

Am I Destined to Be Me?

Approximating the Limits of Social Determinism with Computational Methods

Charlie Rahal, Jiani Yan, and Tom Douglas
Nuffield College/LCDS Working Group, 05.03.2025.



- Some slides about a grant application we are working on.
- Specifically, it's for the ESRC 'Responsive Mode' (£1m).
 - Importantly: an '**AI Highlight Notice**' is currently out for this call.

Person	Role	Expertise
Charlie	PI	CSS, Data Science, Social Data Science
Jiani	Co-I	CSS, CS, SDS, Social-Epidemiology
Tom	Co-I	Philosophy, Determinism, Practical Ethics
Practical Ethicist	Specialist	Philosophy, Determinism, Practical Ethics
Data Scientist	Postdoc (Open)	ML/DL, SDS, CSS
Data Scientist	Postdoc (Open)	ML/DL, SDS, CSS
Hamza	Grant Manager	Research Contracts, Finance

- Please do not share anything from this meeting externally.

<https://doi.org/10.1038/s43588-025-00776-y>

On the unknowable limits to prediction

 Check for updates

The classic dichotomization of prediction error into statically defined 'reducible' and 'irreducible' terms is unable to capture the intricacies of progressive research design where specific types of truly reducible error are being rapidly eliminated at differential speeds. While this is generalizable true – including in the domain of large language models (LLMs) – it is none more apparent than in social systems, where questions regarding the 'predictability' of outcomes are gradually reaching the fore¹. Canonical work that convenes common predictive tasks involving conventional social surveys² shows low power, but rapidly emerging computational approaches that utilize non-standard administrative data and information contained within things such as children's autobiographical essays – often combined with generative artificial intelligence and distributed computing³⁻⁵ – begin to allow high accuracy. We propose an enhanced framework that builds upon existing work⁶ in a way that we believe helpfully decomposes various types of truly reducible 'epistemic' error, which are in theory eliminable (resulting from a lack of knowledge), and isolates residual, 'aleatoric' error (inherent randomness that can never be modeled) that research endeavor will not be able to capture in the limit of human progress. We further emphasize the impossibility of claiming whether things are 'predictable', especially with regards to open, dynamically evolving systems.

Decomposing prediction error into reducible and irreducible terms is not new. The work of Trevor J. Hastie et al.⁶ – a landmark instruction manual in the field – uses the term 'reducible' to taxonomically classify errors into a type that can be eliminated, and the term 'irreducible' to denote errors that cannot be eliminated, but only given current models and information on target variables and feature sets. For example, the insightful mixed-methods work of Ian Lundberg et al.⁷ begins its abstract with 'Why are some life outcomes difficult to predict?'. In that case, they are difficult to predict based on information that is currently measurable, and the 'task' at hand. An outcome that is difficult to predict with one feature set may not be 'unpredictable' in an aleatoric sense: life trajectories may be trivial to predict in the future.

Reflexive terminology is essential, lest conclusions be drawn that there is no space for future positive policy interventions based on ever increasingly accurate predictions that prevent negative outcomes. It is impossible to know whether we have eliminated all reducible epistemic error to approach the practical 'ceiling' of accuracy, to make statements about how predictable outcomes will become, or to understand where we are in the process. Predictability of an outcome can only be asserted when all epistemic errors are eliminated: when 'unmeasured' information is measured and constructed perfectly, features are constructed perfectly, and when functional forms are able to be exactly recovered.

Let y_{true} denote the actual, observable outcome that is measured after constructing the phenomenon of interest; ε is aleatoric error for (that specific) y_{true} . Next, denote x_{true} as the perfectly chosen, constructed and measured feature set from a possibly uncountable, infinite collection of features in the universe, and $f^*(\cdot)$ as the conceptually true underlying function that maps x_{true} onto y_{true} . As opposed to the scenario of perfect prediction described above, $y_{observed}$, $x_{observed}$ and $f(x)$ are defined as the observed target, features, and the model trained on them. Eq. (1) expresses our suggestion for conceptually decomposing predictive error under the assumption of no distributional shift:

$$\begin{aligned}
 y_{true} &= \underbrace{\hat{f}(x_{true})}_{\text{Predictive ceiling}} + \underbrace{\varepsilon}_{\text{Aleatoric error}} \\
 &= \underbrace{[\hat{f}(x_{true}) - f(x_{true}|y_{true})]}_{\text{Model approximation gain (Epistemic)}} \\
 &\quad + \underbrace{[f(x_{true}|y_{true}) - f(x_{true}|y_{observed})]}_{\text{Measurement gain from y (Epistemic)}} \quad (1) \\
 &\quad + \underbrace{[f(x_{true}|y_{observed}) - f(x_{observed}|y_{observed})]}_{\text{Measurement gain from x (Epistemic)}} \\
 &\quad + \underbrace{y_{observed}}_{\text{Current prediction}} + \underbrace{\varepsilon}_{\text{Irreducible (Aleatoric)}}
 \end{aligned}$$

