

Explainable, Transparent, and Trustworthy AI for Biomedicine

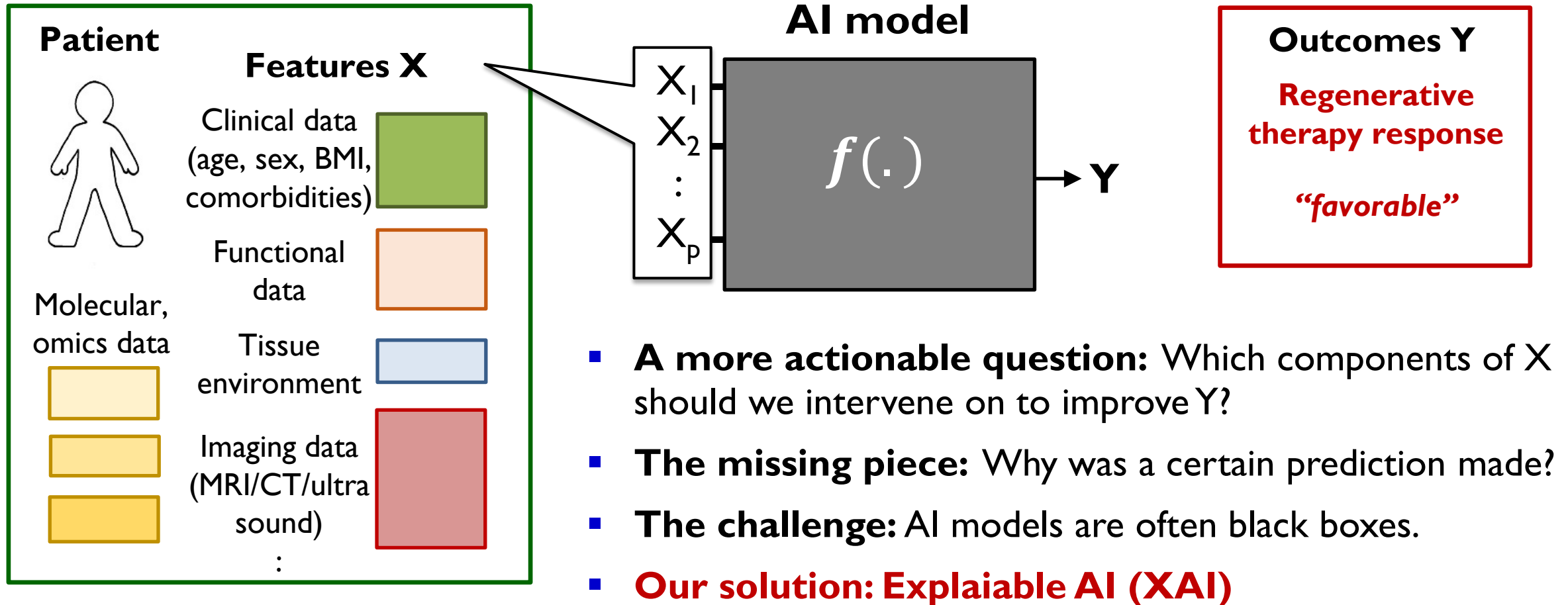
Su-In Lee

Boeing Endowed Professor

Paul G.Allen School of Computer Science & Engineering
University of Washington, Seattle

A common AI problem in biomedicine:

Given features X , predict Y .

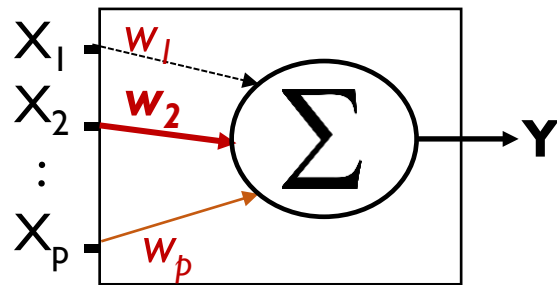


Our solution was to fundamentally advance AI research to make a prediction *with explanations*

- Accuracy vs. interpretability
 - Simple models often lead to lower performance.
 - Complex models are often considered to be a black box.

Linear model

X: Features **Y:** Outcome



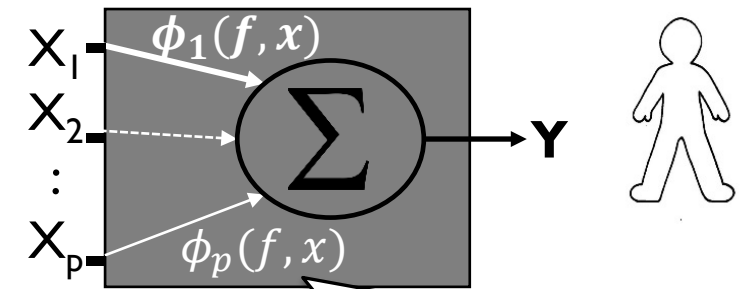
Complex model $f(\cdot)$

Black Box

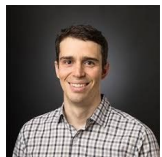


Our approach, SHAP (SHapley Additive exPlanations)

