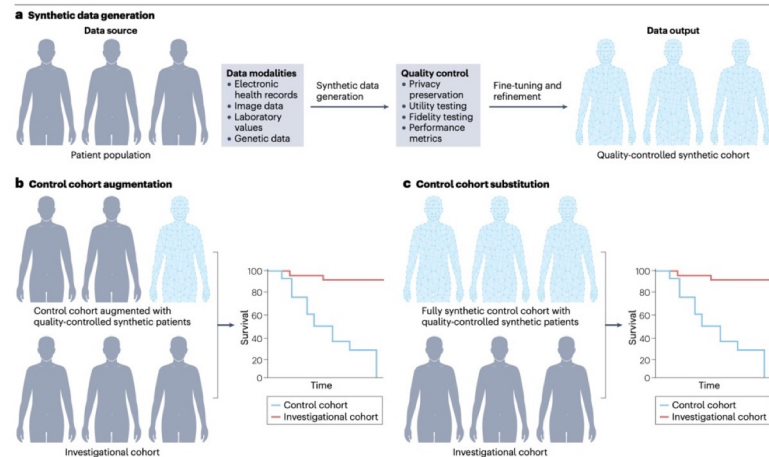
<https://doi.org/10.1038/s41568-026-00912-4>

Review article

 Check for updates

Artificial intelligence-generated synthetic data for cancer research and clinical trials

Jan-Niklas Eckardt^{1,2}, Waldemar Hahn^{3,4}, Arselä Prelaj⁵, Martin Bornhäuser^{1,6,7}, Jan Moritz Middeke^{1,8} & Jakob Nikolas Kather^{1,5,8,9}



Evaluation must benchmark fidelity and utility against real-data references, with strict filtering for clinical plausibility. Privacy auditing is mandatory. Transparent workflows enable reproducibility, trust, and safe downstream use.

Table 1 | Common evaluation metrics for synthetic tabular and image data

Category	Data type	Key metrics	Purpose
Fidelity	Tabular	Statistical tests (e.g., Kolmogorov–Smirnov, χ^2) Quantification of correlation difference (e.g., Jensen–Shannon or Wasserstein distance) Propensity score classifier or distinguishability (area under the curve)	Compares real and synthetic data sets regarding their distributional and correlation similarity
	Image	FID KID Precision and recall for generative models Inception score	Measures perceptual quality, alignment and diversity of generated images
Utility	Tabular	TSTR using task-specific metrics (e.g., AUROC, F1, R^2 , RMSE) Ability to reproduce published results or statistical associations	Evaluates whether synthetic data can substitute for real data in predictive modelling or hypothesis testing
	Image	TSTR on intended tasks (e.g., classification, segmentation, detection) Human expert evaluation in medical contexts	Tests whether synthetic images support real-world tasks
Privacy	Tabular	Membership-inference attack Privacy budget and/or differential privacy DCR	Quantifies risk of patient re-identification or leakage of sensitive attributes
	Image	Membership-inference attack Nearest-neighbour analysis Duplicate detection Model inversion attack (emerging)	Identifies memorization or leakage of real images into the synthetic set

Patients engagement and Usability

AI Scribe

EndPoints

- User trust, transparency, and perceived explainability of AI outputs.
- Impact of the AI tool on user experience, workflow integration, and decision-making.
- Fairness, privacy protection, and bias mitigation across different user groups.

Ambient AI Scribes in Clinical Practice: A Randomized Trial

Authors: Paul J. Lukac, M.D., M.B.A., M.S. ^{ID}, William Turner, B.S. ^{ID}, Sitaram Vangala, M.S. ^{ID}, Aaron T. Chin, M.D. ^{ID}, Joshua Khalili, M.D. ^{ID}, Ya-Chen Tina Shih, Ph.D. ^{ID}, Catherine Sarkisian, M.D., M.S.H.S. ^{ID}, Eric M. Cheng, M.D., M.S. ^{ID}, and John N. Mafi, M.D., M.P.H. ^{ID} [Author Info & Affiliations](#)

Published November 26, 2025 | NEJM AI 2025;2(12) | DOI: 10.1056/AIoa2501000 | [VOL. 2 NO. 12](#)

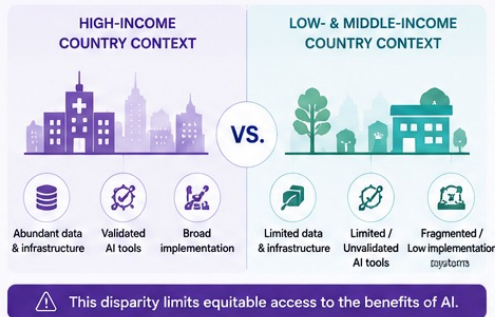
[Copyright © 2025](#)



How OECl can help in addressing “AI Gaps” or Gaps with AI

1. THE CHALLENGE: MAJOR GAPS PERSIST

AI in cancer care is unevenly applied, especially in low- and middle-income countries (LMICs). Many tools are developed in highly resourced settings, limiting their applicability where needs are greatest.



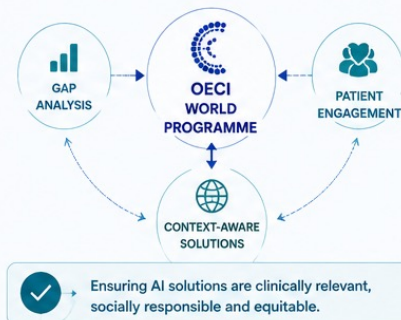
2. IMPACT ACROSS THE CANCER CARE CONTINUUM

Key areas underdeveloped or poorly implemented in LMICs



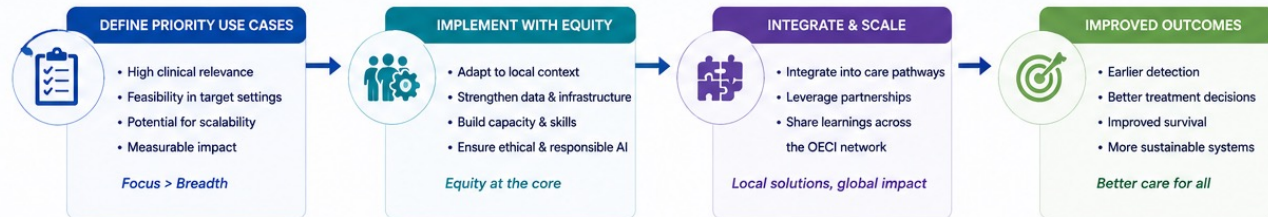
3. THE OECl RESPONSE: CONTEXT-AWARE GAP ANALYSIS

Conduct a structured, context-aware gap analysis to identify high-impact, scalable use cases tailored to diverse healthcare environments.



4. FROM VISION TO ACTION: PRIORITIZE, IMPLEMENT, DELIVER IMPACT

Focus on priority use cases that are clinically relevant, feasible and have the greatest potential impact on patient outcomes.



GUIDING PRINCIPLES

- Clinical relevance
- Equity & fairness
- Safety & reliability
- Transparency & trust
- Sustainability