As x_{true} and y_{true} are enhanced information sets over and above $x_{observed}$ and $y_{observed}$, they could taxlogically reduce what may currently have been termed 'irreducible'. Predictive accuracy could also improve through better construct validity⁸: the extent to which the model's predictions or features accurately capture the theoretical construct or concept the model is intended to measure or represent. This may not necessarily involve the collection of additional information, but simply a re-framing of a variable or new utilization of

existing information. Incorporation of previously unmeasured information (from $x_{observed}$ towards x_{true}) or better measurement of existing features (from $x_{observed}$ and $y_{observed}$ towards x_{true} and y_{true}) will also reduce epistemic error, as will reductions in model approximation error (from f towards f^*). See supplementary information for a fully delineated decomposition of Eq. (1). To illustrate how these different forms of error reduction work in practice, we depict our generalized framework in Fig. 1. This demonstrates how

- Published on Tuesday!
- It inspires a lot of our work in this area.
- Careful delineation of prediction error.
- Major application to 'open' social systems.
- Big part of this grant ('Stage 3').

nature computational science

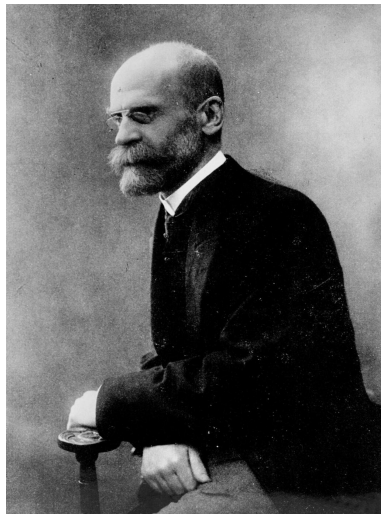
“Am I Destined to Be Me?
Cursed to Repeat it?
Never to be Free?
What Does it Mean?”

Health (2020),
‘These Days’, DISCO4: 2, 12.



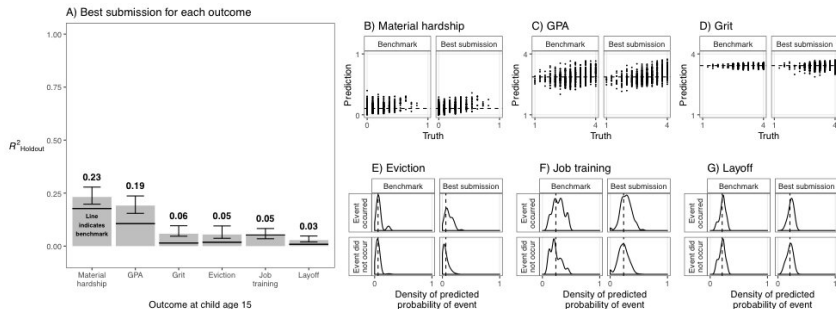
“A social fact is any way of acting, whether fixed or not, capable of exerting over the individual an external constraint”

Émile Durkheim (1895),
‘The Rules of Sociological Method’.



If this notion were true in its strongest form, one would expect high predictability across social systems with minimal stochasticity in individual lives.

But, I've seen this before!



But, I've seen this before!

The FFC is a wonderful study! But:

- There are problems with the dataset.
 - This is in terms of both dimensionality and missingness.
- There are problems with the scope ('target variables').
- There are problems with external validity.
- There are problems with 'time to event'.
- There are problems with interpretation of it.
- There are problems with the undeveloped ethical aspects of it.

It re-invigorated the field. Lundberg et al. (2024) also immensely readable.

What are we proposing? Not limited to:

- Computational re-examination with largest datasets.
- Rigorous analysis of multiple dimensions and facets.
- New frameworks for normative ethics and regulation.
- Further development and integration of new AI tools.
- Development of standalone libraries.
- Decomposition of epistemic errors based on our earlier work.
- Interactive tools for engagement.
- International comparison with focus on UKRI-funded data.

