

“Am I Destined to be Me?”

A broad ‘research agenda’ presented at the Methods Hub, Faculty of Social Sciences

University of Hong Kong

9th February, 2026

Charles Rahal



- How accurately can you predict the result of a flipped coin?

- How accurately can you predict the result of a flipped coin?
- What about a spin on the roulette wheel?

- How accurately can you predict the result of a flipped coin?
- What about a spin on the roulette wheel?
- Now, consider the following quote...



“We may regard the present state of the universe as the effect of its past and the cause of its future. An intellect which at a certain moment would know all forces that set nature in motion, and all positions of all items of which nature is composed, if this intellect were also vast enough to submit these data to analysis, ... nothing would be uncertain.”

‘Essai philosophique sur les probabilités’

– Laplace (1814)



“We may regard the present state of the universe as the effect of its past and the cause of its future. An intellect which at a certain moment would know all forces that set nature in motion, and all positions of all items of which nature is composed, if this intellect were also vast enough to submit these data to analysis, ... nothing would be uncertain.”

‘Essai philosophique sur les probabilités’

– Laplace (1814)

Do you think that Laplace’s Demon – in theory – could predict the value of a flipped coin?

Laplace's Demon is a physical impossibility...

Laplace's Demon is a physical impossibility...

- **Strong or monocausal forms** of both **social determinism** and **genetic determinism** are largely rejected

Laplace's Demon is a physical impossibility...

- **Strong or monocausal forms** of both **social determinism** and **genetic determinism** are largely rejected
- But **weaker, qualified versions** of each remain widely accepted and actively researched.

Laplace's Demon is a physical impossibility...

- **Strong or monocausal forms** of both **social determinism** and **genetic determinism** are largely rejected
- But **weaker, qualified versions** of each remain widely accepted and actively researched.
- The key distinction is between *determinism as exclusivity* versus *determinism as influence*.

Laplace's Demon is a physical impossibility...

But: some things are surely more “predictable” than others, right?

Laplace's Demon is a physical impossibility...

But: some things are surely more “predictable” than others, right?

- How about the spin of the roulette wheel?

Laplace's Demon is a physical impossibility...

But: some things are surely more “predictable” than others, right?

- How about the spin of the roulette wheel?
 - Pretty hard...

Laplace's Demon is a physical impossibility...

But: some things are surely more “predictable” than others, right?

- How about the spin of the roulette wheel?
 - Pretty hard...
- Predicting whether you'll see a sports car in Hong Kong on your way home?

Laplace's Demon is a physical impossibility...

But: some things are surely more “predictable” than others, right?

- How about the spin of the roulette wheel?
 - Pretty hard...
- Predicting whether you'll see a sports car in Hong Kong on your way home?
 - Not impossible... (we can get some data and model this, though!)

Laplace's Demon is a physical impossibility...

But: some things are surely more “predictable” than others, right?

- How about the spin of the roulette wheel?
 - Pretty hard...
- Predicting whether you'll see a sports car in Hong Kong on your way home?
 - Not impossible... (we can get some data and model this, though!)
- Whether somebody will die in the next 10 years?

Laplace's Demon is a physical impossibility...

But: some things are surely more “predictable” than others, right?

- How about the spin of the roulette wheel?
 - Pretty hard...
- Predicting whether you'll see a sports car in Hong Kong on your way home?
 - Not impossible... (we can get some data and model this, though!)
- Whether somebody will die in the next 10 years?
 - Quite easy with the right data (e.g. age) and methods?

Laplace's Demon is a physical impossibility...

But: some things are surely more “predictable” than others, right?

- How about the spin of the roulette wheel?
 - Pretty hard...
- Predicting whether you'll see a sports car in Hong Kong on your way home?
 - Not impossible... (we can get some data and model this, though!)
- Whether somebody will die in the next 10 years?
 - Pretty easy with the right data and methods?

(n.b: even if the Demon wasn't impossible, Laplacian determinism isn't necessarily incompatible with “free will”)

Because if things are predictable, we can “intervene”!

We can predict that people are going to die:

- So lets intervene and help them live longer, healthier lives!

The presence of high blood pressure predicts more heart attacks...

- So lets intervene and give some statins!

These aren't controversial ideas, right?

Because if things are predictable, we can “intervene”!

We can predict that people are going to die:

- So lets intervene and help them live longer, healthier lives!

The presence of high blood pressure predicts more heart attacks..

- So lets intervene and give some statins!

These aren't controversial ideas, right?

But what if we could predict everything, in order to positively intervene?

Researchers are increasingly turning to “prediction challenges”...

Fragile Families Challenge Mass Collaboration Paper Published

May 29, 2020

If hundreds of scientists created predictive algorithms with high-quality data, how well would the best predict life outcomes? Not very well. The paper summarizing the methods and results of the FFCWS Challenge led by Matt Salganik and Ian Lundberg has been published in [**Proceedings of the National Academy of Science**](#). [**Socius**](#) has also published [a special collection](#) of 12 articles describing participants' approaches to predicting these six outcomes as well as 3 articles describing methodological and procedural insights from running the Challenge.