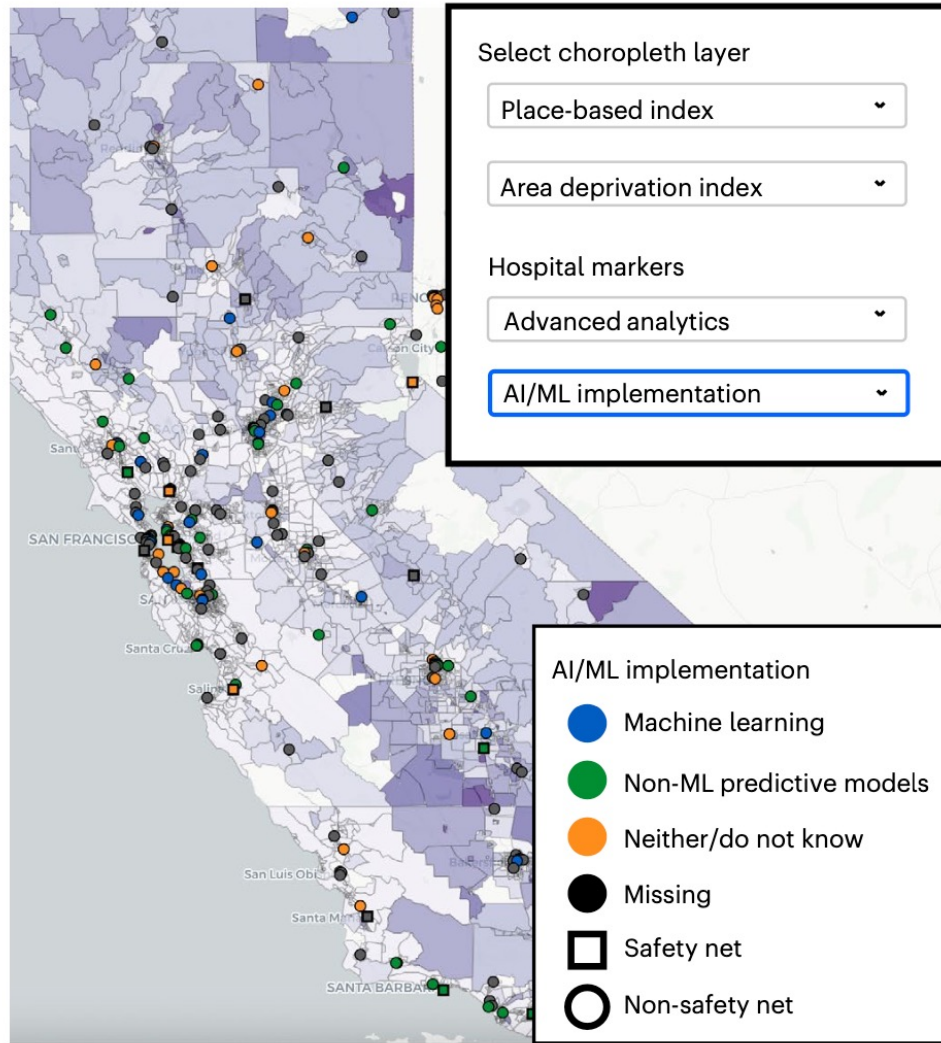
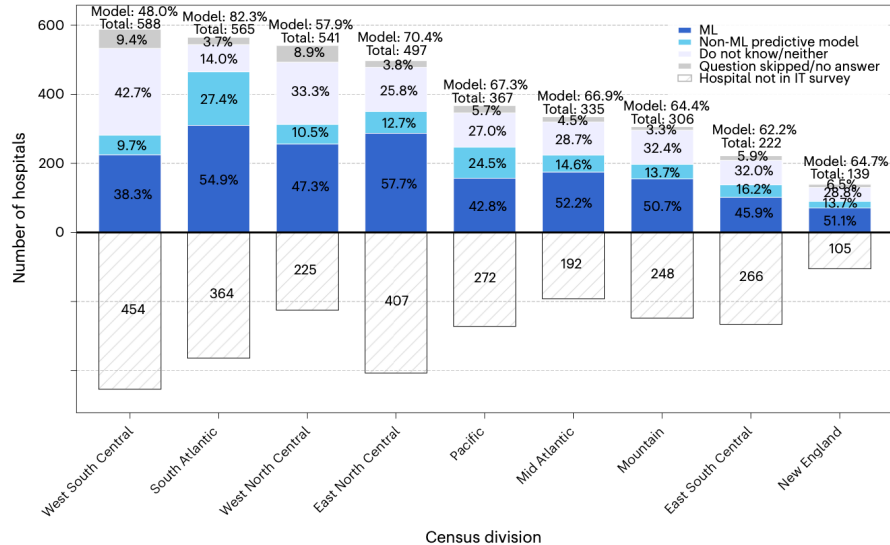
IMPLEMENTAION GAP

The landscape of AI implementation in US hospitals

Received: 23 August 2025

Yeon-Mi Hwang¹, Madelena Y. Ng^{1,2}, Malvika Pillai^{1,3}, Michelle P. Sahai¹ & Tina Hernandez-Boussard^{1,4}

Accepted: 22 October 2025





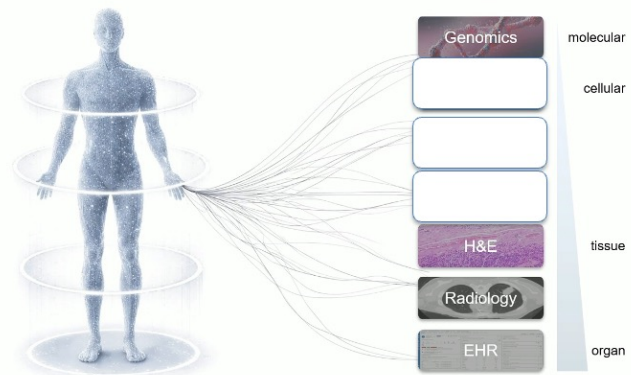
PREDICTIVE
AI TOOLS



FOUNDATION
MODELS

INTEGRATED CLINICAL BIOMARKERS

THE VISIBLE LAYER

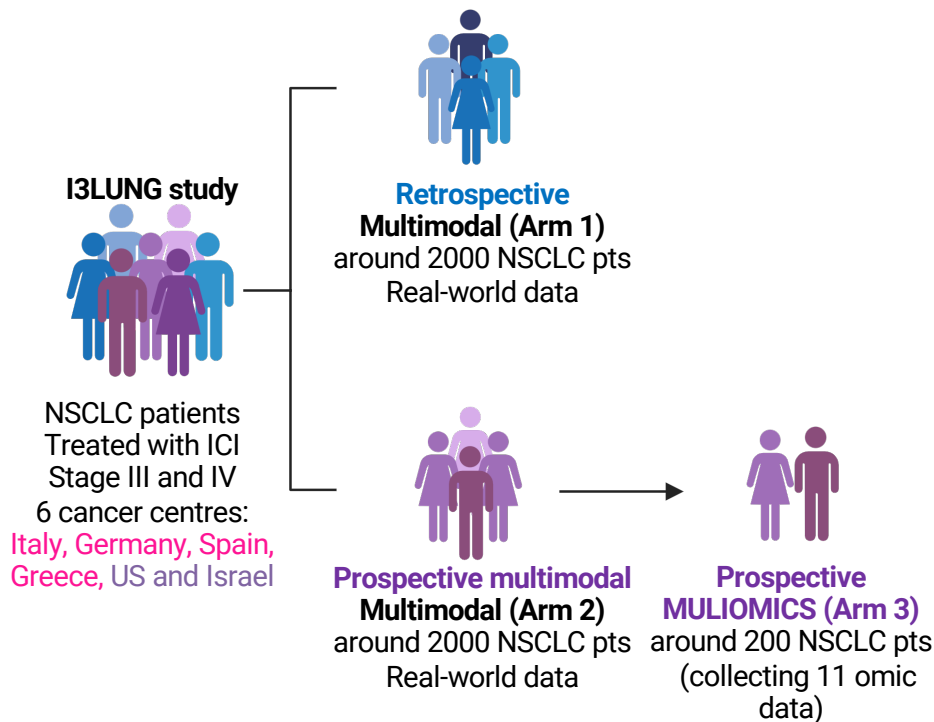


Immune prediction: the best use case for AI
(more patients compared to target)

IMMUNOTHERAPY PREDICTION OVER THE LAST 7 YEARS WITH AI

Reference	Cancer type (full modalities)	Data modalities				Analysis	Data split
		CLIN 	GEN/ TRANSC 	PD-L1 IHC/ DP 	RAD CT/ PET 		
Yoo <i>Nature 2025</i> (SCORPIO)	Pan-cancer n>9000 [NSCLC-IO=666]					Machine Learning	CV TEST EXVAL
Rakaee <i>JAMA Oncol 2024</i>	NSCLC-IO n=958					Deep Learning	TRAIN TEST EXVAL
Saad <i>Lancet Dig.Health 2023</i>	NSCLC-IO n=976					Ensemble Deep Learning	TRAIN TEST
Chang <i>Nat Cancer 2024</i> (LORIS)	Pan-cancer n=2881 [NSCLC-IO=1456]					Logistic Regression	CV TEST
Jee <i>Nature 2024</i> (MSK-CHORD)	Pan-cancer n=24950 [NSCLC-IO=7809]					Random Survival forest, NLP-Multimodal	TRAIN TEST EXVAL
Xiang <i>Nature 2024</i> (MUSK)	NSCLC-IO n=118					Vision-language FM	CV
Vanguri <i>Nat Cancer 2022</i>	NSCLC-IO n=247 (81)					DyAM Multimodal integration	CV
Captier <i>Nat Commun 2025</i>	NSCLC-IO n=317 (80)					Multimodal Late Fusion	CV