For a particular prediction

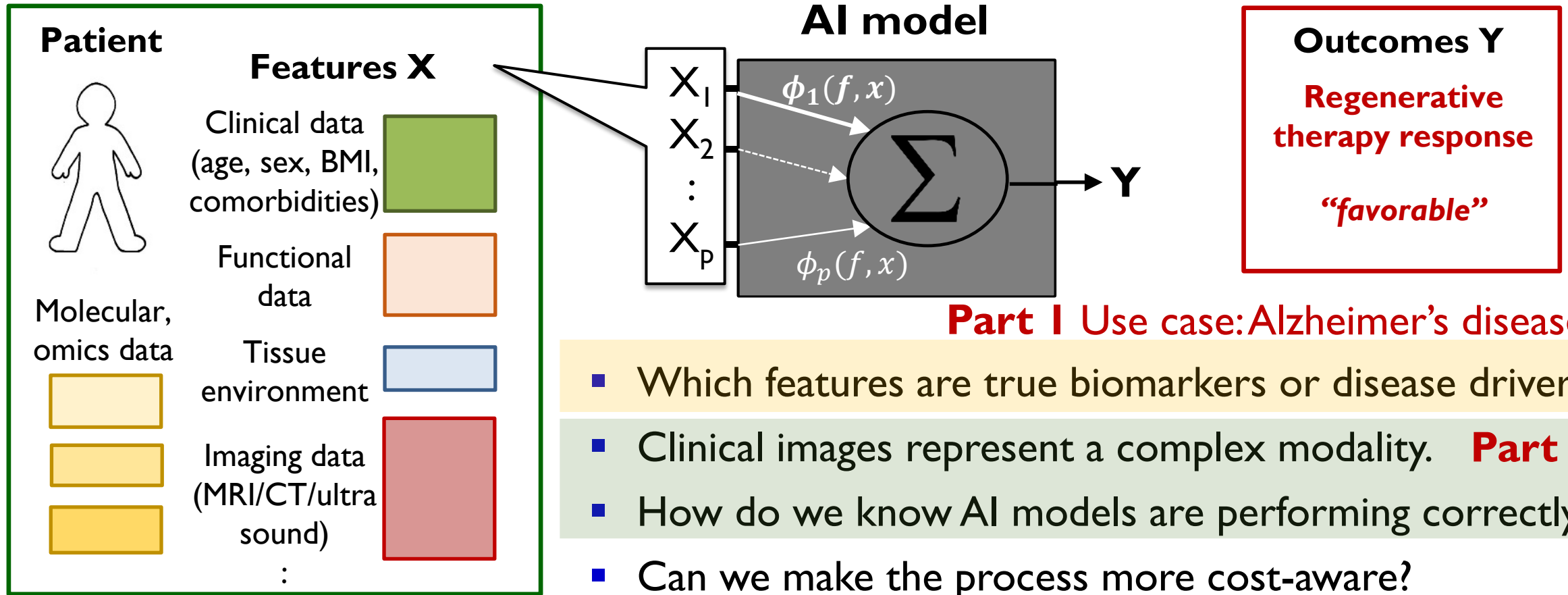


SHAP can estimate feature importance for a particular prediction for any model.



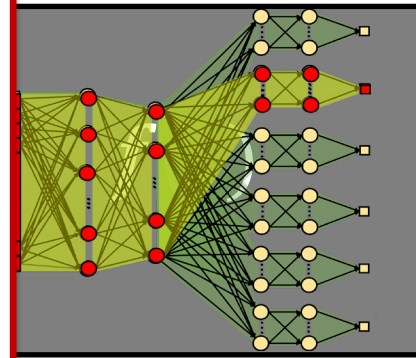
Scott, CSE PhD'19

Explainable AI (XAI): Accurately predicting an outcome is vital, but the critical question revolves around *why*.



mechanistic explanation (AD) phenotypes

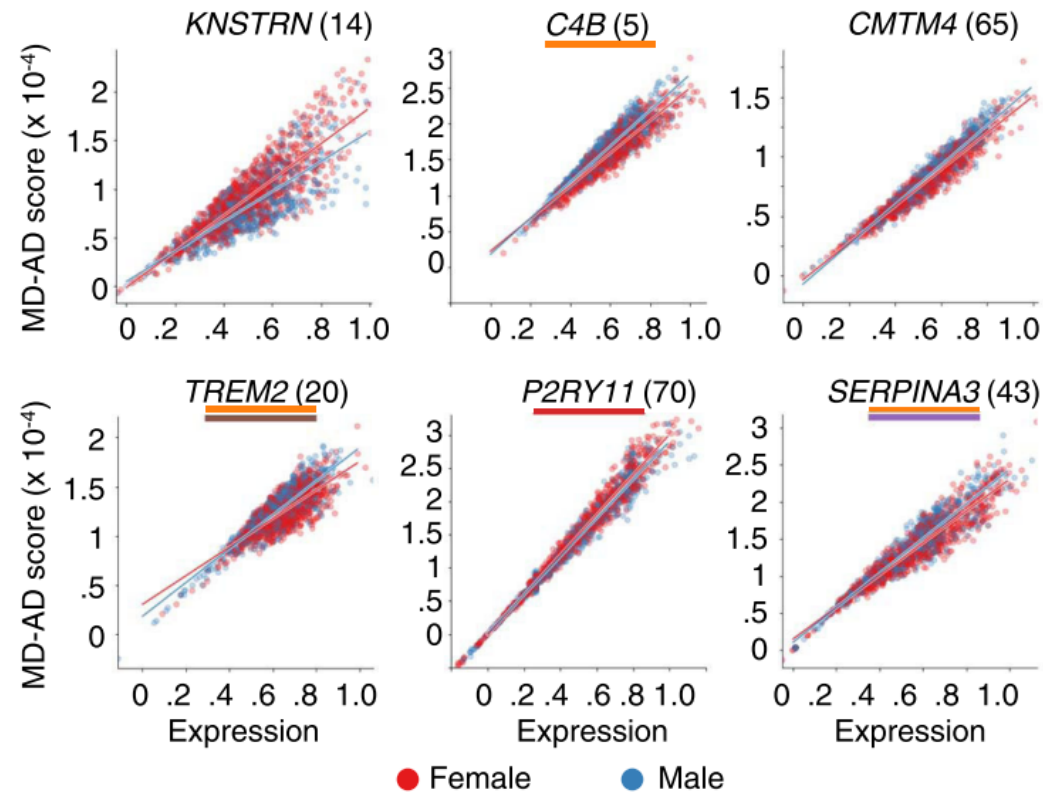
AI model



Outcomes Y

Neuropathological phenotypes:

$A\beta$, Tau, CERAD score,
Plaque counts, Braak
stage, tangle counts

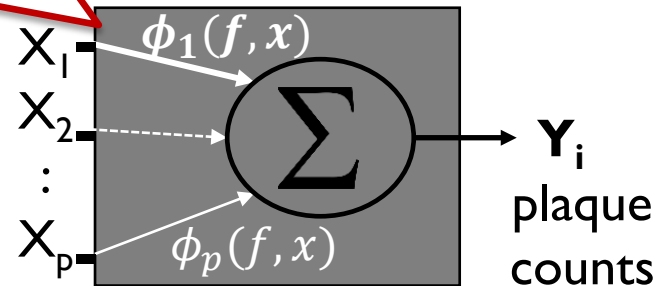


MSBB
($n=879$)

Demographic
Data

ROSMAP
($n=542$)