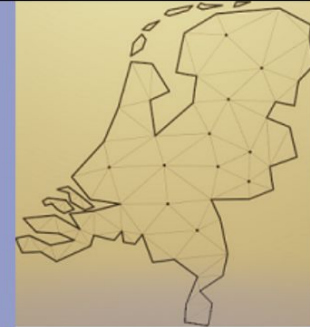


Researchers are increasingly turning to “prediction challenges”...

[FAQ](#)[PAPER](#)[GUIDES](#)[ORGANIZERS ▼](#)

Welcome to the website of the **Predicting Fertility data challenge (PreFer)**. The aim of the challenge is to measure current predictability of fertility outcomes in the Netherlands to advance our understanding of fertility

The challenge has now ended. We are analysing the results and will be sharing updates on this website. If you have questions about PreFer, please check the [FAQ](#) or [contact us](#)



Researchers are increasingly turning to “prediction challenges”...



December 1, 2025



PARTICIPATE IN THE ‘GOOD LIFE’ DATA CHALLENGE!

The LIVES Centre is launching the ‘Good Life’ Data Challenge, a large-scale collaboration using the Swiss Household Panel (SHP) to address a key question: What predicts the feeling of having lived a happy,

meaningful, and interesting (psychologically rich) life thus far? Full details of the call for participation can be found by following the link below:

<https://www.centre-lives.ch/fr/appel/call-participation-good-life-data-challenge>

We invite researchers from across the social sciences to submit a theory-driven proposal (600-800 words) by **February 15, 2026**, using the online form provided in the call. Proposals may use any variables from the 1999-2025 waves of the SHP to predict responses to three new items included in the 2025 wave (data will only become available in 2026), in which respondents provide retrospective assessments of happiness, meaning, and psychological richness.

This relates to the idea of “predictability hypotheses”:



The Oxford Handbook of the Sociology of Machine Learning

Christian Borch (ed.), Juan Pablo Pardo-Guerra (ed.)

Search in this book

Buy on Amazon

Disclosure

Contents

- Front Matter
- Part I The Past, Present, and Future of Machine Learning in Sociology
- ▼ Part II Machine Learning as a Methodological Toolbox

CHAPTER

13 Predictability Hypotheses: A Metatheoretical and Methodological Introduction

Austin van Loon

<https://doi.org/10.1093/oxfordhb/9780197653609.013.4> Pages 227–244

Published: 18 December 2023

Annotate Cite Permissions Share ▼

Abstract

Perhaps owing to decades of metatheoretical work done during a time of methodological and computational limitations, it is a widespread belief that machine learning and deductive hypothesis testing are fundamentally incompatible. This chapter argues that the more flexible estimators provided by machine learning are in fact invaluable tools for testing pre-specified social theories, especially sociological ones. This chapter defines a class of hypotheses called *predictability hypotheses*, or general theoretical statements that posit some unspecified (though potentially causal) relationship between two or more theoretical constructs. It then points to several examples of predictability hypotheses already being used by social scientists and describes how and when they are useful. It then considers some methodological implications for testing predictability hypotheses and concludes by making predictions about their future in sociology.

Metrics

[View Metrics](#)

Email alerts

- New books
- New reference articles
- New articles
- Activity related to this book

[Sign up for marketing](#)

More from Oxford Academic

- Social Research and Statistics
- Social Sciences
- Social Theory
- Sociology
- Books
- Journals

Social data scientists increasingly make statements about how ‘predictable’ things are.

PNAS

RESEARCH ARTICLE

| SOCIAL SCIENCES

 OPEN ACCESS

The origins of unpredictability in life outcome prediction tasks

Ian Lundberg  ^{a,1}, Rachel Brown-Weinstock ^b, Susan Clampet-Lundquist ^c, Sarah Pachman ^d, Timothy J. Nelson ^b, Vicki Yang ^b, Kathryn Edin ^{b,d,1}, and Matthew J. Salganik ^{b,d,e}

Abstract

Why are some life outcomes difficult to predict? We investigated this question through in-depth qualitative interviews with 40 families sampled from a multidecade longitudinal study. Our sampling and interviewing process was informed by the earlier efforts of hundreds of researchers to predict life outcomes for participants in this study. The

Health data scientists increasingly make statements about how ‘predictable’ things are.



Machine Learning with Applications

Volume 15, March 2024, 100516



Machine learning for an explainable cost prediction of medical insurance

Ugochukwu Orji^a, Elochukwu Ukwandu^b  

insurance has the advantage. when an individual considers the medical insurance policy in a financial plan, it can reduce the trouble. The safety assurance against rising healthcare charges is a remarkable benefit of medical insurance.

- **Insurance coverage for healthcare emergencies** - health emergencies are **unpredictable** and can occur at random. With a good medical insurance policy, one need not worry about the expenses, but rather concentrate and full recovery from the illness.
- **Improvements in medical technology and capacity** - according to a study by Weisbrod (1991), the expansion of health insurance coverage over the years has created incentives to

We argue that they should not be doing this!

Imagine you have a prediction task...