I3LUNG Project arms



Objective: Developing AI-based tools by using multimodal data for clinical decision making in NSCLC patients treated with immunotherapy



Prelaj et al. Under Revision in Nature Medicine

Retrospective Arm and Cohorts=2396 patients

Cohort collection

Study population



2396 patients, III-IV stage NSCLC ,
Immune checkpoint inhibitors
6 clinical cancer centers:
INT, GHD, SZMC, VHIO; UOC, MH

Collected baseline data



2252 clinical data demographic;
treatment; lab; molecular



1723 Genomics



936 H&E slides



1705 CT scans

Cohort creation

Cohort 1
Curative setting

Cohort 2
First metastatic line

Cohort 3
Later metastatic lines

used in analysis (n=2075)



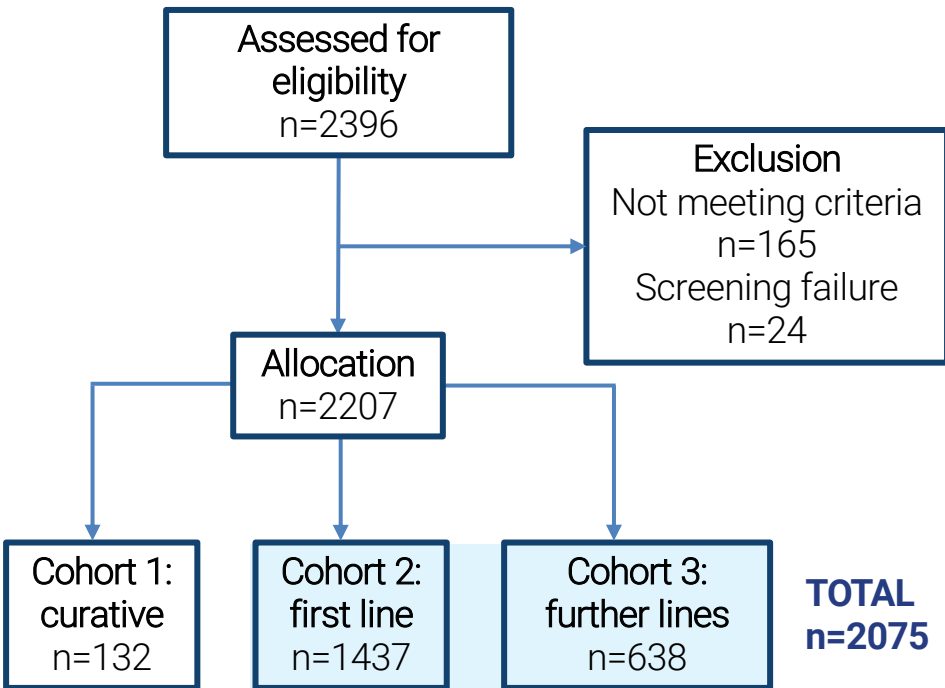
Train set
(85%, n=1550)



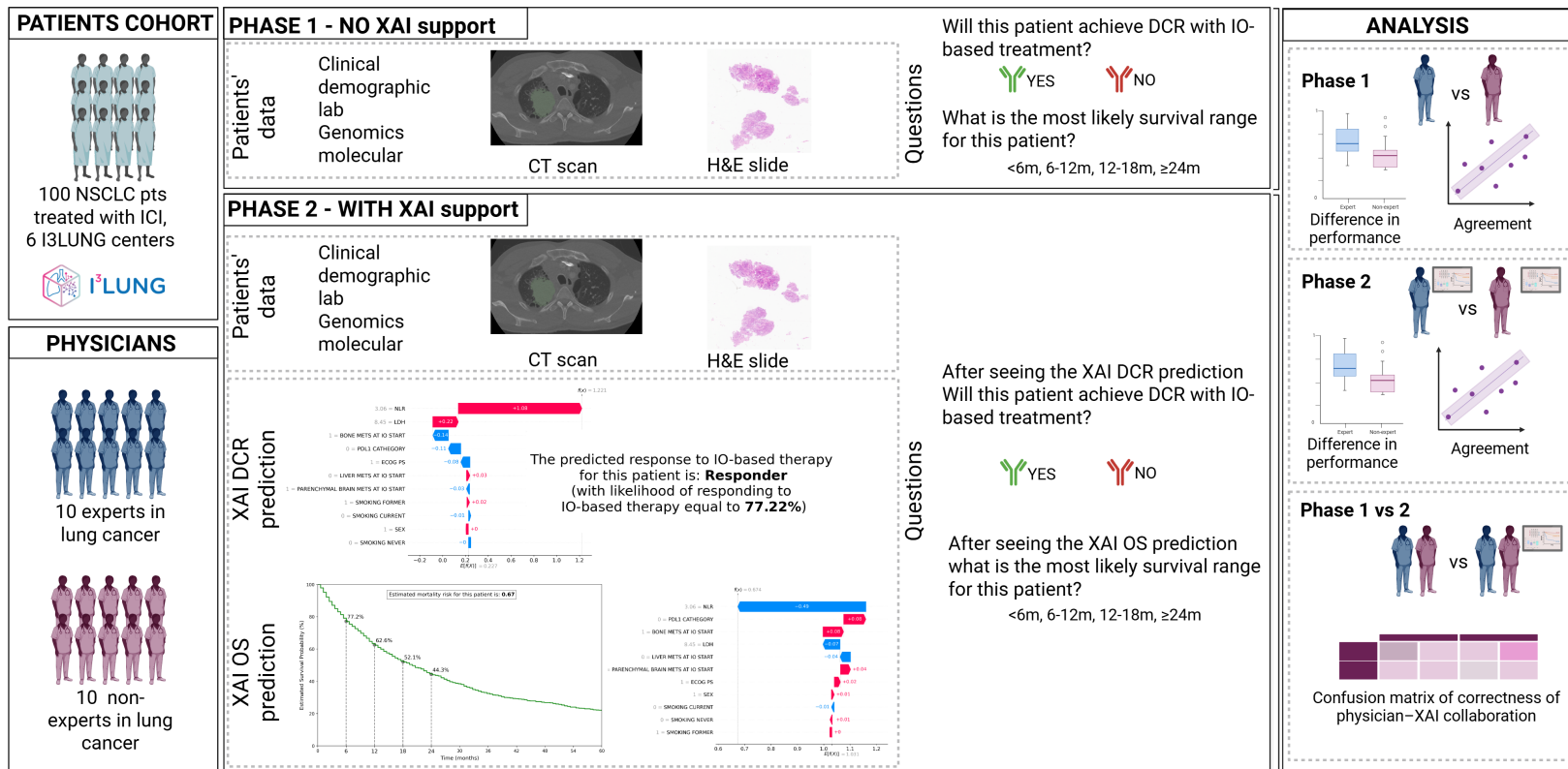
Test set
(15%, n=274)



Ex. val. set (UOC)
(n=251)



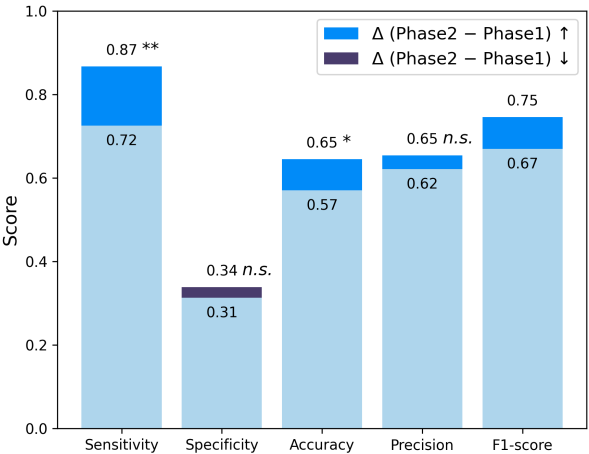
I3LUNG: Clinical Usability study (Expert vs Non-Expert)