What are we proposing? Four stages:

	Output	Description	Panels	Census	Registry	Biobanks	Time Use
Stage 0	1	Normative Ethics	✓	✓	✓	✓	✓
	2	Regulatory Policy	✓	✓	✓	✓	✓
Stage 1	3	Proof-of-Concept	✓	✓	✓	✗	✗
	4	Pre-registration	✓	✓	✓	✗	✗
Stage 2	5	Autowrangler	✓	✓	✓	✓	✓
	6	Predictability Index	✓	✓	✓	✓	✓
	7	Headline Paper	✓	✓	✓	✓	✓
Stage 3	8	Error Decomposition	✗	✗	✓	✓	✓

Stage 0: Normative Ethics/Regulatory Policy

1. Rawlsian Fairness Principles:
 - Apply a maximin strategy ensuring the least advantaged benefit.
2. Kantian Categorical Imperatives:
 - Formulate universal policies: individuals as ends, not means.
3. Virtue Ethics Criteria:
 - Define measurable virtues via multi-criteria decision analysis.
4. Utilitarian Cost-Benefit Analysis:
 - Optimize net utility by balancing total benefits against costs.

Stage 0: Normative Ethics/Regulatory Policy

5. Procedural Justice Integration::

- Employ statistical tests to verify impartial, consistent procedures.

6. Normative Accountability Measures:

- Implement audit trails and metrics to certify ethical compliance.

7. Standardized Ethical Impact Assessments:

- Develop indices to quantitatively evaluate ethical implications.

8. Adaptive Bayesian Updating:

- Incorporate iterative feedback: refine normative ethical models.

Stage 1: Social Data as Text and Images

- Advance CNNs and LLMs for prediction in social contexts.
- Pick six intersecting variables from:

- Panels

- Registers

- Censuses

and:

- Socio-economics

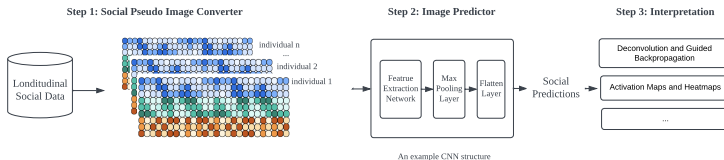
- Lifecourse

- Demographics

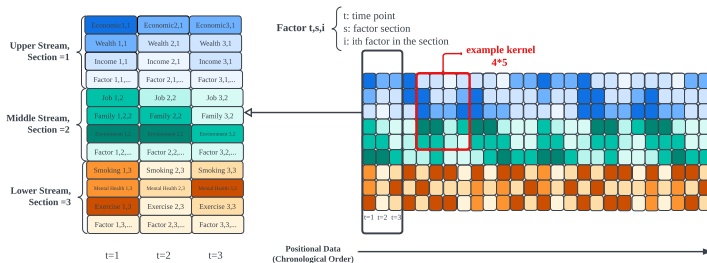
- Pre-register a 'domain expert' and SuperLearner baseline.
- International/multi-dataset context allows tests of external validity.
- Response to Salganik et al. (2024): 'breaking free from the X matrix'.
- Deep model fusion based models which learn from multiple datasets.

A simple stacked variable representation

A. The schematic workflow



B. An example of converted social aimage



Extension to Recurrence Matrices and GAF

- For each person $i \in \{1, \dots, N\}$ and each variable $j \in 1, \dots, d\}$, define the individual recurrence plot (with threshold function θ) as:

$$R_{t,t'}^{(i,j)} = \theta\left(|x_{i,t}^{(j)} - x_{i,t'}^{(j)}|\right), \quad t, t' \in \{1, \dots, T\}. \quad (1)$$

- Multidimensional Recurrence Plots (MRP) for person i are defined as:

$$MRP_{t,t'}^{(i)} = \prod_{j=1}^d R_{t,t'}^{(i,j)}. \quad (2)$$

- This aggregates the recurrence information across all d variables.
- We will also explore representations as Gramian Angular Fields.

In terms of social data as text:

- For prediction tasks, the key is to design a transformation:

$$g : \mathbb{R}^{n \times d} \rightarrow \mathcal{L}, \quad (3)$$

that converts tabular data $X \in \mathbb{R}^{n \times d}$ into a textual representation $\ell = g(X)$.

- Three ideas for how to enhance this.