For example, let's imagine we are predicting educational attainment.

Pick your favourite model evaluation metric. This could be 'accuracy', Brier Skill Score, ROC-AUC, R^2 or anything else.

Here is your model performance:

$$\mathbf{BSS = 0.85}$$

The best possible BSS=1, where we have perfectly predicted everything.

Imagine you have a prediction task...

For example, let's imagine we are predicting educational attainment.

Pick your favourite model evaluation metric. This could be 'accuracy', Brier Skill Score, ROC-AUC, R^2 or anything else.

Here is your model performance:

$$\text{BSS} = 0.85$$

The best possible BSS=1, where we have perfectly predicted everything.

Then...

- 1. What's the 0.15, the gap between your current prediction and the best possible prediction?**

Imagine you have a prediction task...

For example, let's imagine we are predicting educational attainment.

Pick your favourite model evaluation metric. This could be 'accuracy', Brier Skill Score, ROC-AUC, R^2 or anything else.

Here is your model performance:

$$\text{BSS} = 0.85$$

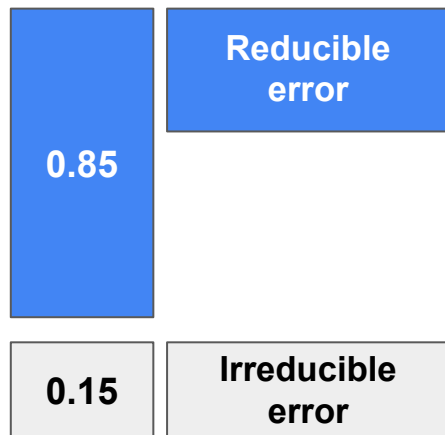
The best possible BSS=1, where we have perfectly predicted everything.

Then...

- 1. What's the 0.15, the gap between your current prediction and the best possible prediction?**
- 2. Are we ever able to achieve the 'perfect' score as defined by any common metric?**

Existing error decomposition (e.g., Hastie et al. 2001)

ROC-AUC



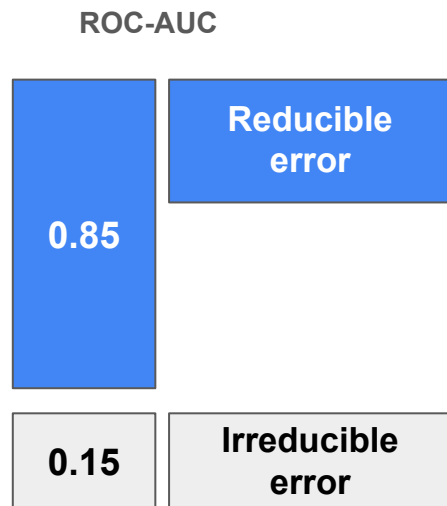
$$\begin{aligned}\text{EPE}_k(x_0) &= \text{E}[(Y - \hat{f}_k(x_0))^2 | X = x_0] \\ &= \sigma^2 + [\text{Bias}^2(\hat{f}_k(x_0)) + \text{Var}_{\mathcal{T}}(\hat{f}_k(x_0))] \end{aligned} \quad (2.46)$$

$$= \sigma^2 + \left[f(x_0) - \frac{1}{k} \sum_{\ell=1}^k f(x_{(\ell)}) \right]^2 + \frac{\sigma^2}{k}. \quad (2.47)$$



How can we better systematically think about prediction errors?

Existing error decomposition (e.g., Hastie et al. 2001)



From  to  :

What determines accuracy?

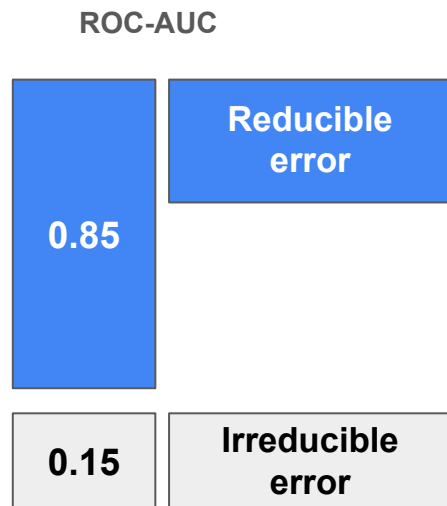
1. Algorithms;
2. Dimensionality;
3. Construction of the target feature

What is 'irreducible' today may not be irreducible tomorrow!



How can we better systematically think about prediction errors?

Existing error decomposition (e.g., Hastie et al. 2001)



From  to  :

What determines accuracy?

1. Algorithms;
2. Dimensionality;
3. Construction of the target feature

What is 'irreducible' today may not be irreducible tomorrow!

Everything is EVOLVING!