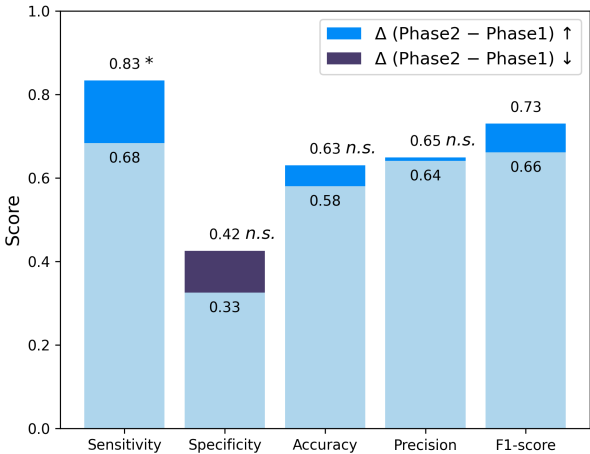
AI-ON-Lab

Clinical usability: IO efficacy prediction and user feedbacks

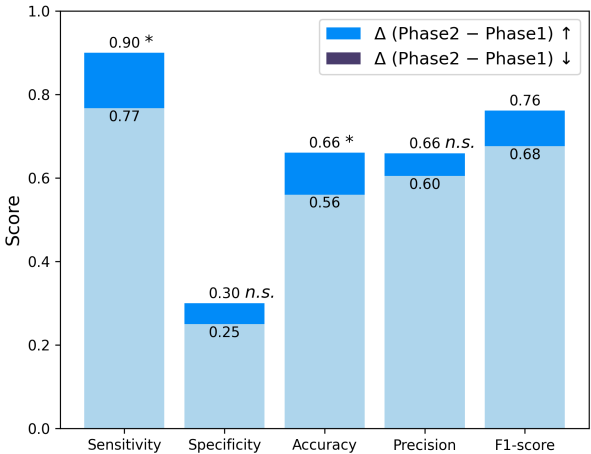
All physicians



Lung Cancer Experts



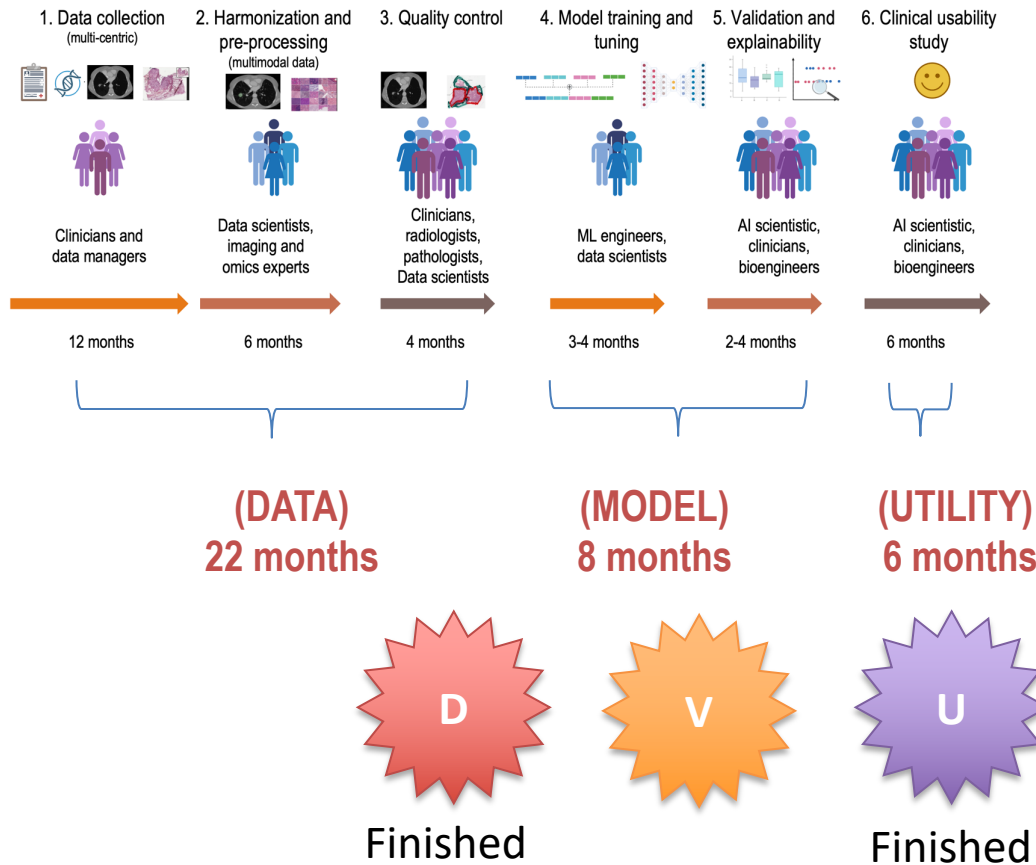
General Oncologist
Non-experts



XAI increases the probability of predicting the true DCR by 37% and true OS by 36% for all physicians (DCR: 23% for experts vs 53% non-experts) (OS: 14% for experts vs 61% non-experts)