Clinical
data



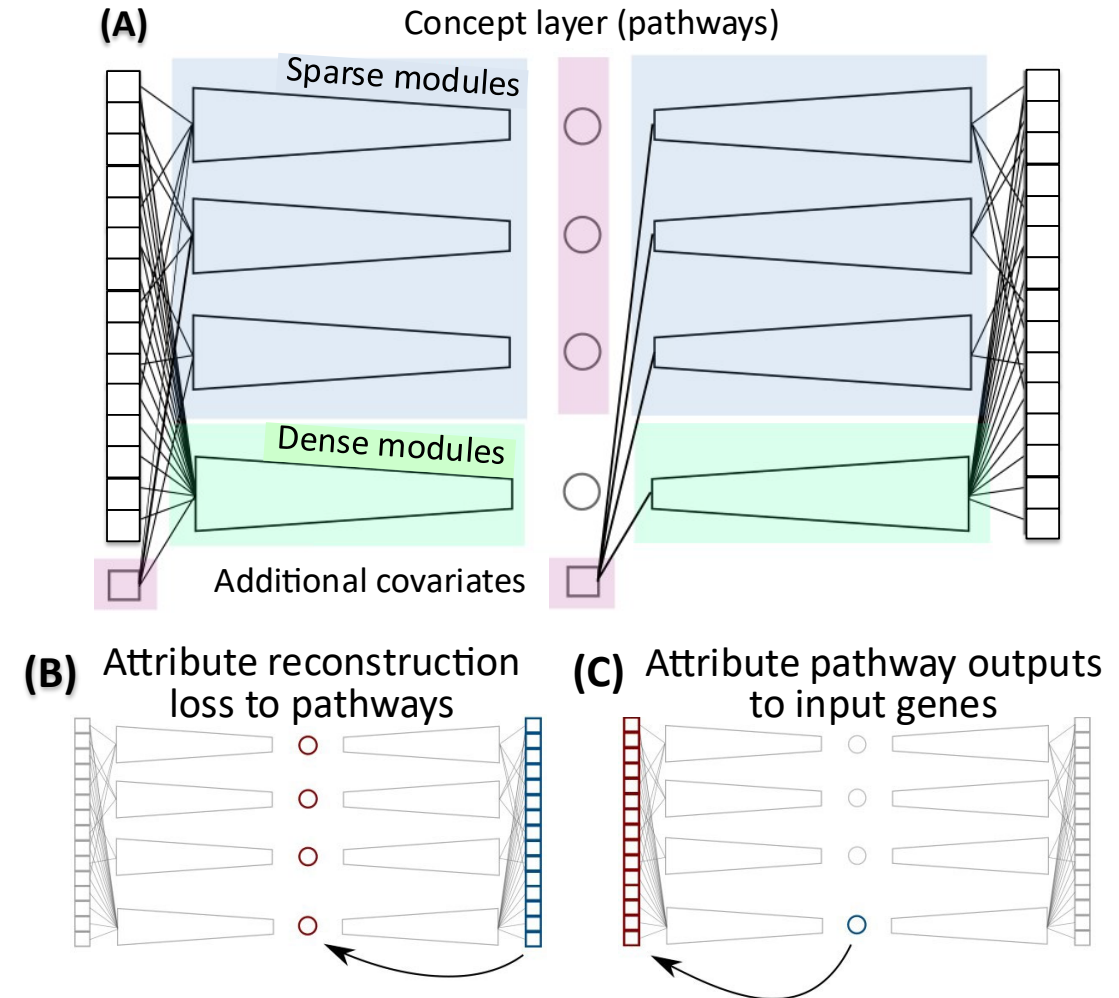
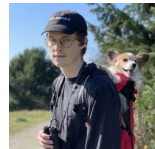
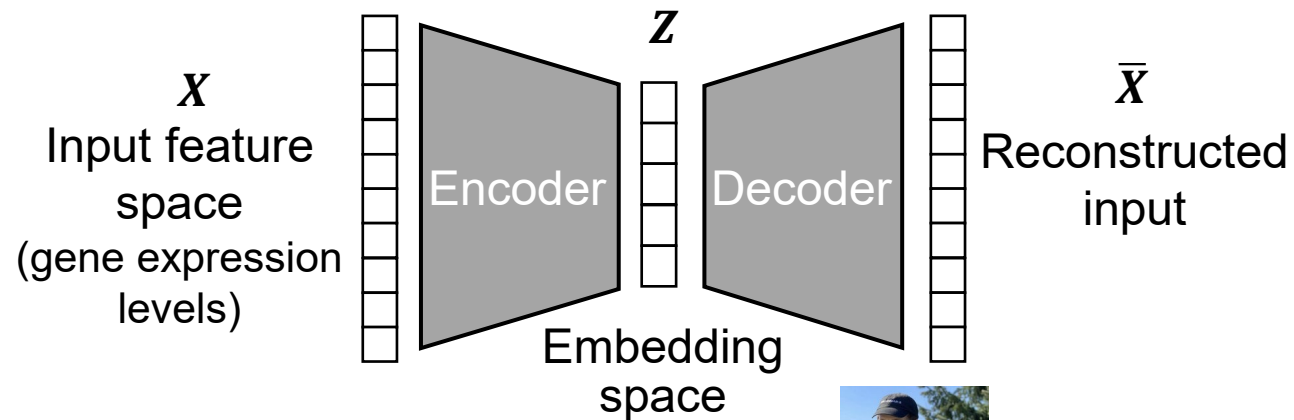
- Estimate each gene's contribution to AD neuropathologies
 - Previously unknown **sex-differential immune response (*microglia activity*)** in AD



CSE PhD'22

Biologically interpretable AI modeling further advances data-driven discovery

- Individual genes are not as interpretable as functional units (e.g., pathway)
- Unsupervised modeling enables the incorporation of unlabeled data
 - XAI can pinpoint important genes that explain the expression variation within the dataset



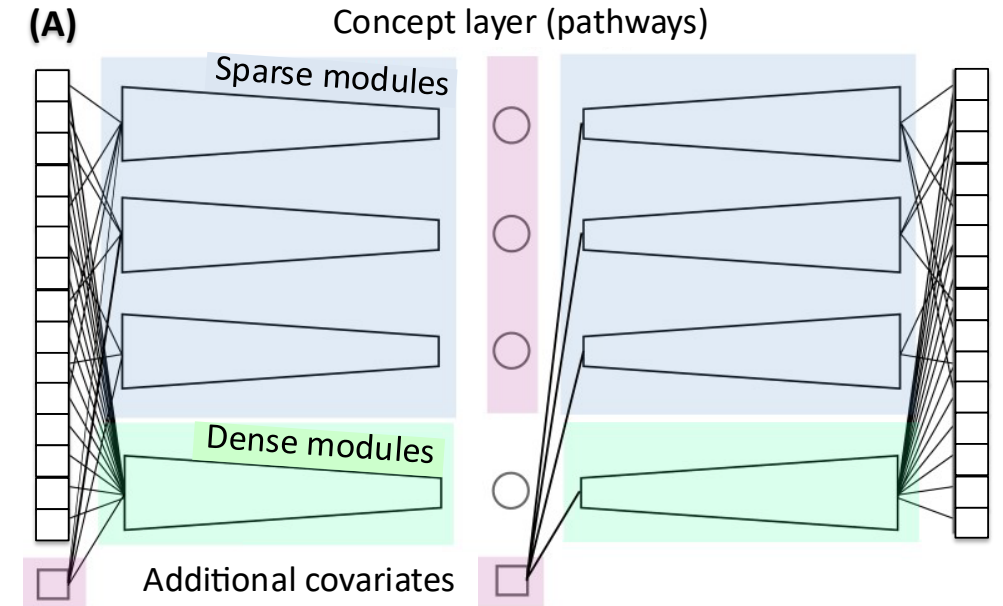
Biologically interpretable AI modeling identifies experimentally validated AD therapeutic targets

- We applied our approach to extended bulk RNAseq datasets from AD study cohorts
- We identified **mitochondrial complex I** as a potential mediator for tolerance to A β toxicity
 - *In vivo* validation in a transgenic *C. elegans* model expressing A β done by Matt Kaeberlein's lab

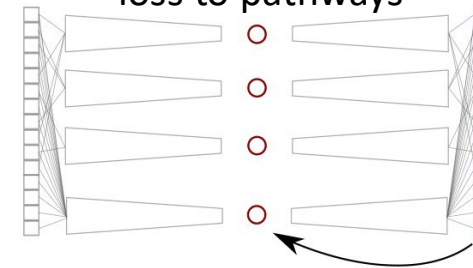
A promising pharmacological avenue!