Including data, algorithms and the construction of the target feature itself



How can we better systematically think about prediction errors?

nature computational science

Explore content ▾

About the journal ▾

Publish with us ▾

[nature](#) > [nature computational science](#) > [correspondence](#) > article

Correspondence | Published: 04 March 2025

On the unknowable limits to prediction

[Jiani Yan](#)  & [Charles Rahal](#) 

[Nature Computational Science](#) **5**, 188–190 (2025) | [Cite this article](#)

621 Accesses | **1** Citations | **7** Altmetric | [Metrics](#)

Social (data) scientists must conceptualise error differently/more reflexively:

Aleatoric error

Inherent randomness can never be modelled
(genuine stochasticity)



Social (data) scientists must conceptualise error differently/more reflexively:

Aleatoric error

Inherent randomness can never be modelled
(genuine stochasticity)



Epistemic error

In theory eliminable in the limit of human progress
(resulting from a lack of knowledge)



A typical prediction task workflow

1. **Form a research question:** e.g., *'How predictable is GPA/income?'*
2. **Read relevant paper & theories, identify covariates:** e.g., *social determinants of health.*
3. **Find appropriate/ collect data for the research question:** e.g., *biobanks, surveys.*
4. **Use proper algorithm(s) to run the prediction task:** e.g., *a regression, or an LLM.*
5. **Analyse outcomes:** e.g., *any of your favourite model evaluation metrics.*

A typical prediction task workflow

1. **Form a research question:** e.g., *'How predictable is GPA/income?'*
2. **Read relevant paper & theories, identify covariates:** e.g., *social determinants of health.*
3. **Find appropriate/ collect data for the research question:** e.g., *biobanks, surveys.*
4. **Use proper algorithm(s) to run the prediction task:** e.g., *a regression, or an LLM.*
5. **Analyse outcomes:** e.g., *any of your favourite model evaluation metrics.*

***Once data, research object and covariates are specified, the research design is confirmed;
The aleatoric error of the research object target feature becomes fixed.***

Table 2: Evaluation of Model Performance in Death Prediction over Ten Seeds

Metrics	HRS			SHARE			ELSA		
	SL	LightGBM	LR	SL	LightGBM	LR	SL	LightGBM	LR
IMV	0.192	0.195	0.111	0.073	0.077	-0.055	0.056	0.060	0.068
ROC-AUC	0.814	0.816	0.823	0.789	0.795	0.453	0.823	0.828	0.837
PR-AUC	0.687	0.691	0.702	0.508	0.522	0.172	0.491	0.499	0.521
Pseudo R ²	0.275	0.280	0.295	0.195	0.207	-0.041	0.208	0.221	0.243
FFC R ²	0.602	0.605	0.613	0.443	0.451	0.279	0.394	0.404	0.421
Test IP	0.317	0.317	0.317	0.202	0.202	0.202	0.155	0.155	0.155



Is 0.823 the predictive ceiling? Can I ever improve it?

Example taken from Yan, J. (2025), '*Revisiting the social determinants of health with explainable AI: a cross-country perspective*', American Journal of Epidemiology

$$\begin{aligned} y_{\text{true}} &= \underbrace{f^*(\mathbf{x}_{\text{true}})}_{\substack{\text{Predictive} \\ \text{ceiling}}} + \underbrace{\varepsilon}_{\substack{\text{Aleatoric} \\ \text{error}}} \\ &= \underbrace{[f^*(\mathbf{x}_{\text{true}}) - f(\mathbf{x}_{\text{true}}|y_{\text{true}})]}_{\substack{\text{Model approximation gain} \\ \text{(Epistemic)}}} \\ &\quad + \underbrace{[f(\mathbf{x}_{\text{true}}|y_{\text{true}}) - f(\mathbf{x}_{\text{true}}|y_{\text{observed}})]}_{\substack{\text{Measurement gain from } y \\ \text{(Epistemic)}}} \\ &\quad + \underbrace{[f(\mathbf{x}_{\text{true}}|y_{\text{observed}}) - f(\mathbf{x}_{\text{observed}}|y_{\text{observed}})]}_{\substack{\text{Measurement gain from } \mathbf{x} \\ \text{(Epistemic)}}} \\ &\quad + \underbrace{y_{\text{predicted}}}_{\substack{\text{Current prediction}}} + \underbrace{\varepsilon}_{\substack{\text{Irreducible} \\ \text{(Aleatoric)}}} . \end{aligned}$$