From Data To Implementation

Our experience on human and time resources: I3LUNG use case

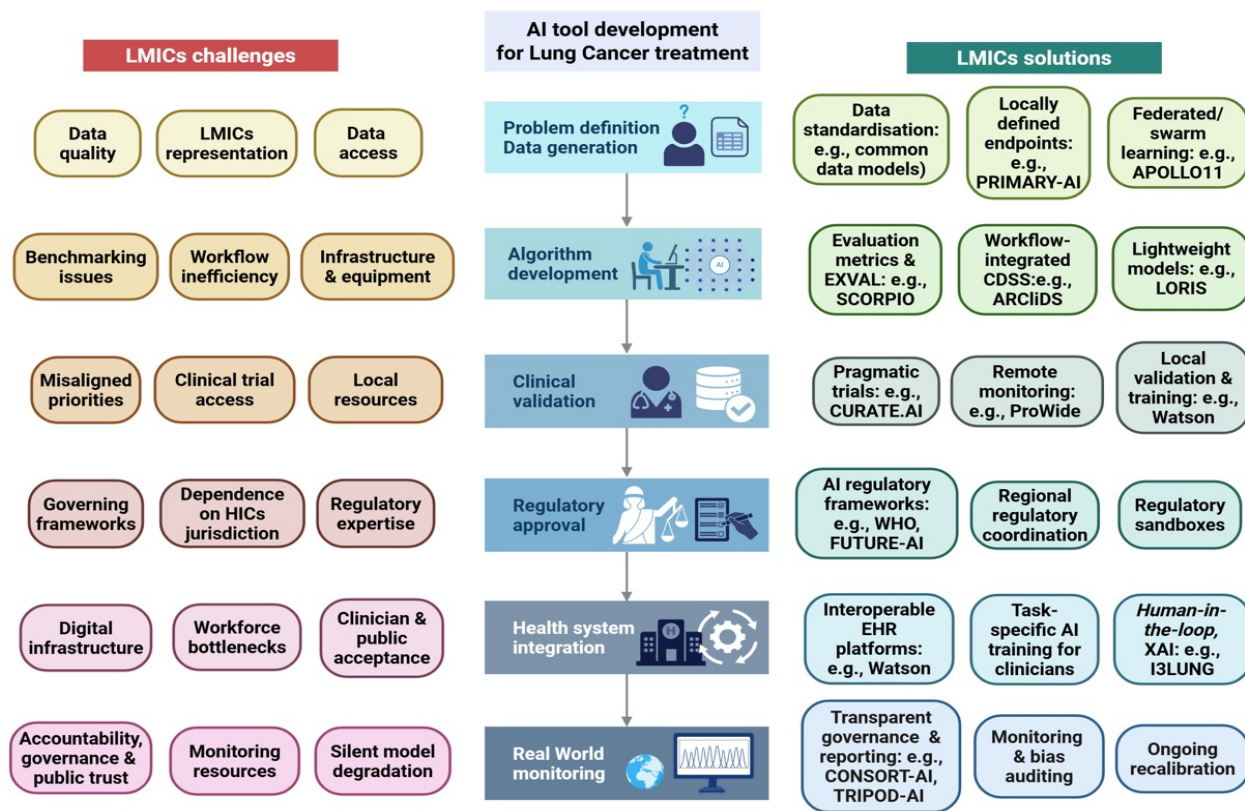


7. Implementation and data interoperability



ORIENTING LOW AND MIDDLE INCOME COUNTRY PRIORITIES

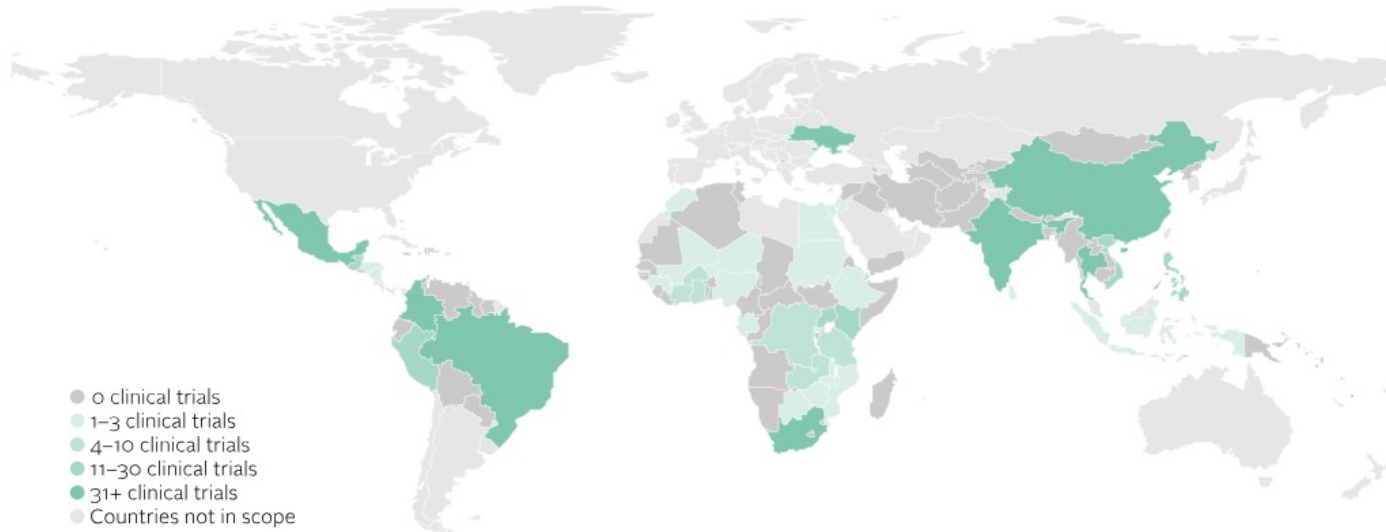
Addressing LMICs Challenges for Equitable AI:



LMICs Underrepresentation in Clinical Trials: A Global Challenge

Over half (70/113) the countries in scope have no active clinical trials

Clinical trials carried out in countries covered by the Index, are concentrated in a small cohort of countries, such as China, Brazil, Mexico, South Africa and Thailand. While these countries are all LMICs covered in the Index analysis, all five are classified as upper-middle-income countries, with few companies' conducting clinical trials in low-income countries.



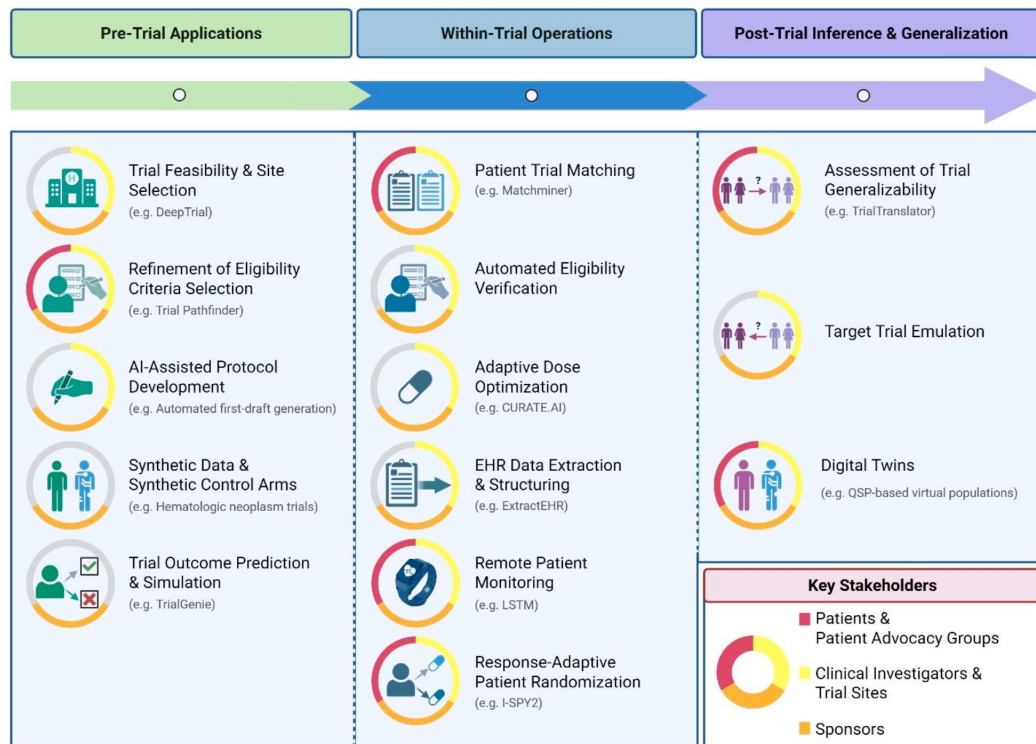
Authors: Andrea Villa^{1*}, Ashley L. Eadie^{2*}, David Synnott^{3,4}, Rebecca Romanò¹, Max Piffoux⁵, Evelyn Yi Ting Wong⁶, Naomi Schneinerman⁷, Daniel Tao Shao Weng⁶, Filippo Guglielmo Maria de Braud¹, Miriam Koopman⁸, Jarushka Naidoo^{3,9}, Madhusmita Behera¹⁰, Selen Bozkurt¹⁰, Susan Halabi¹¹, Rodrigo Dienstmann^{12,13}, Loic Verlingue^{5,14}, Arsela Prelaj^{1**}, Ravi B. Parikh^{2,10**}

AI-driven site selection can identify underserved regions with high cancer burden and low trial access.

Automated patient screening, eligibility matching, and EHR data extraction can reduce workload and resource requirements.

Support capacity building across the OECI network

AI can help emerging cancer centres participate in clinical trials and translational research.