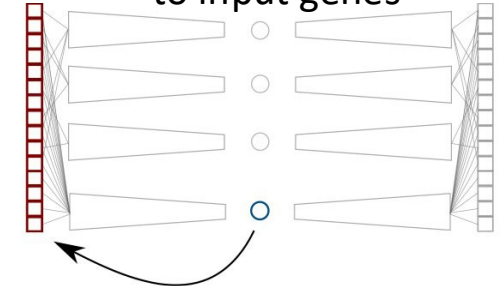
Capsaicin



(B) Attribute reconstruction loss to pathways



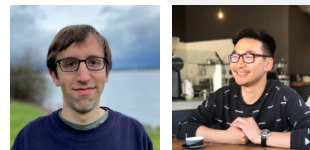
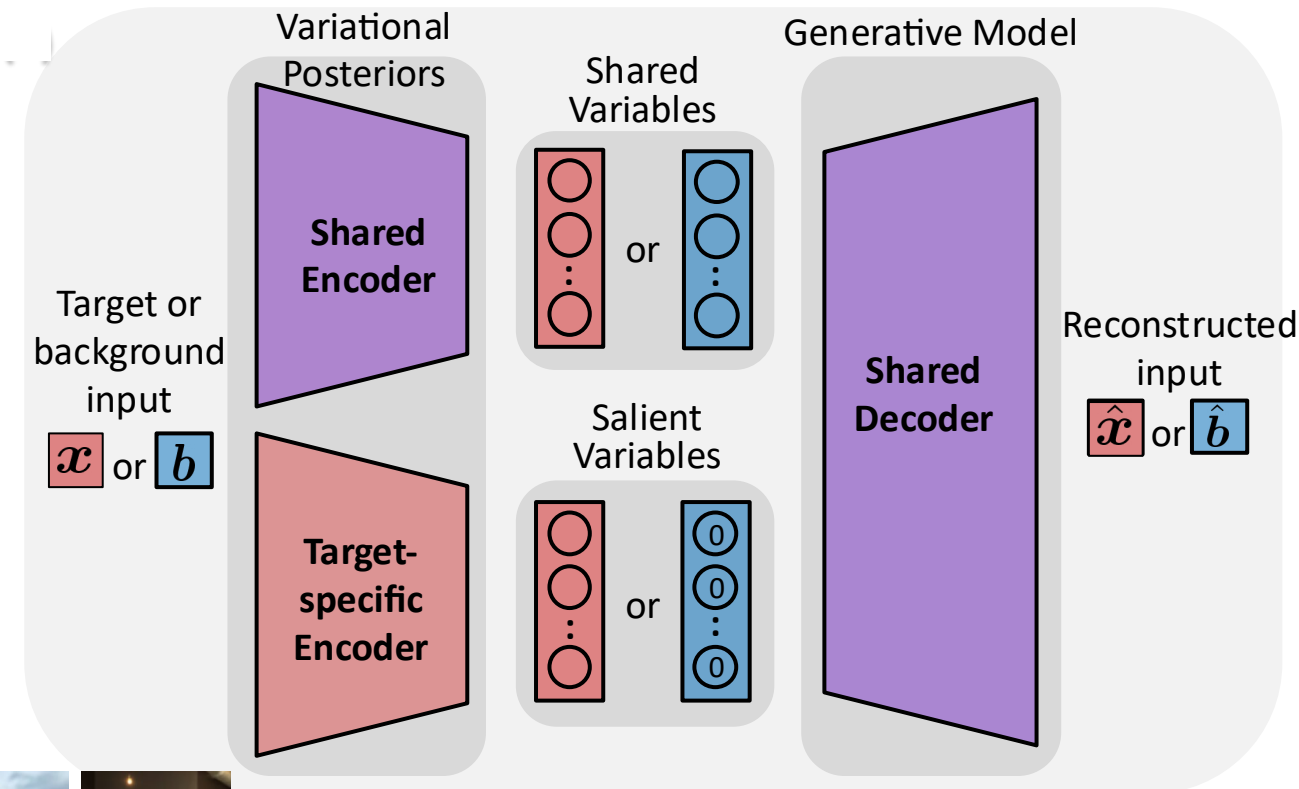
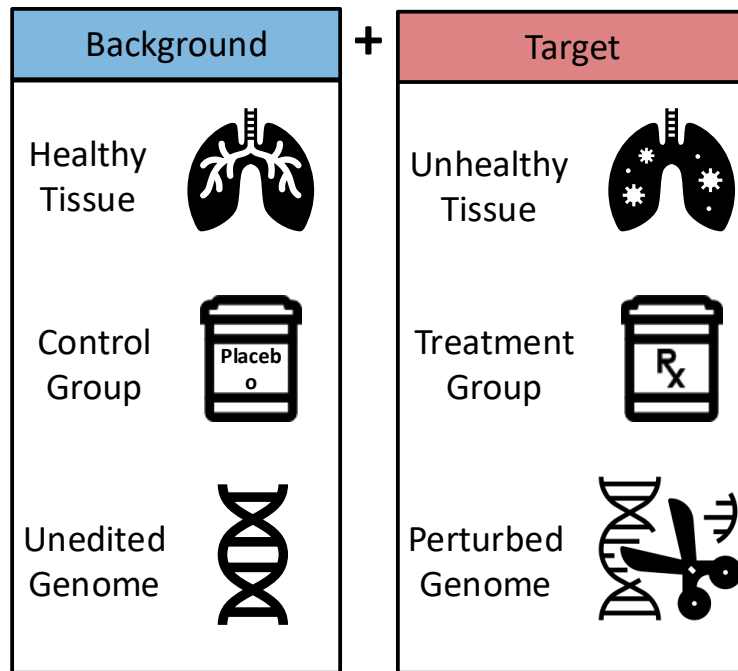
(C) Attribute pathway outputs to input genes



Contrastive modeling enhances interpretability

- Single-cell datasets are often collected to investigate differences in cellular state between **background** cells and those under specific **treatments**

(A) Contrastive Analysis



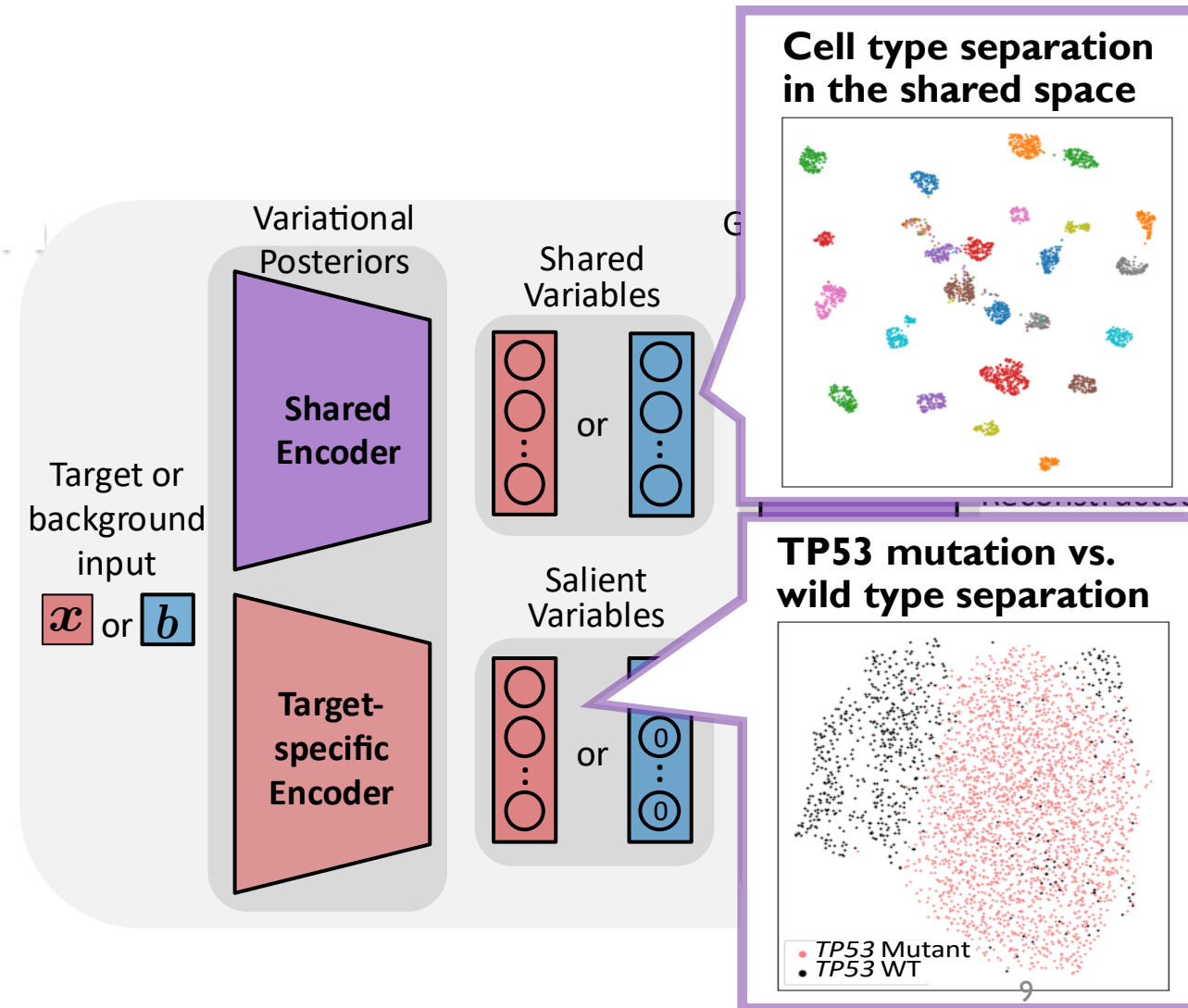
Contrastive modeling enhances interpretability

- Cancer cells treated with idasanutlin vs. untreated as background

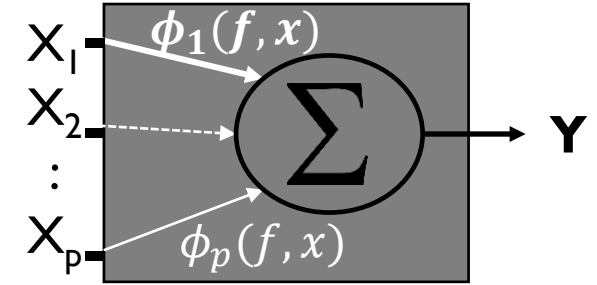
- Cells behave differently in salient space depending on their TP53 mutation status

Important implications for personalized medicine!

- How about AD vs. control brain tissue?
 - What drives neurodegeneration (in collaboration with Jessica Young)
 - What drives biological aging process? (Jessica Young & Suman Jayadev)



Outline – Two parts



- Part 1 – Identifying disease-driving genes and processes
 - Unveiling neurodegenerative disease insights with explainable AI
[*Nature Comm*'21 ; *Genome Biology*'23; *Nature Methods*'23]
- Part 2 – Auditing AI models
 - Feature attributions [*Nature MI*'21 featured in *Nature*'22]
 - Counterfactual explanations (via generative AI) [*Nature BME*'25 cover; *Lancet*'24]
 - Concept-based explanations (via foundation models) [*Nature Medicine*'24]



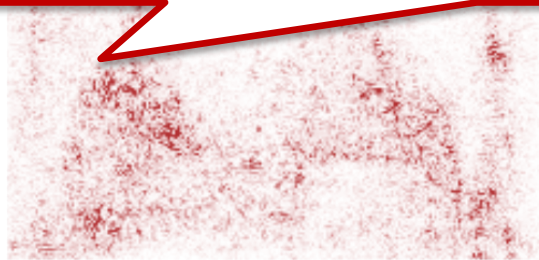
Our explainable AI techniques are generalizable to a wide range of biomedical problems, including regenerative medicine.

Auditing AI for COVID-19 detection using XAI

- Many published AI models that detect COVID-19

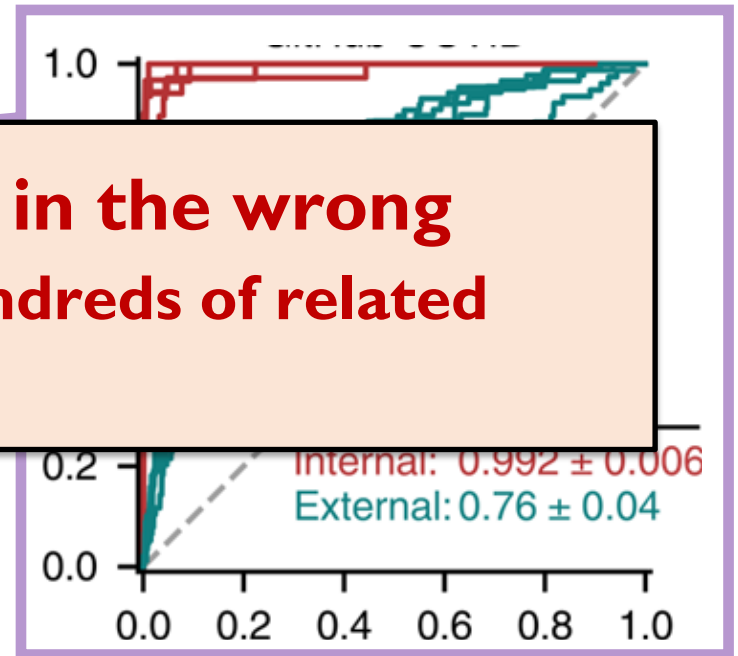
XAI helped us to stop the field from moving in the wrong direction – There were 6 published papers and hundreds of related models out there that learned the shortcuts.

Many kinds of analyses for model auditing presented in the paper!



99th

0th



- ✓ Clear lung bases predict negative COVID-19 status
- ✗ laterality markers should not predict negative status
- ✗ medical devices should not predict negative status

MSTP / CSE PhD



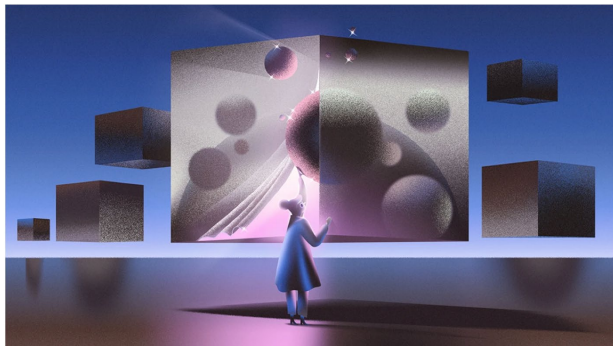
Our AI auditing work featured in *Nature*

- “Breaking into the black box of artificial intelligence” *Nature* Outlook

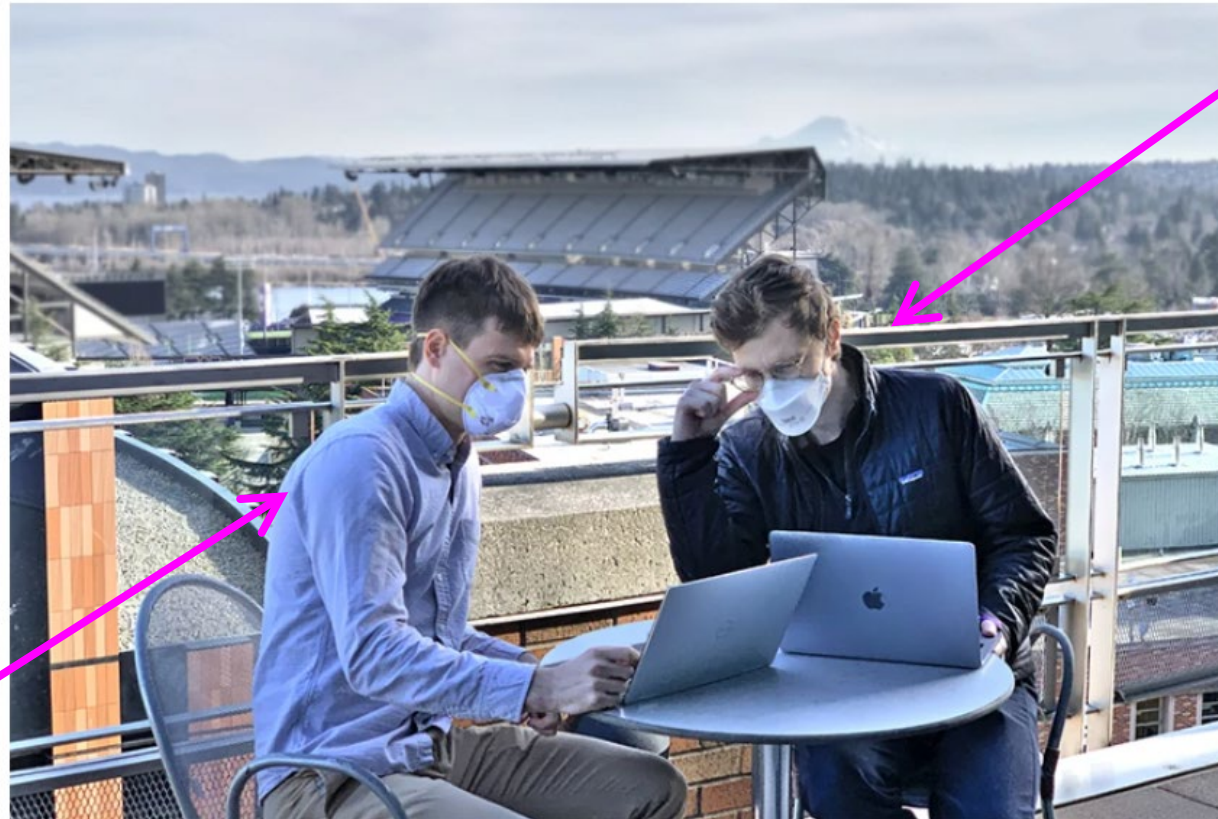
Breaking into the black box of artificial intelligence

Scientists are finding ways to explain the inner workings of complex machine-learning models.

By Neil Savage



UW MSTP/CSE PhD student **Alex DeGrave**



UW MSTP / CSE PhD
Joe Janizek (Just
matched to Stanford)

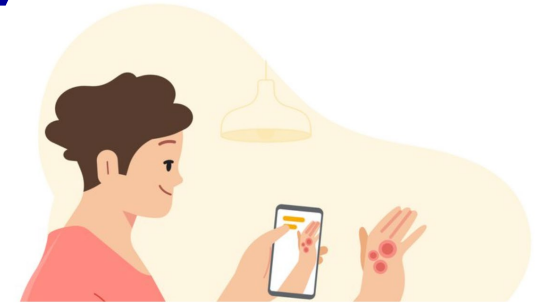
Alex DeGrave and Joseph Janizek are students on the Medical Scientist Training Program at the University of Washington, in Seattle. Credit: Alex DeGrave

DeGrave,* Janizek* et al. ***Nature Machine Intelligence***, 2021



Further digging into the flaws in the reasoning processes of clinical AI – dermatology

- Auditing AI models to predict skin cancer
 - Five models – 2 academic models, 2 commercial devices, and 1 competition winner
- Technical challenges – saliency maps often do not work

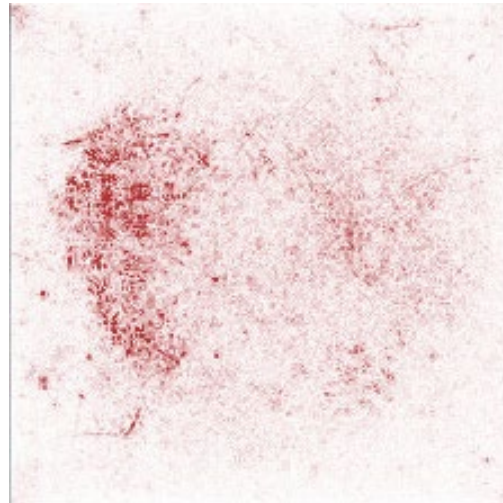


Original image

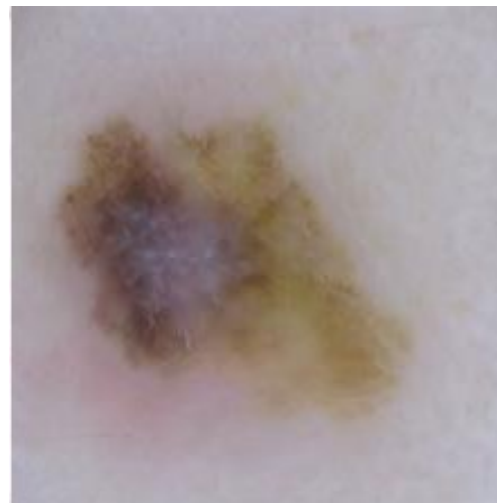


Predicted: benign

Saliency map



Modified image



Predicted: malignant

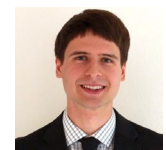
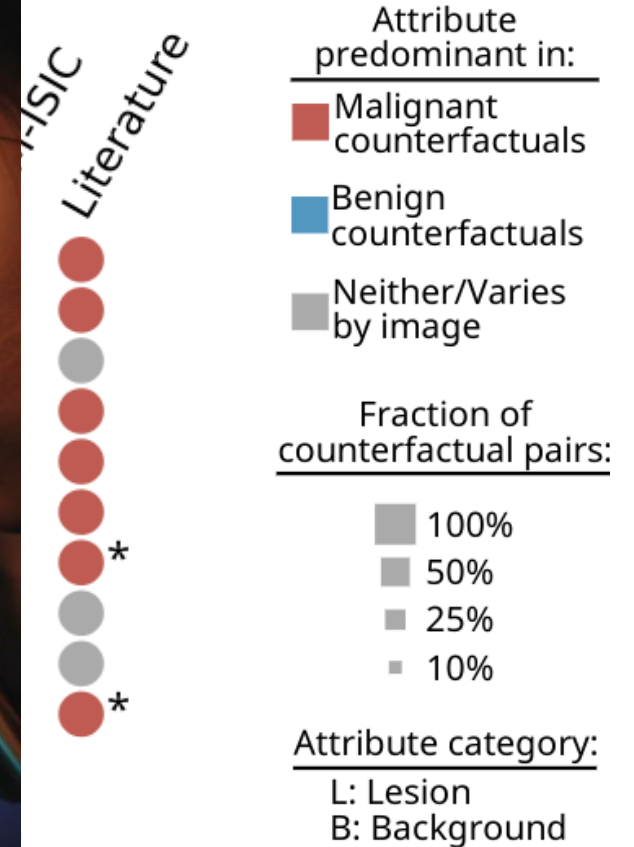
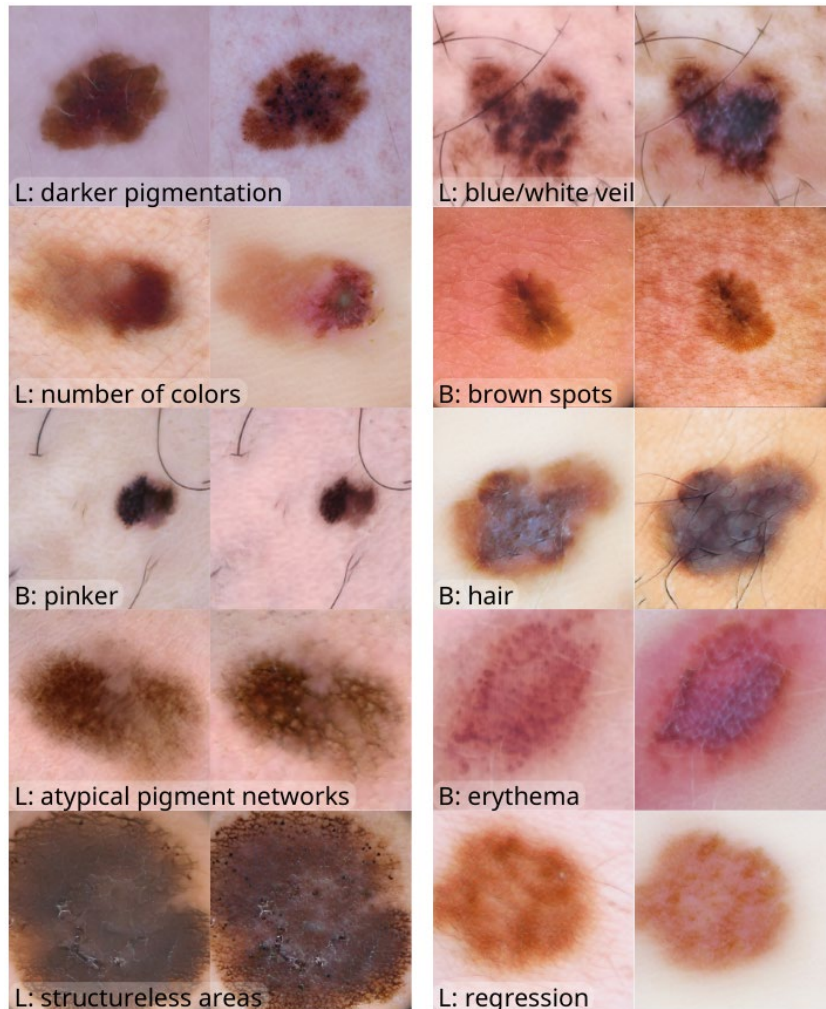
Our solution:

- Generate **counterfactual images from the AI model**
- Systematic characterization by experts: Drs. Roxana Daneshjou, and Zhuo Ran Cai (Stanford)

DeGrave et al. (*Nature Biomedical Engineering*)

Kim et al. (*Nature Medicine*, 2024)

How do dermatology AI systems make decisions on dermoscopic images



Alex,
MSTP /
CSE PhD

The *Lancet* perspective

“The clinical potential of counterfactual AI”

by Su-In Lee and Eric Topol

Digital medicine

The clinical potential of counterfactual AI models

Clinicians frequently use conditional reasoning for treatment decisions by envisioning potential outcomes for patients. This is counterfactual thinking, exploring “what if” scenarios. Developments in generative artificial intelligence (AI) enable us to simulate this patient-level reasoning at the data level, opening new opportunities for science and health care. We term this approach counterfactual AI.

This approach is exemplified by use of counterfactual images in dermatology. Using AI, original skin images were modified to resemble melanoma guided by the decision-making process of a particular AI-based dermatological classifier. Dermatologists were then tasked with identifying clinically relevant features in the counterfactual images of melanoma and normal conditions. This process elucidated the reasoning processes of five AI-based dermatological classifiers. This data-centric counterfactual AI aligns the reasoning processes of AI classifiers with human clinicians’ intuition, establishing a new approach to auditing clinical AI classifiers. Model auditing provides insights into the performance of deployed clinical AI classifiers for patients, clinicians, regulators, and data scientists.

AI model focuses on (figure). Yet they provide only a partial view of the inner workings of complex AI models, impeding efforts to identify flaws in clinical AI reasoning processes. Counterfactual AI expands the scope of explainable AI by providing counterfactual images that elicit specific outcome predictions from complex AI classifiers (figure), enabling humans to grasp more comprehensive insights into the reasoning processes of these classifiers. Collaborating with clinicians, counterfactual AI could unearth previously unnoticed image attributes. Research indicates that by partnering with AI methods capable of automatically annotating images with an array of semantically meaningful concepts, counterfactual AI can systematically probe AI classifiers about how these concepts affect their decision-making processes.