Algorithm

Y

Construct Validity

Representation

X

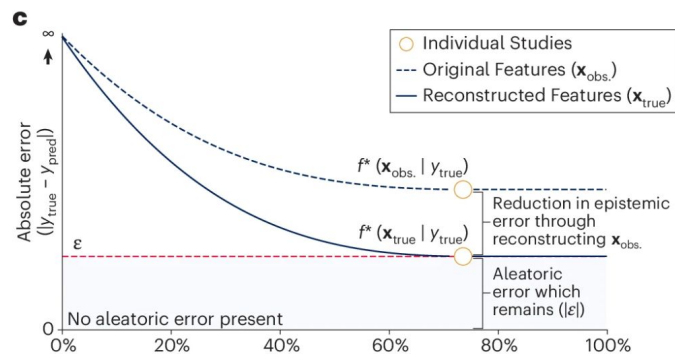
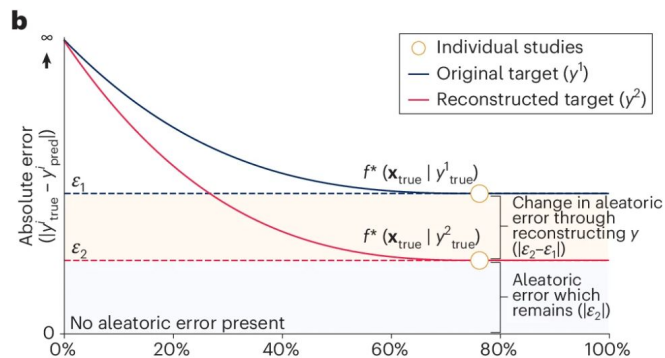
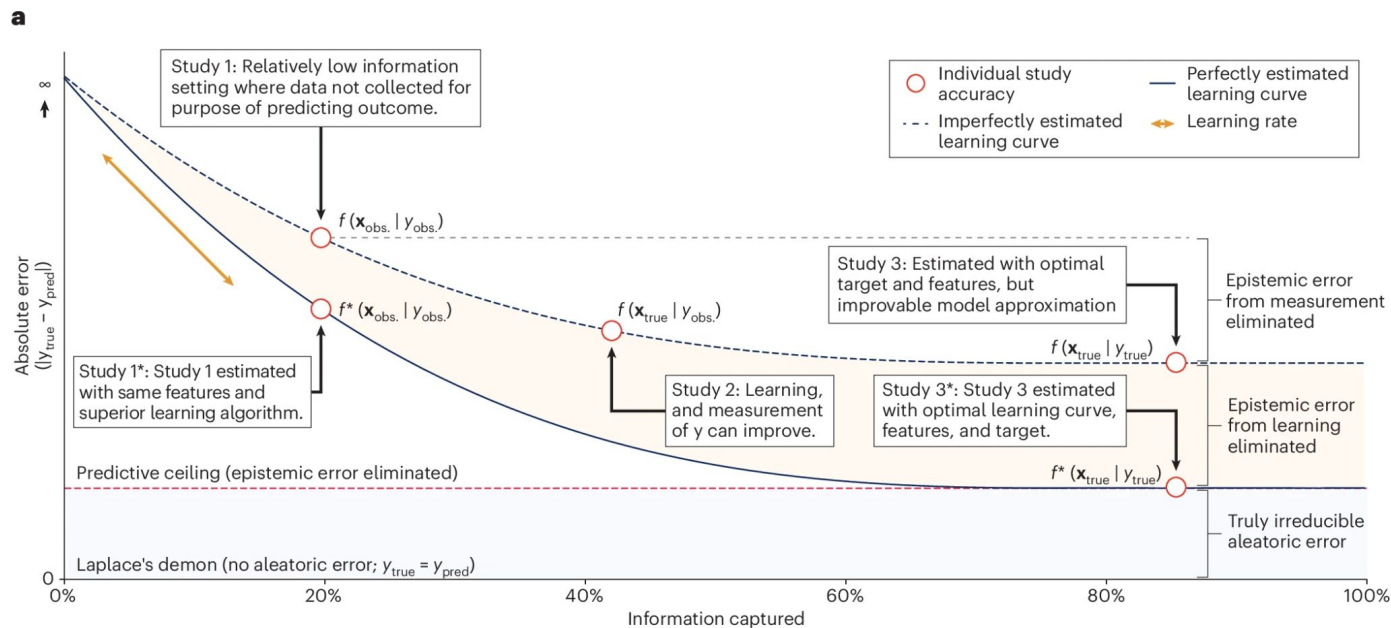
Relevant information inclusion

We claim that people should be responsibly using terms while thinking about the limits of human endeavour:

We can only say how predictable things are *conditional* on:

- 1. How our ‘target’ feature is constructed;**
- 2. Our feature set and it’s dimensionality**
- 3. The algorithm used to make the prediction**

We emphasize the importance – and develop a new style – of ‘learning curve’.



What can we do? Comprehensive algorithms and datasets

nature

View all journals | Search | Log in

Explore content | About the journal | Publish with us

Sign up for alerts | RSS feed

nature > articles > article

Article | Open access | Published: 17 September 2025

Learning the natural history of human disease with generative transformers

Artem Shmatko, Alexander Wolfgang Jung, Kumar Gaurav, Søren Brunak, Laust Hvas Mortensen, Ewan Birney, Tom Fitzgerald & Moritz Gerstung

Nature (2025) | Cite this article

146k Accesses | 3 Citations | 1446 Altmetric | Metrics

Abstract

Decision-making in healthcare relies on understanding patients' past and current health states to predict and, ultimately, change their future course^{1,2,3}. Artificial intelligence (AI) methods promise to aid this task by learning patterns of disease progression from large corpora of health records^{4,5}. However, their potential has not been fully investigated at scale. Here we modify the GPT⁶ (generative pretrained transformer) architecture to model the progression and competing nature of human diseases. We train this model, Delphi-2M, on data from 0.4 million UK Biobank participants and validate it using external data from 1.9 million Danish individuals with no change in parameters. Delphi-2M predicts the rates of more than 1,000 diseases, conditional on each individual's past disease history, with accuracy comparable to that of existing single-disease models. Delphi-2M's generative nature also enables sampling of synthetic future health trajectories, providing meaningful estimates of potential disease burden for up to 20 years, and enabling the training of AI models that have never seen actual data. Explainable AI methods⁷ provide insights into Delphi-2M's predictions, revealing clusters of co-morbidities within and across disease chapters and their time-dependent consequences on future health, but also highlight biases learnt from training data. In summary, transformer-based models appear to be well suited for predictive and generative health-related tasks, are applicable to population-scale datasets and provide insights into temporal dependencies between disease events, potentially improving the understanding of personalized health risks and informing precision medicine approaches.