FPOCUS ON OPEN CHALLENGES OF PERSONALIZED HEALTH CARE SYSTEM

Lung-RADS: lung cancer screening

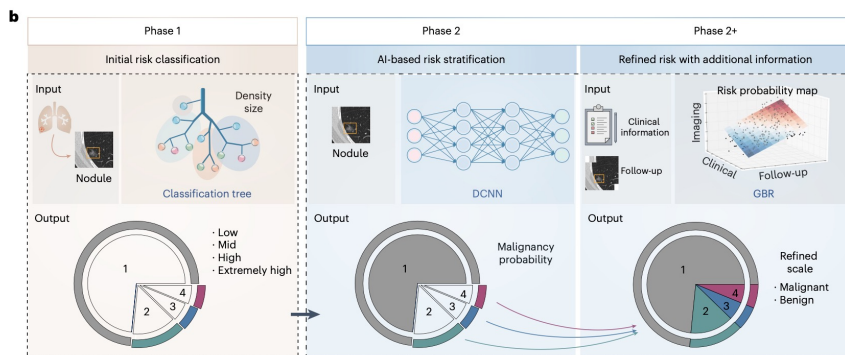
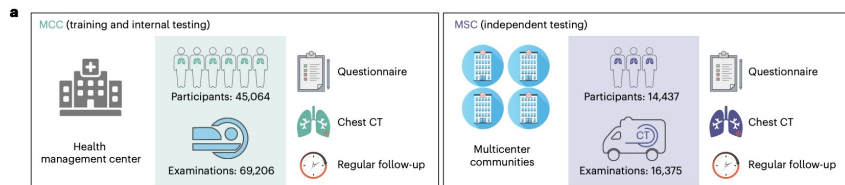


nature medicine

Developed using **45,064 participants** from a medical check-up cohort and validated on **14,437 participants** from mobile CT units across Western China

Internal dataset: AUC 0.918 vs 0.881 for single-dimension models; Independent dataset: Sensitivity 87.1%

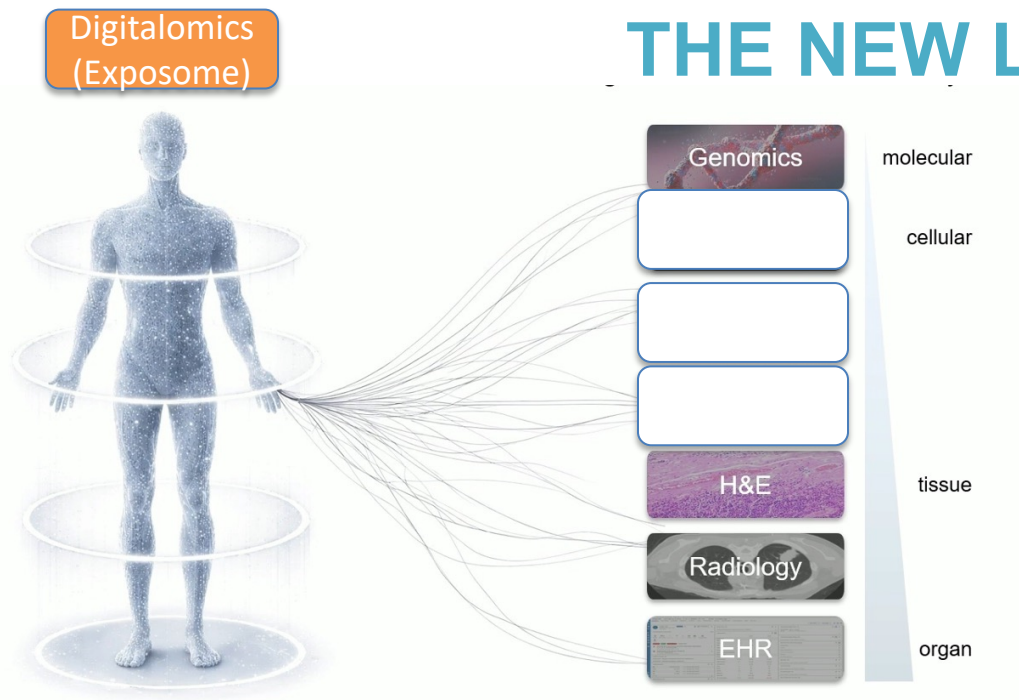
Negative predictive value: 99%, reducing risk of missed cancers



c

Category	1: low risk	2: mid risk	3: high risk	4: extremely high risk
Findings	Refined four-category risk scale			
Management	Continue annual screening with LDCT in 12 months	6-month CT	3-month CT	Immediate clinical assessment

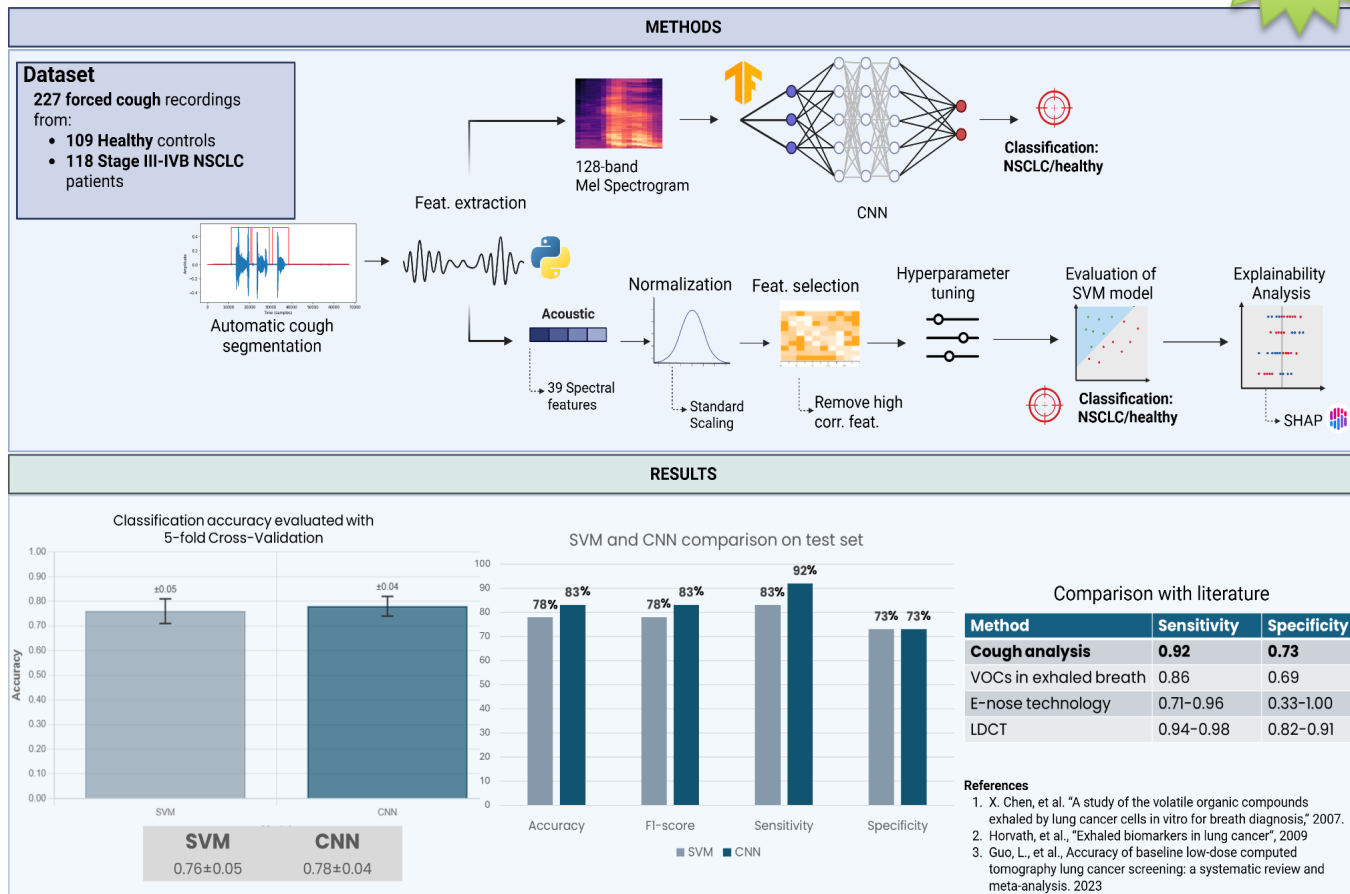
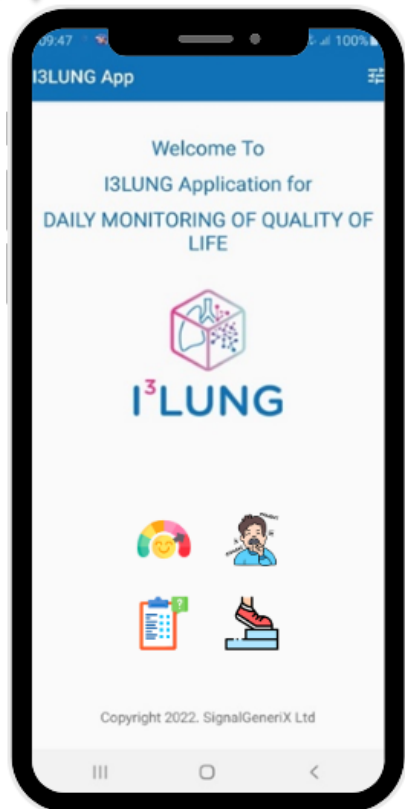
INTEGRATED CLINICAL + EXPOSOME BIOMARKERS