1. Predictive Narrative Encoding. A high correlation between x_i and x_j (or theoretical relationship) that is indicative of an outcome might be encoded as:

‘The simultaneous increase in x_i and x_j robustly signals a higher likelihood of y ’

- 1.1 A further extension of this would be using an alternate (or same?) LLM to serialize the text.
2. Learn g by optimizing a differentiable surrogate for loss. Train g_θ jointly with a predictor h_ϕ by minimizing:

$$\min_{\theta, \phi} \mathbb{E} \left[L \left(h_\phi \left(g_\theta(X) \right), y \right) \right]. \quad (4)$$

Automatically discover textual constructs that best encapsulate the predictive signals in X .

Deep Model Fusion for Social Prediction

3. Assume an image-like representation:

$$\phi_{\text{img}}(X_i) \in \mathbb{R}^{H \times W \times C}, \quad (5)$$

a text embedding:

$$\phi_{\text{text}}(s_i) \in \mathbb{R}^{d_t}. \quad (6)$$

and incorporate a traditional superlearner on tabular data X for person i :

$$g_{SL}(X_i) \in \mathbb{R}^{d_{SL}}, \quad (7)$$

where g_{SL} is an ensemble of classical models.

- Define a fusion module f to obtain the joint representation:

$$F_i = f\left(\text{CNN}(\phi_{\text{img}}(X_i)), \phi_{\text{text}}(s_i), g_{SL}(X_i)\right) \in \mathbb{R}^{d_f}. \quad (8)$$

- Typical fusion strategies include:

- Concatenation:** Combine features from all modalities.
 - Cross-modal Attention:** Interactive weights between modalities.

Stage 2: Indexing Social Predictability

- New algorithms for pre-processing of social data at scale.
 - Different representations amenable to new algorithms as developed in Stage 1.
- Apply a range of algorithms to a vast set of normalized data.
- Health and genetic dimension added from UKB+CKB.
 - Allows consideration of GxE.
- Add in behavioral dimension from Time-Use-Surveys.
 - This allows better testing of TTE.

Universal Multi-Task Predictive Modeling

- Each variable X_i in a high-dimensional dataset $\mathbf{X} = \{X_1, X_2, \dots, X_p\}$ is both a predictor/outcome.
- Learn a collection of functions

$$f_i : \mathcal{X}_{-i} \rightarrow \mathbb{R}, \quad i = 1, \dots, p,$$

where $\mathcal{X}_{-i}$ denotes all variables except X_i , so that the prediction error

$$\mathbb{E} [L(f_i(\mathbf{X}_{-i}), X_i)]$$

is minimized for an appropriate loss function L .

- **Multi-Task Learning/Masked Feature Modeling** possible strategies.
- Another new idea: integrate a **DAG** as a regularizer a scale.

Predictability Index

[About](#)[Privacy Policy](#)

Selector

Domain

Health outcomes

Dataset

UKB

Target Variable

BP

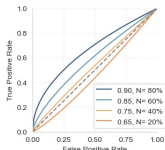
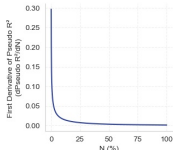
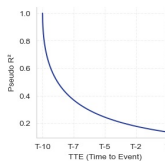
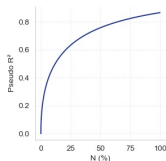
Time to Event

T-1

Introduction

Lorem ipsum dolor sit amet, consectetur adipiscing elit. Sed eu rutrum ipsum. Nulla a tortor sagittis, lacinia ipsum nec, pretium leo. Nulla blandit molestie accumsan. In hac habitasse platea dictumst. Maecenas quis metus pretium orci scelerisque rhoncus quis eu leo. Proin bibendum venenatis est, in tincidunt lectus laoreet semper. Quisque pellentesque vulputate eros vitae euismod.

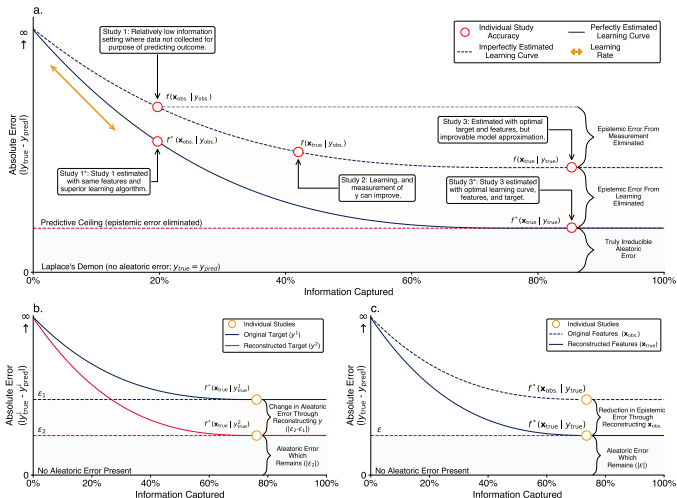
Model Performance

Ranks for selected domain

TTE	Data	Var	Pseudo R^2
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	UKB	BP	0.50
T-0	CKB	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	0.33

LEVERHULME
TRUSTOXFORD
POPULATION
HEALTHUK
RIEconomic
and Social
Research CouncilUNIVERSITY OF
OXFORD

Stage 3: Decomposing Epistemic Error



1. Learning Curve Analysis: Incremental Information Refinement:

$$L(I) = \mathbb{E} \left[|y_{\text{true}} - f_I(\mathbf{x})| \right], \quad (9)$$

$$\Delta_\epsilon \approx L(I_{\text{observed}}) - L(I_{\text{enhanced}}). \quad (10)$$

2. Approximate Δ_ϵ from perturbations /w **first-order Taylor expansion**:

$$\Delta_\epsilon \approx \nabla_{\mathbf{x}} f(\mathbf{x}, y) \cdot \delta \mathbf{x} + \nabla_y f(\mathbf{x}, y) \cdot \delta y, \quad (11)$$

where:

- $\Delta_\epsilon = f(\mathbf{x} + \delta \mathbf{x}, y + \delta y) - f(\mathbf{x}, y)$ is the change in the predictive error due to small changes in the inputs.
- $\nabla_{\mathbf{x}} f(\mathbf{x}, y)$ and $\nabla_y f(\mathbf{x}, y)$ are gradients (Jacobians) of f w.r.t. $\mathbf{x}$ and y .
- δx (δy): small perturbation/measurement error in $\mathbf{x}$ (y).

3. Epistemic error can be partially captured by the uncertainty in the model parameters **via bootstrap**:

$$\text{Var} (f(x_{\text{observed}})) \approx \frac{1}{M} \sum_{i=1}^M \left(f^{(i)}(x_{\text{observed}}) - \bar{f}(x_{\text{observed}}) \right)^2, \quad (12)$$

where $\bar{f}(x_{\text{observed}}) = \frac{1}{M} \sum_{i=1}^M f^{(i)}(x_{\text{observed}})$ is the mean prediction.

- This might be capturing Aleatoric Error, though, and not actually identifying epistemic error (further thinking needed).

4. **Latent Variable Modeling** can be used to “invert” the measurement error process/estimate error reduction through improved measurement:

$$x_{\text{observed}} = x_{\text{true}} + \delta x, \quad (13)$$

$$y_{\text{observed}} = y_{\text{true}} + \delta y, \quad (14)$$

Estimate the true feature set through EM algorithm or variational inference:

$$\Delta_{\text{measurement}} \approx \|f(x_{\text{observed}}) - f(\hat{x}_{\text{true}})\|. \quad (15)$$

- This is one way to consider the effect of construct validity (of X).
- However, it assumes knowledge of the DGP.

5. **Controlled Experiments/Simulations** with Synthetic Data To assess the impact of measurement errors, we simulate synthetic data using a known true model:

$$y_{\text{true}} = f^*(x_{\text{true}}) + \epsilon, \quad (16)$$

where ϵ represents inherent aleatoric noise. We introduce measurement errors as:

$$x_{\text{obs}} = x_{\text{true}} + \delta x, \quad y_{\text{obs}} = y_{\text{true}} + \delta y, \quad (17)$$

where δx and δy are perturbations capturing imperfections in data acquisition (measure or construct).

6. Also: concentration inequalities (e.g., Hoeffding's or McDiarmid's inequality) to establish probabilistic bounds.
7. Also: differential geometric methods like the Fisher information matrix.

Thanks for your attention!

Things we'd specifically like to discuss:.

1. Is this too risky for the ESRC?
 - Are there ethical concerns?
2. What can we do that goes even further beyond:
 - 2.1 BOLT?
 - 2.2 TabLLM?
3. Compliance related constraints to merging datasets on TREs?
4. Is this too much/not enough social science?