Download PDF

Associated content

This AI tool predicts your risk of 1,000 diseases – by looking at your medical records
Shamini Bundell & Nick Petric Howe
Nature | Nature Podcast | 17 Sept 2025

Which diseases will you have in 20 years? This AI accurately predicts your risks
Gemma Conroy
Nature | News | 17 Sept 2025

AI uses medical records to accurately predict onset of disease 20 years into the future
Yonghui Wu
Nature | News & Views | 30 Sept 2025

Sections

Figures

References

Abstract

Main

A transformer model for health records

Modelling multi-disease incidences

Sampling future disease trajectories

Explaining Delphi-2M predictions

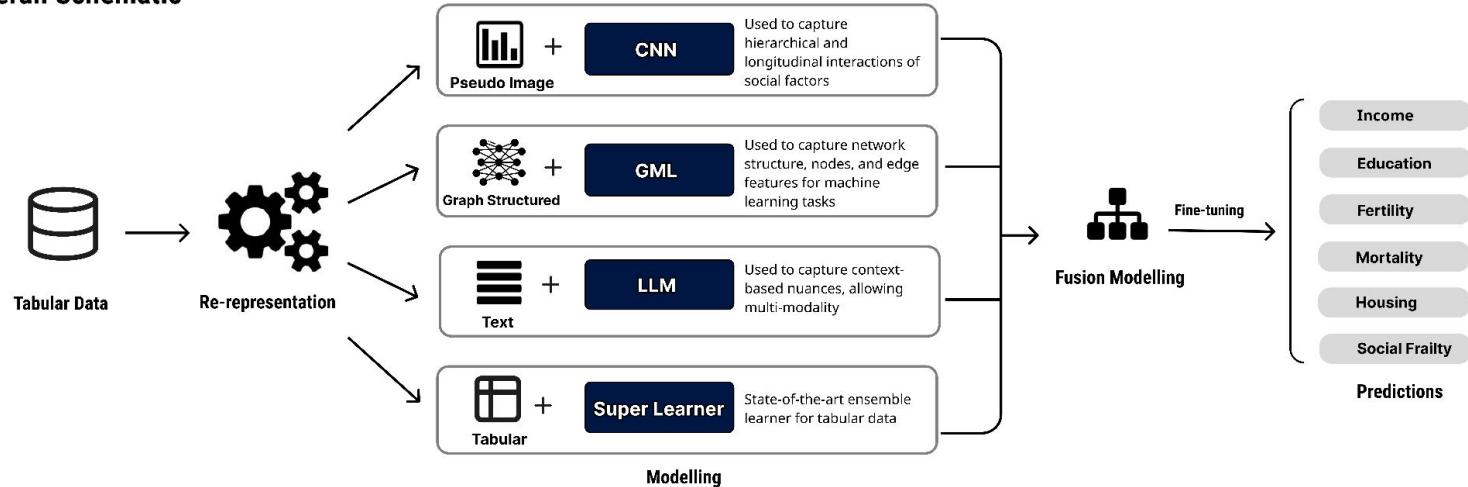
External validation and bias assessment

Discussion

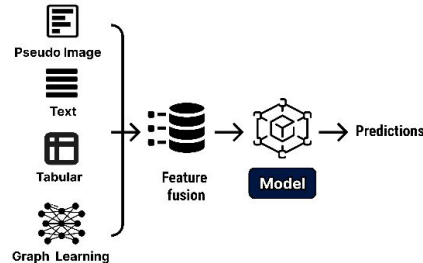
Methods

What can we do? Comprehensive algorithms and datasets

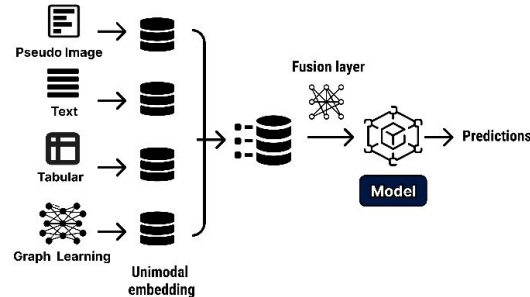
a. Overall Schematic



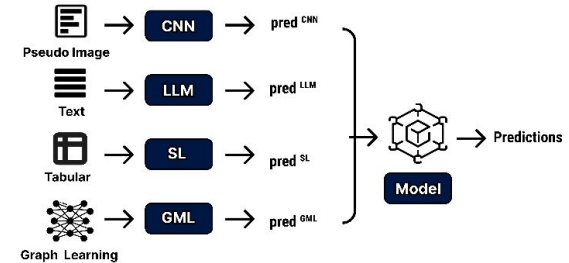
b. Early Fusion



c. Intermediate Fusion

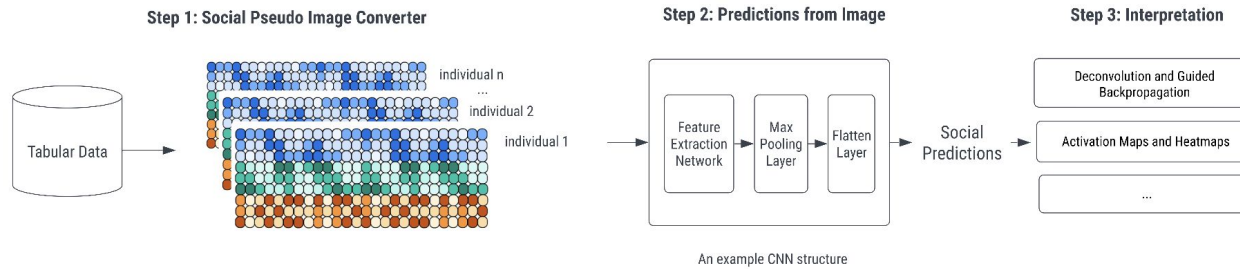


d. Late Fusion

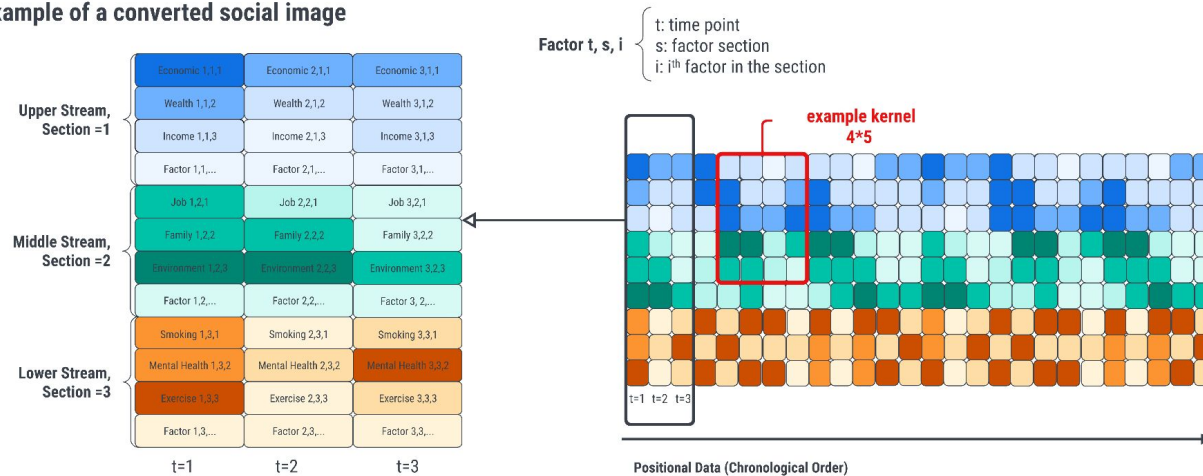