New Biomarkers collected through devices (via wearable or via app), generally ignored before using AI in healthcare

With this *expotype* we are able to further increase the prediction (no estimation is possible at the moment)

Cough Analysis: Prescreening Tools



Cough analysis in High Risk for Lung Cancer population

STUDY COHORTS

Healthy control

Healthy subjects
Median age: 45 (25-68), non-smokers mainly
109 cough recordings

At-risk group

Age ≥ 55
Current smokers ≥ 15 pack-years
Former smokers (last 20 yrs)
234 cough recordings

Surgical

ONGOING

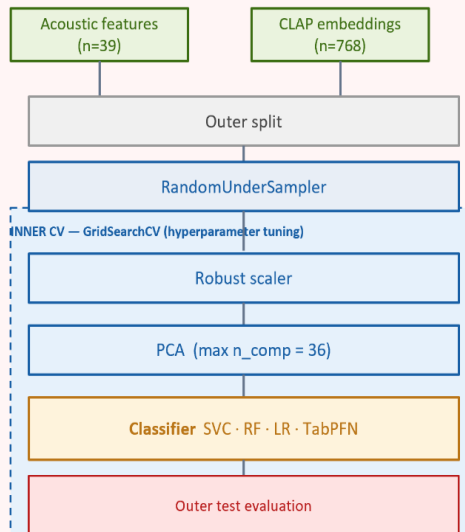
Patients with pulmonary nodules
undergoing pre-surgical staging
15 cough recordings so far

NSCLC patients

Stage III–IVB
Undergoing immunotherapy
150 cough recordings

ANALYSIS PIPELINE EXAMPLE

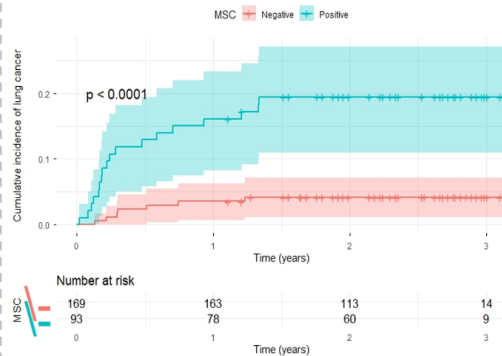
OUTER CV — stratified 5-fold



Cough Analysis

Daily coughs are recorded and a-posteriori analyzed to relate them with trend of therapy

MSC Preliminary results



Method (Preliminary results)	Sens.	Spec.
Cough analysis (at-risk vs NSCLC)	0.82	0.70
Cough analysis (healthy vs NSCLC)	0.92	0.73

Method (From Literature)	Sens.	Spec.
VOCs in exhaled breath	0.86	0.69
E-nose technology	0.71–0.96	0.33–1.00
LDCT screening	0.94–0.98	0.82–0.91



HOW WE CAN USE ETHICALLY GENERATIVE AI FOR CLINICAL DECISION MAKING GLOBALLY

Safety of a large language model-based clinical decision support system in African primary healthcare

An EMR-integrated LLM clinical decision support system was assessed across **16 primary care clinics in Kenya**

Hallucinations were uncommon (**3.4%**) and recommendations were highly concordant with local guidelines (**99%**).

Implementation challenge: Despite guideline alignment, clinicians adopted only a minority of AI suggestions, while some potentially harmful recommendations (**7.8%**) were incorporated into final documentation.

The system corrected or mitigated pre-existing clinical risks in approximately **8% of encounters**, demonstrating potential to improve care quality.



An Unsafe LLM Response Overridden

The case:

A 28-year-old woman presented, following treatment for candidiasis/UTI five days prior, on this occasion with right lower quadrant (RLQ) abdominal pain and an elevated heart rate (102). She had tenderness in the RLQ, but no distention, guarding or rebound tenderness. A complete blood count was unremarkable except for an MCV of 78.

The LLM response:

The AI Consult suggested further evaluation of possible abdominal and gynaecological pathologies, including appendicitis, ovarian torsion and ectopic pregnancy. However, the tool argued that appendicitis should be the primary differential, and suggested an ultrasound.



The outcome:

The clinician, independent of the AI Consult output, ordered a pregnancy test, which turned out to be positive. The clinician made a provisional diagnosis of ectopic pregnancy, and the patient was urgently referred to specialized care at an inpatient hospital.



A harmful LLM response enacted

The case:

A 5-year-old child, with normal vital signs, presents with a mild non-productive cough, runny nose, fevers and abdominal pain with four episodes of vomiting on the morning of the presentation. A stool sample showed mucous and moderate pus cells.

The LLM response:

The AI Consult's interpretation of the symptoms and stool results was that they suggested a gastrointestinal infection. It recommended prescribing oral rehydration solution (ORS), as well as an antibiotic such as metronidazole or amoxicillin.



The outcome:

The clinician took heed of the advice to prescribe amoxicillin, even though the presentation did not warrant it. Moreover, amoxicillin is not the preferred antibiotic for a GI infection and is likely to cause additional GI symptoms.



Overfitting to the prompt

The case:

A 21-month-old child presented with typical cold symptoms. There were no concerning findings in the history or examination. The clinician planned to prescribe OTC symptomatic treatments.

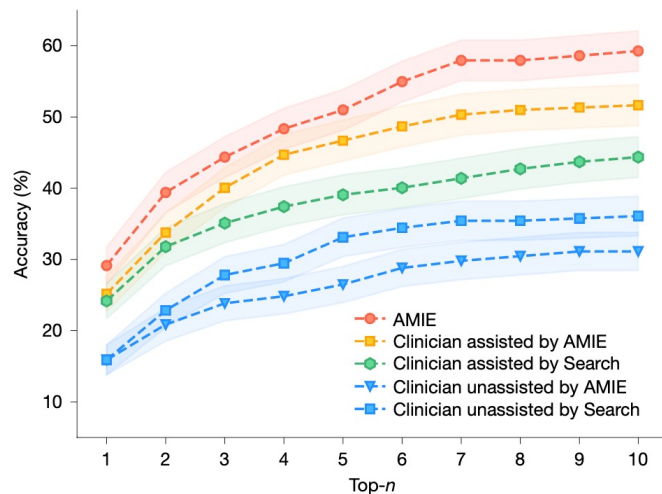
The LLM response:

In addition to confirming the clinician's initial plan, the LLM recommended, "over-the-counter throat lozenges if the child is old enough."

Within the AI Consult prompt is a one-shot example that involves viral pharyngitis. The prompt encourages the model to also consider symptomatic treatments that may not be part of standard guidelines, such as [explicitly mentioned] throat lozenges for cases of viral pharyngitis. In our post-deployment testing, we found that unprompted models do not make the mistake of recommending throat lozenges to a 1-year-old, who is clearly not old enough for that treatment to make sense. This is a good example of how detailed prompt instructions may have unintended consequences in model outputs.



LLMs for CDM (AMIE)



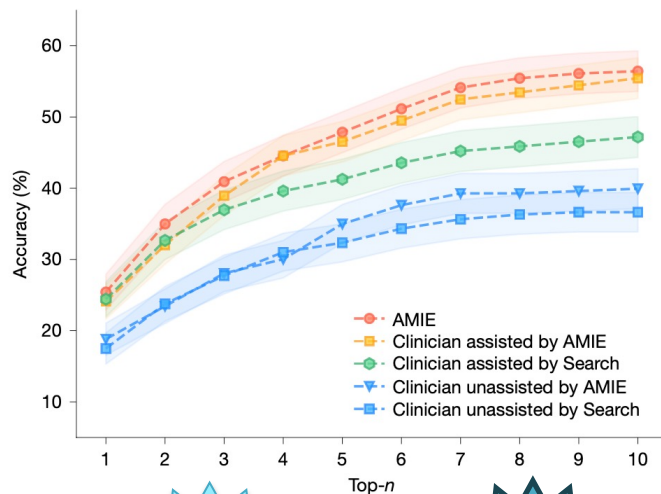
nature

Explore content ▾ About the journal ▾ Publish with us ▾

[nature](#) > [articles](#) > [article](#)

Article | [Open access](#) | Published: 09 April 2025

Towards accurate differential diagnosis with large language models



?