Counterfactual AI in medicine faces ethical concerns and challenges related to fairness, data quality, and generalisability. Obtaining high-quality, diverse datasets is difficult. Generalising to new data is also problematic, particularly across diverse patient populations and health-care settings. Moreover, ethical and regulatory issues, including patient privacy concerns about the use of training data, must be



Further reading

DeGrave AJ, Cai ZR, Janizei Daneshjou R, Lee SI. Audit inference processes of machine image classifiers by leveraging generative AI and the expertise of physicians. *Nat Biomed* 2023; published online December 12. <https://doi.org/10.1038/s41551-023-01160-9>

Nature Reviews Bioengineering, 2025

Medical AI transparency in the entire life cycle

nature reviews bioengineering

<https://doi.org/10.1038/s44222-025-00363-w>

Review article

Check for updates

Transparency of medical artificial intelligence systems

Chanwoo Kim^{1,2}, Soham U. Gadgil^{1,2} & Su-In Lee¹✉

Abstract

Medical artificial intelligence (AI) systems hold promise for transforming healthcare by supporting clinical decision-making in diagnostics and treatment. The effective deployment of medical AI requires trust among key stakeholders – including patients, providers, developers and regulators – which can be built by ensuring transparency in medical AI, including in its design, operation and outcomes. However, many AI

Sections

Introduction

Data transparency

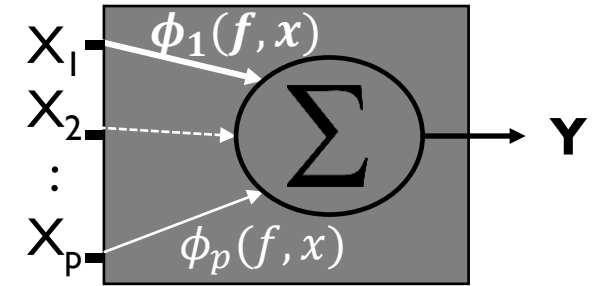
Model transparency

Deployment transparency

Outlook

Outline – Two parts

- Part 1 – Identifying disease-driving genes and processes
 - Unveiling neurodegenerative disease insights with explainable AI
[*Nature Comm*'21; *Genome Biology*'23; *Nature Methods*'23]
- Part 2 – Auditing AI models
 - Feature attributions [*Nature MI*'21 featured in *Nature*'22]
 - Counterfactual explanations (via generative AI) [*Nature BME*'25 cover; *Lancet*'24]
 - Concept-based explanations (via foundation models) [*Nature Medicine*'24]



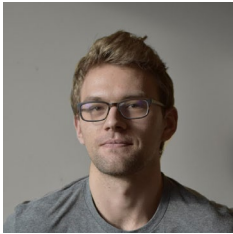
Our explainable AI techniques are generalizable to a wide range of biomedical problems, including regenerative medicine.

AI for bioMedical Sciences (AIMS) Lab

UW MSTP



Nicasia
Beebe-Wang
(CSE PhD)



Ian Covert
(former CSE
PhD)



Wei Qiu
(CSE PhD)



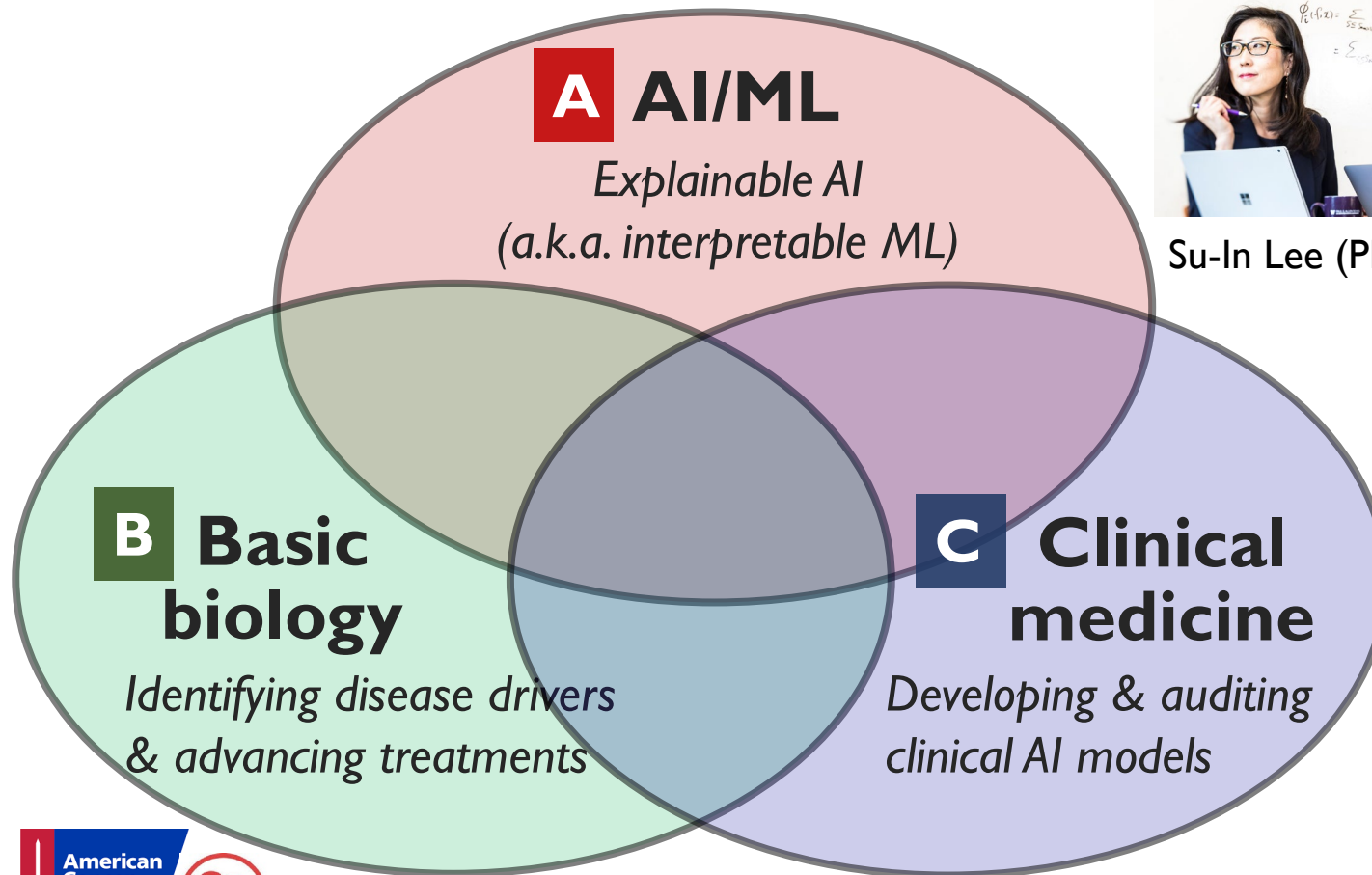
Chris Lin
(CSE PhD)



Mingyu Lu, MD
(CSE PhD)



Patrick Yu
(CSE PhD)



Su-In Lee (PI)



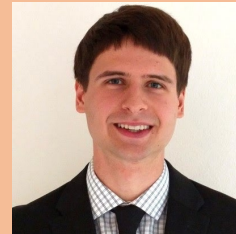
Hugh Chen
(former CSE
PhD)



Joe Janizek
(MSTP, CSE
PhD; matched
to Stanford)



Ethan
Weinberger
(CSE PhD)



Alex DeGrave
(MSTP, CSE
PhD)



Chanwoo Kim
(CSE PhD)



Soham Gadgil
(CSE PhD)



Previous members: Ben Logsdon (postdoc), Safiye Celik (CSE PhD'18), Scott Lundberg (CSE PhD'19), Parmita Mehta (CSE PhD'20), Gabe Erion (MSTP, CSE PhD'21; now Harvard Medical School for residency),