What can we do? Comprehensive algorithms and datasets

a. The schematic workflow



b. An example of a converted social image



What can we do? Comprehensive algorithms and datasets

Predictability Index

LEVERHULME
TRUST

OXFORD
POPULATION
HEALTH



Economic
and Social
Research Council



Uehiro Oxford
Institute

About Privacy Policy

Domain

Health Outcomes

Dataset

China Kadoorie Biobank (CKB)

Target Variables

Body Mass Index

Time to Event

T-1

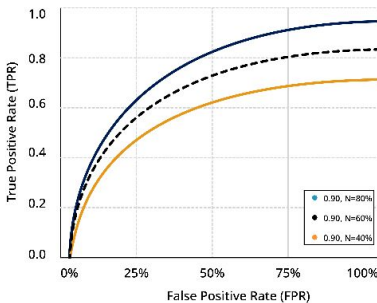
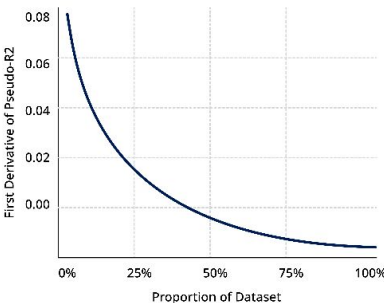
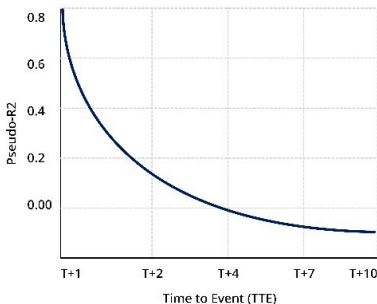
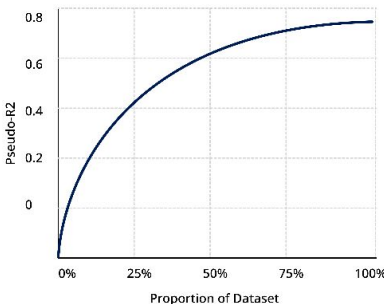
Data Representation Method

Fusion-Based Foundation Model

Introduction

Welcome to the online Predictability Index! We are a group of computational social scientists and philosophers at the University of Oxford interested in how accurately we can predict people's lives. Please don't hesitate to contact us with any thoughts or comments you have, or just to chat; predictability_index@demography.ox.ac.uk. Find us on GitHub at github.com/predictability_index.