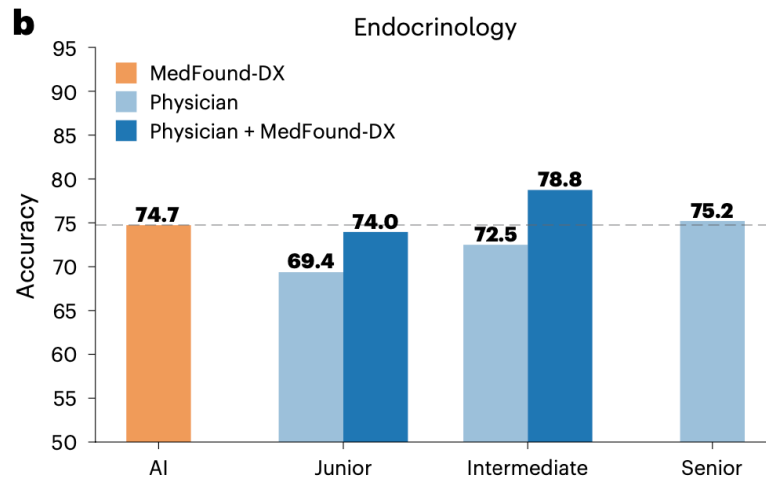
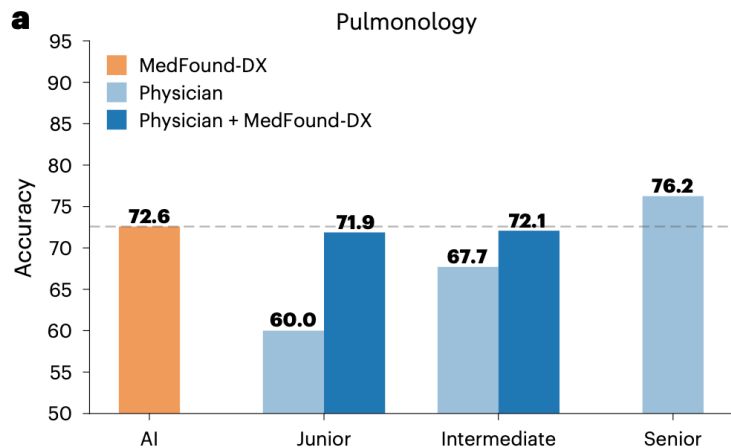
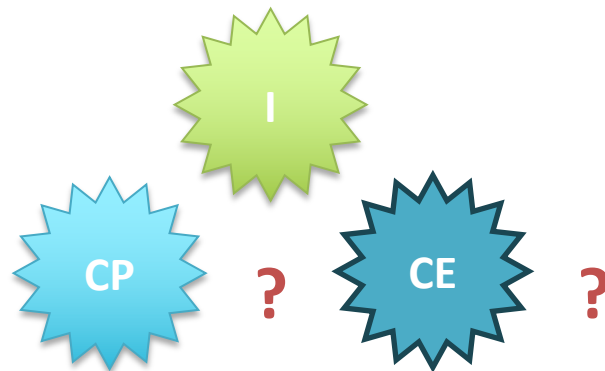


?

A generalist medical language model for disease diagnosis assistance

These results demonstrate that our LLM-based diagnostic generalist outperformed junior and intermediate physicians in both specialties and was similar to senior physicians.

MedFound-DX-PA: a BIOLLMs



Limitation: they did not compare MedFound+Senior

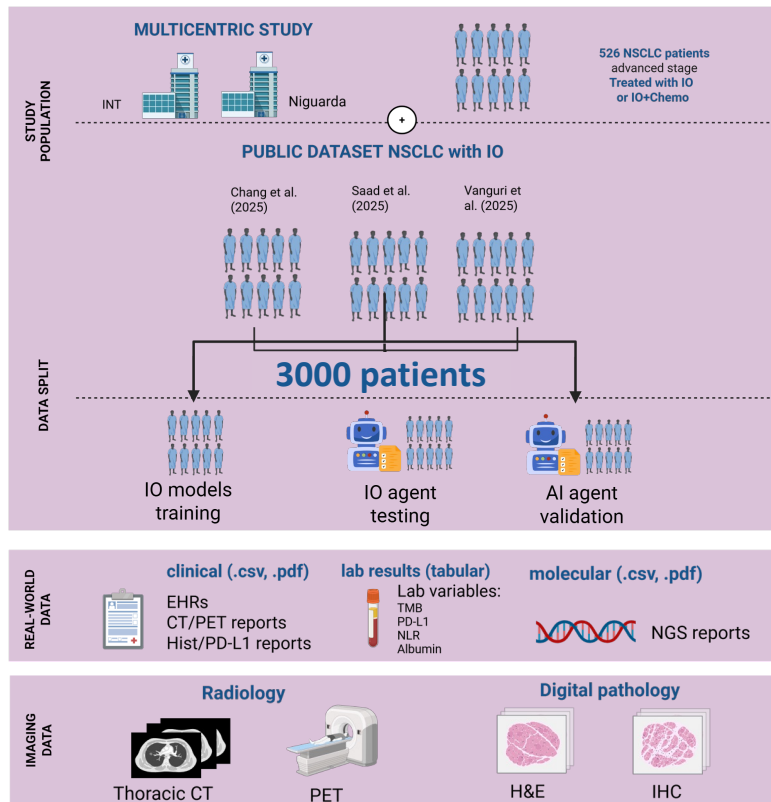


AGENTIC AI
(All in ONE
MODELS in a
Dynamic
Fashion)

AI Agents

Treatment-based AI Agents

Collaborative Human-Agent for Therapeutic Decision-Making In Cancer Immunotherapy Prediction



You are the Lead Oncologist.

Construct a final, detailed treatment plan based on the reports from your team. Synthesize all the above information into a final response. Output a structured, professional medical report. Each statement you make should be supported by the input data. Cite the input source of each statment (i.e. "Patient has IV stage NSCLC") Follow EXACTLY this format and these headings:

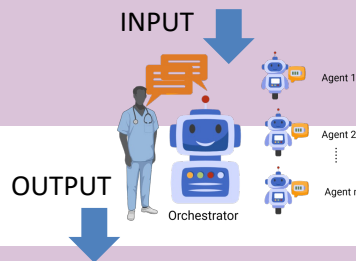
****Key Findings****

****Guideline-based Evidence****

****Predictive Tools****

****Regimen:****

****Contraddictions****



****Key Findings****

- The patient is a 63-year-old female with stage IV, adenocarcinoma of the lung (T1bN0M1b) and multiple metastatic sites (liver, bone, adrenal, pleural) with no brain metastases (Key Findings).
- PD-L1 tumor proportion score (TPS) is 10 % (PD-L1 = 10 % → category 1-49 %) (Key Findings).
- ECOG performance status is 1 (Key Findings).
- **Guideline-based Evidence**** : NCCN 2024 and ESMO 2024 guidelines recommend
- **Predictive Tools****
- ASTEP predicts a higher benefit for ICI-chemo over chemo-alone (CATE = -0.072, confidence = moderate) (Predictive Tools).
- LORIS data are unavailable; however, the ASTEP result aligns with guideline recommendations and supports the chosen regimen.

****Regimen:** Pembrolizumab + Carboplatin + paclitaxel (triple-combination).**





ESAC LinkedIn



ESAC Website

Become a Member!
Shape the Future!



EDUCATION AND PATIENT
ENGAGEMENT IN AI

**THE EUROPEAN
INTERDISCIPLINARY SOCIETY
OF ARTIFICIAL INTELLIGENCE
FOR CANCER RESEARCH**

Acknowledgments

<https://Ai-onlab.com>



Arsela.prelaj@istitutotumori.mi.it

 [@arselaprelaj.bsky.social](https://www.instagram.com/arselaprelaj.bsky.social)



i³LUNG



A project funded by
the European Union

<https://i3lung.eu>