Model Performance



Ranks

TTE	Data	Feature	Pseudo-R2
T-0	UKB	BMI	0.71
T-1	UKB	BMI	0.63
T-0	CKB	BMI	0.61
T-2	CKB	BMI	0.55
T-0	UKB	BP	0.53
T-1	CKB	BMI	0.50
T-0	UKB	ADL	0.48
T-0	UKB	ADL	0.44
T-0	UKB	CRP	0.42
T-1	UKB	CRP	0.35
T-1	UKB	ADL	0.22

Can we ever know where we are in the process?

- We have up to seven new approaches for approximating epistemic error across data, model, and measurement.
- These methods include learning curves (Mohr & van Rijn, 2021), Taylor expansions with the law of propagation of uncertainty, and bootstrap variants.
- We plan to develop a large-scale integrated simulation bench that uses user-defined data-generating mechanisms.

But what about the philosophical and ethical implications of research and policy such as this?

But what about the philosophical and ethical implications of research and policy such as this?

(I would argue this is dramatically and chronically under-appreciated)

1. Responsibility distortion under high-precision prediction



- High-precision prediction may be taken to support a deterministic view of human behaviour.
- This can weaken attributions of moral agency and responsibility for past actions.
- Conversely, prediction may justify holding individuals responsible for anticipated futures.
- Such shifts risk altering normative standards used in law, welfare, and public policy.

“With great power comes great responsibility”

2. Stigmatisation/inequality arising from predictive risk classification



- Predictive classification can label individuals or groups as “high-risk”.
- Such labels may entrench stigma and shape self-perception or external treatment.
- Existing structural inequalities may be reproduced through data and modelling choices.
- Interventions based on prediction may differentially burden already disadvantaged groups.

“With great power comes great responsibility”

3. Epistemic overconfidence in the use of predictive systems



- Apparent accuracy can obscure irreducible uncertainty and model limits.
- Decision-makers may over-rely on predictions in policy or intervention design.
- This risks legitimising actions not proportionate to the available evidence.
- Errors may be interpreted as individual failure rather than predictive limitation.

“With great power comes great responsibility”

4. Privacy/informational risks in longitudinal predictive modelling



- Longitudinal and fused data representations enable sensitive inference.
- Predictions may reveal attributes not explicitly disclosed by individuals.
- This challenges norms of consent, data minimisation, and informational control.
- Predictive capacity may outpace existing ethical and regulatory safeguards.

“With great power comes great responsibility”

Takeaway Message #1:

The future surely holds immense promise, and is an exciting area of research



- Health data scientists commonly screen on predictive risk for disease.
 - Why shouldn't social scientists predict traits like social frailty?
 - Especially if there are avenues for positive social intervention?
- **Current** methods and data show relatively low power.
 - But this is rapidly changing.
 - Partly due to the rise of LLMs and GAI.
 - Partly due to the rise of high performance computing
 - Partly due to the rise of population-level datasets.

“Until we have reason to believe that we have eliminated error through better measurement and construction of features and outputs — in addition to the elimination of learning error — we should only tentatively discuss how predictable things are with current data and technology.”

Takeaway Message #2:

Let's be careful about absolute statements when talking of current tools/tech!



- Aleatoric error is unmodelable; fixed at the point of confirming your research question.
- Aleatoric error levels can vary depending on the research object; some topics are easier to predict than others.
- Epistemic error can be eliminated through
 - Better capture of information
 - Better representativity
 - Better variable construction
 - Better relevant information inclusion
 - Better algorithms

“ Until we have reason to believe that we have eliminated error through better measurement and construction of features and outputs — in addition to the elimination of learning error — we should only tentatively discuss how predictable things are with current data and technology.”

Takeaway Message #3:

All must appreciate ethical risks of such high-precision approaches.



- **Responsibility distortion**
Prediction shifts moral judgment from realised actions to anticipated futures.
- **Stigma and inequality**
Risk labels may entrench disadvantage and reproduce structural bias.
- **Overconfidence in prediction**
Apparent accuracy can mask uncertainty and legitimise unwarranted intervention.
- **Privacy risk**
Rich life-course data enable intrusive inference beyond consent.

“ Until we have reason to believe that we have eliminated error through better measurement and construction of features and outputs — in addition to the elimination of learning error — we should only tentatively discuss how predictable things are with current data and technology.”

Thank you for your time and attention!

Please consider joining our open and online seminar series, Metrics and Models:

