

On the Responsible use of Pseudo-Random Number Generators

Charlie Rahal

LCDS and Nuffield College, University of Oxford

Max Planck Institute for Demographic Research, November 26th 2025



What is a Pseudo Random Number Generator?

"A pseudo-random number is a value produced by a deterministic algorithm (a PRNG) that is intended to mimic the statistical properties of true randomness. Given an internal state (often initialized from a seed), a PRNG computes a sequence of numbers that appear unpredictable for many applications but are fully reproducible if the algorithm and state are known. PRNGs are essential for simulations, sampling, any many other types of algorithms."

- GPT-5 mini

PRNGs are just one type of RNG!

	'RN'	'RN*' and 'Hardware'	'RN*' and 'Quantum'	'RN*' and 'Pseudo'	'RN*' and 'Quasi'
Health Sciences	0.406	0.000646	0.000407	0.00109	0.004201
Life Sciences	0.292	0.000516	0.001336	0.00179	0.001202
Physical Sciences	0.225	0.005430	0.007717	0.00726	0.003884
Social Sciences	0.092	0.000216	0.000242	0.00055	0.002420

A cursory scientometric analysis of Randon Numbers (RN*). Data comes from OpenAlex API. All numbers are % of the entire scientific record.

- More 'PRNG papers' in CS (4222) than others combined.
- %PRNG papers since 1970: 0.0008% → 0.0048%.

PRNGs occur everywhere.

Computational Sciences:

- Numerical Integration
- Cryptography
- Genetic Algorithms
- Signal Processing
- Deep Learning
- Weather Forecasting
- Procedural Content
- ...

‘Applied’ (Health/Social) Sciences:

- RCTs/Survey Sampling
- Bootstrapping
- Epidemiological Modeling
- Econometric Estimation
- Data Imputation
- Topic Modeling
- ABMs
- ...

MT64: ubiquitous PRNG implementation (R, Stata, Python, ...)

prgn_stream_demo.py

```
1 import random
2
3 floats = [random.random() for _ in range(5)]
4 print(floats)
```

[0.36381, 0.20691, 0.50947, 0.33315, 0.87402]

prgn_stream_demo.py

```
1 import random
2
3 floats = [random.random() for _ in range(5)]
4 print(floats)
```

[0.4215, 0.01561, 0.9984, 0.7674, 0.274]

prgn_stream_demo.py

```
1 import random
2
3 random.seed(123456789) # deterministic
4
5 floats = [random.random() for _ in range(5)]
6 print(floats)
```

[0.64140, 0.542189, 0.99317, 0.84325, 0.81173]

So, why do we care about PRNGs?

An invisible source of uncertainty in the scientific record: PRNGs!

So, why do we care about PRNGs?

An invisible source of uncertainty in the scientific record: PRNGs!

- **Q:** Has anyone ran a program twice, with different results?

So, why do we care about PRNGs?

An invisible source of uncertainty in the scientific record: PRNGs!

- **Q:** Has anyone ran a program twice, with different results?
- **Q:** Has anyone here ever set a 'seed'? Which seed?

So, why do we care about PRNGs?

An invisible source of uncertainty in the scientific record: PRNGs!

- **Q:** Has anyone ran a program twice, with different results?
- **Q:** Has anyone here ever set a 'seed'? Which seed?
 - Maybe you prefer `set.seed(42)`?

So, why do we care about PRNGs?

An invisible source of uncertainty in the scientific record: PRNGs!

- **Q:** Has anyone ran a program twice, with different results?
- **Q:** Has anyone here ever set a 'seed'? Which seed?
 - Maybe you prefer `set.seed(42)`?
 - Or `numpy.random.seed(123)`?

So, why do we care about PRNGs?

An invisible source of uncertainty in the scientific record: PRNGs!

- **Q:** Has anyone ran a program twice, with different results?
- **Q:** Has anyone here ever set a 'seed'? Which seed?
 - Maybe you prefer `set.seed(42)`?
 - Or `numpy.random.seed(123)`?
- Why did you do this?

So, why do we care about PRNGs?

An invisible source of uncertainty in the scientific record: PRNGs!

- **Q:** Has anyone ran a program twice, with different results?
- **Q:** Has anyone here ever set a 'seed'? Which seed?
 - Maybe you prefer `set.seed(42)`?
 - Or `numpy.random.seed(123)`?
- Why did you do this?
 - Maybe to eliminate variation in algorithms with PRNGs?

So, why do we care about PRNGs?

An invisible source of uncertainty in the scientific record: PRNGs!

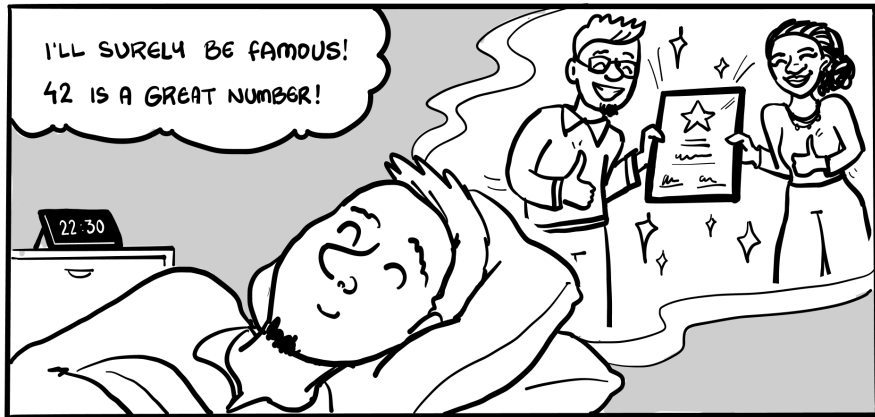
- **Q:** Has anyone ran a program twice, with different results?
- **Q:** Has anyone here ever set a 'seed'? Which seed?
 - Maybe you prefer `set.seed(42)`?
 - Or `numpy.random.seed(123)`?
- Why did you do this?
 - Maybe to eliminate variation in algorithms with PRNGs?
 - This seems to be the current 'best practice' advice.

So, why do we care about PRNGs?

An invisible source of uncertainty in the scientific record: PRNGs!

- **Q:** Has anyone ran a program twice, with different results?
- **Q:** Has anyone here ever set a 'seed'? Which seed?
 - Maybe you prefer `set.seed(42)`?
 - Or `numpy.random.seed(123)`?
- Why did you do this?
 - Maybe to eliminate variation in algorithms with PRNGs?
 - This seems to be the current 'best practice' advice.
- This is very well intentioned: it allows reproducibility. Yay!

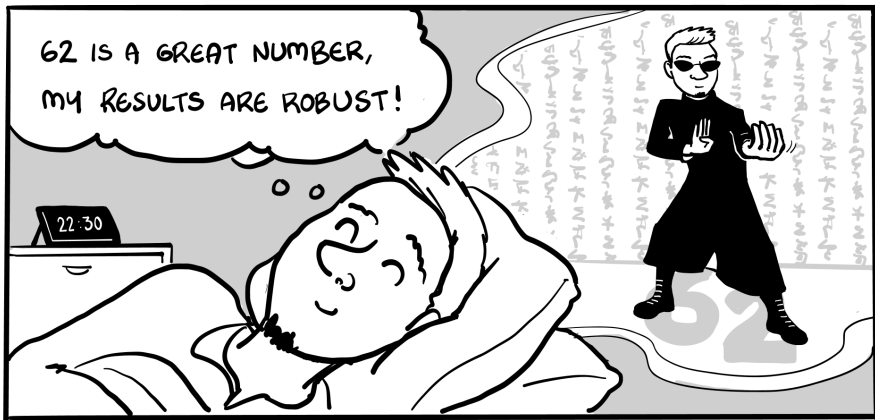


















THE SEED OF DOUBT: A RANDOM TALE OF REPRODUCIBILITY



*See 'Rahal, C., Yan, J. And Verhagen, M. (2025), 'On the Responsible Use of Pseudo Random Number Generators', arXiv, pp. 1–26'

DRALON BH LADRA SORVALA

“If any research design is susceptible to this kind of variation, there is a problem”.

– Anonymous Colleague, Oxford (2019).

The General Premise

- **Problem Statement:** By setting **one** seed in applied algorithmic pipelines, we ignore the variation of our estimand as a function of how PRNGs were instantiated, This is computationally un-intensive, but scientifically dangerous.

The General Premise

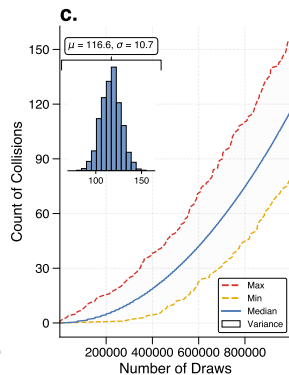
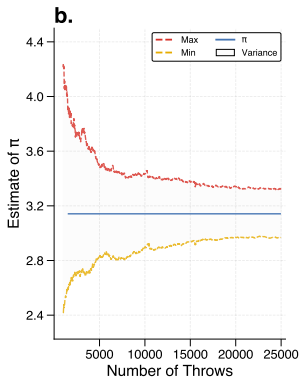
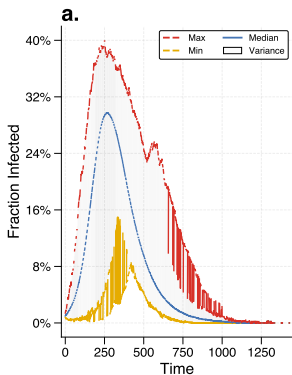
- **Problem Statement:** By setting **one** seed in applied algorithmic pipelines, we ignore the variation of our estimand as a function of how PRNGs were instantiated, This is computationally un-intensive, but scientifically dangerous.
- **Solution:** Visualize the outcome space of a **large number** (10k? 100k?) of seeds simultaneously. This is computationally intensive, but scientifically responsible.

The General Premise

- **Problem Statement:** By setting **one** seed in applied algorithmic pipelines, we ignore the variation of our estimand as a function of how PRNGs were instantiated, This is computationally un-intensive, but scientifically dangerous.
- **Solution:** Visualize the outcome space of a **large number** (10k? 100k?) of seeds simultaneously. This is computationally intensive, but scientifically responsible.
- **Replication:** Similar to specification curves and p-hacking, we can retrospectively consider whether an original result is in the tail/IQR of the distribution of possible outcome space.

What do we do?

- Through **many** replications and simulations (20?), we:
 - Show some simple simulated examples.
 - Show seed variability in data preprocessing
 - Show seed variability in modelling
 - Show seed variability in prediction
 - Show the potential compounding of these effects
 - Show implications for fairness
 - Show how this affects LLMs, too.
- Then provide guidance to hopefully improve scientific practice.
- All code and data is open source and available on GitHub.



What's going on here?

- These examples sample from `random` or `np.random()` directly.
 - No model fitting, no data processing, just raw sampling.
 1. Simple draws from a SIR model.
 2. A replication of Buffon's infamous needle problem.
 3. Replicating some recent work on random number collisions.

What's going on here?

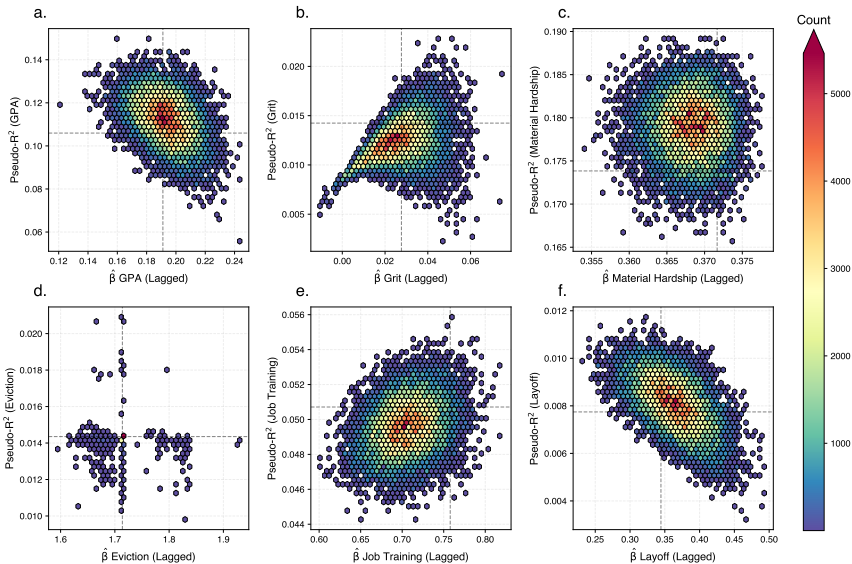
- These examples sample from `random` or `np.random()` directly.
 - No model fitting, no data processing, just raw sampling.
 1. Simple draws from a SIR model.
 2. A replication of Buffon's infamous needle problem.
 3. Replicating some recent work on random number collisions.
- Even in these simple cases, seed variability is **non-trivial**.

What's going on here?

- These examples sample from `random` or `np.random()` directly.
 - No model fitting, no data processing, just raw sampling.
 1. Simple draws from a SIR model.
 2. A replication of Buffon's infamous needle problem.
 3. Replicating some recent work on random number collisions.
- Even in these simple cases, seed variability is **non-trivial**.
- What are the implications for more complex pipelines?

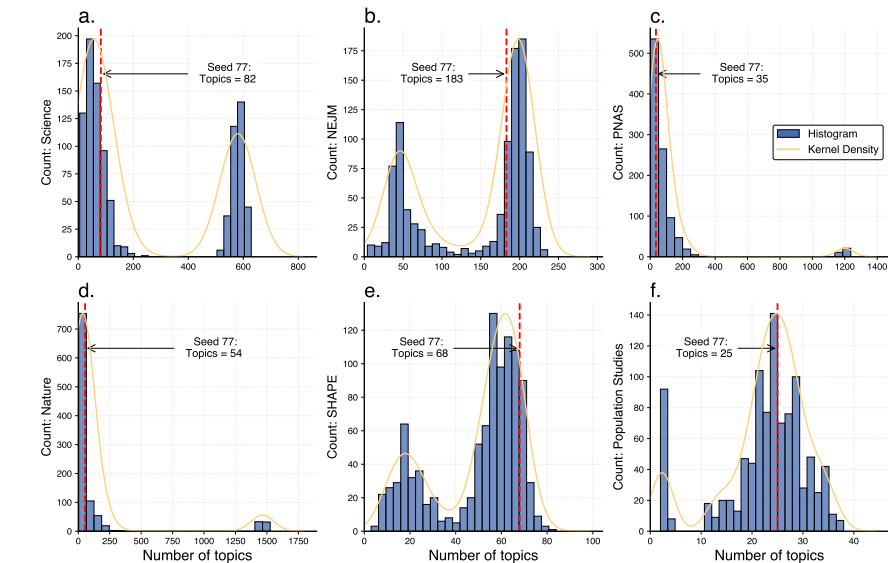
More Complex Replications

- Lets move into more complex and direct replications.
 - Major focus on the quantitative/computational social sciences.
 - Machine/Deep Learning
 - Sociology
 - Digital Humanities
 - Criminology
 - Medicine
 - Computer Vision
 - Econometrics
 - Time Series Analysis
 - Urban Segregation
 - LLMs
- Both recent and classical papers.
 - For each, we'll continue to give intuition as to what's going on.



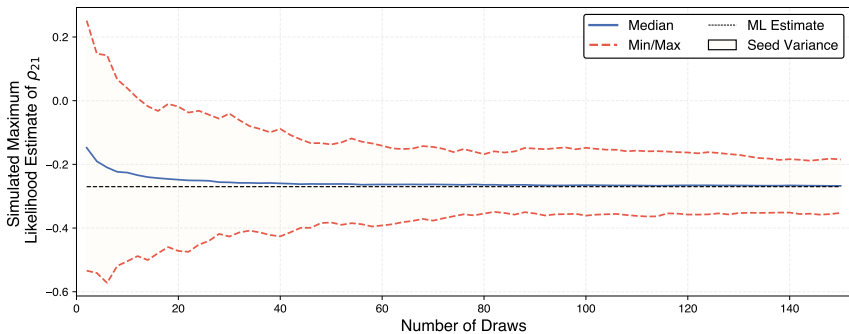
What's going on here?

- Amelia fills in missing data by drawing values from distributions.
- These distributions are based on observed data patterns.
- The imputed values reflect the uncertainty in the missing data.
- This is a form of Multiple Imputation.
- This results in entirely different models being built.
- This has implications for both inference and prediction.



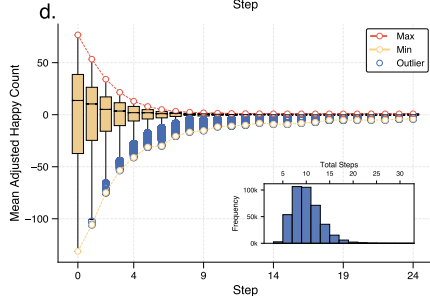
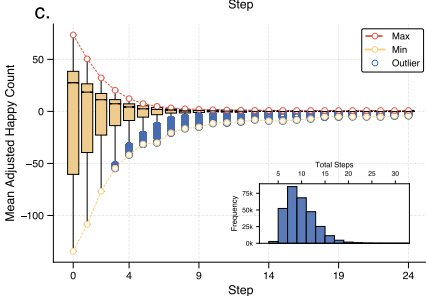
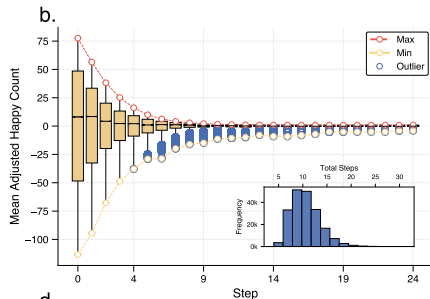
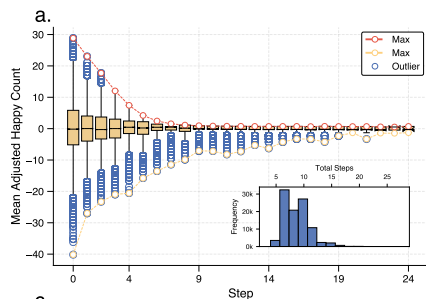
What's going on here?

- BERTopic has many process: one particularly volatile to seeds.
- It commonly involves dimension reduction via UMAP.
 - **UMAP**: causes different low-dimensional embeddings.
- Note: CUMML implementation of UMAP gives variable outputs for spectral initialisation *even when setting the seed*.
- If using HDBScan for clustering, no seed variability.
- In general here, the seed variance looks **enormous**.



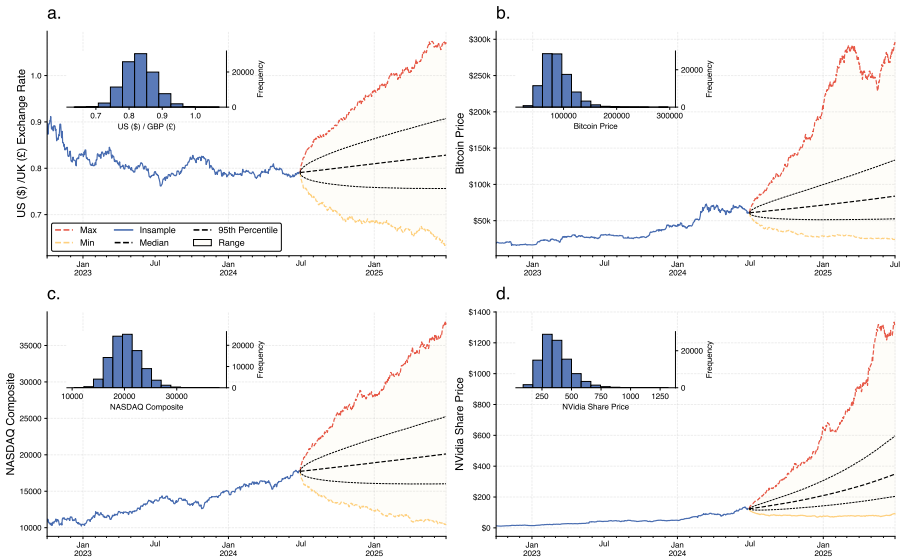
What's going on here?

- The `mvprobit` estimator requires evaluating M integrals jointly.
- This can be rewritten into a multivariate normal distributions.
 - Then, a sequence of normal distributions with certain properties.
- `mvprobit` samples from this to get at the underlying likelihood.
- Rewrite each individual normal to its CDF, which allows you to sample by just entering a uniform value between $[0,1]$.
 - This random value is seed-dependent.
- Where CDF is steep on a small interval, seeds matter a lot.



What's going on here?

- There are multiple (!) sources of PRNG in the Schelling Model:
 1. The initialisation of the agents on the grid.
 2. The shuffling of unhappy agents and empty cells.
- Agents might start more/less segregated from the beginning.
 - This affects the speed of total segregation over time.
- Even if the model reaches similar segregation, configuration of agents and trajectories will differ.
- As shown in the previous slide, parameterisation matters (a lot).



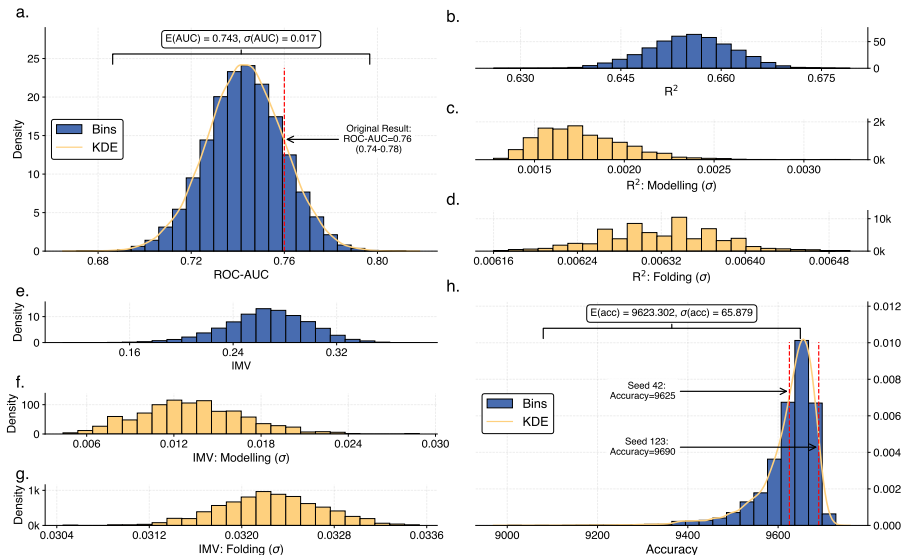
What's going on here?

- Seeds modify sequences of PRNG draws, affecting y_t .
- Moments of previous changes $(\mu, \sigma/2)$ control magnitude.
- Draws of $\mathcal{N}(\mu, \sigma/2)$ impact asset price trajectory.
- The seed affects long-term trends by influencing both direction and size of early random steps.

This can be formally written as:

$$y_t = y_{t-1} + \epsilon_t \quad \text{where} \quad \epsilon_t \sim \mathcal{N}(\mu, \sigma/2) \quad (1)$$

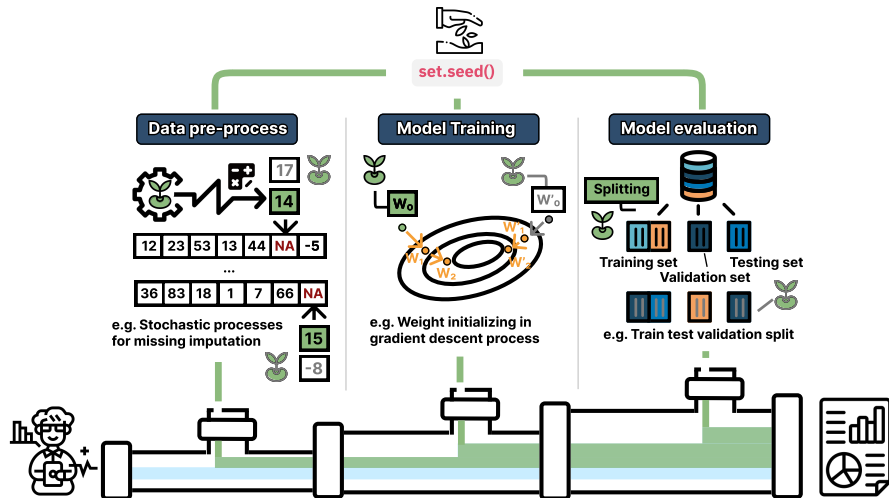
- Previously thought that nothing outperforms RW for ForEx.
- RWs commonly used as a benchmark for structural models.

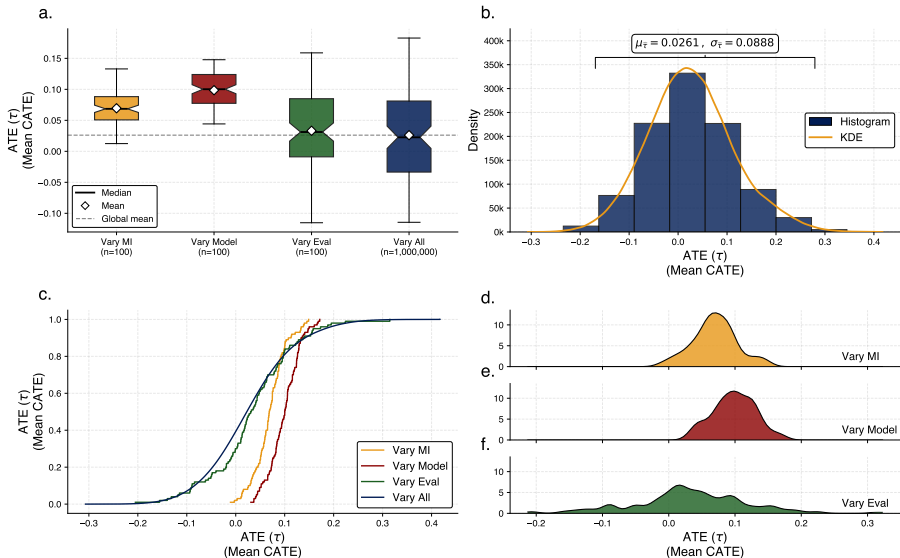


What's going on here?

In addition to modelling, ML/DL pipelines involve stochastic folding.

- This happens with every form of CV.
- The only exception is LOOCV.
- The fewer the splits, the more the seed matters.
- We are able to decompose types of seed effect.
 - This is done with Titanic and housing examples.
- The final example shows no folding CV, only model variability.



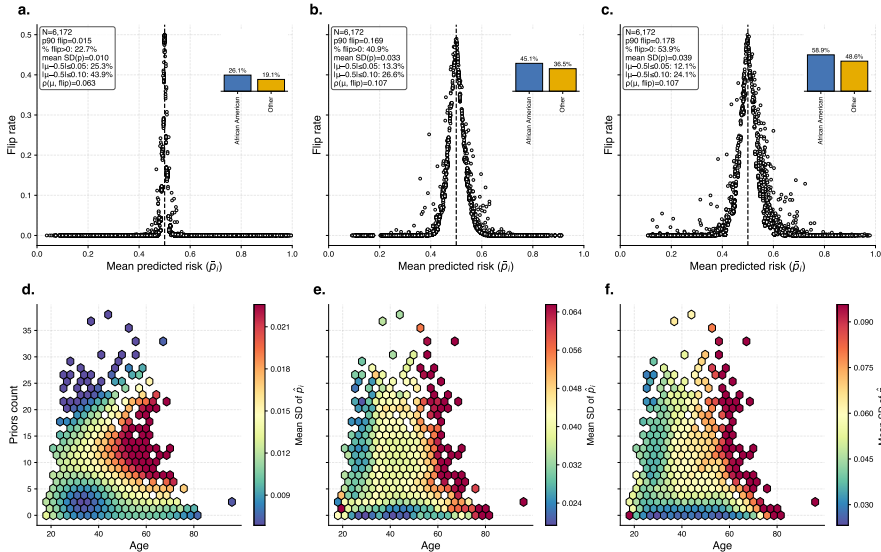


Compound Simulation: Core Idea

- Assess how randomness from imputation, model fitting, and evaluation affects downstream ATE estimates.
- Importantly, it shows how seed variability compounds.
- Pipeline components:
 - Simulate one dataset with MAR data.
 - IterativeImputer (MICE-style stochastic draws),
 - CausalForestDML (cross-fitted).
- Compare five seed-varying schemes: `baseline_fixed`, `vary_imp`, `vary_model`, `vary_eval`, `vary_all`.

Simulation Parameters

- $N = 600$
- $d = 16$
- $\tau = 0.10$
- $\text{miss_rate} = 0.50$
- $R_{\text{imp}} = 100$
- $R_{\text{model}} = 100$
- $R_{\text{eval}} = 100$
- $M_{\text{imputations}} = 1$
- $\text{n_estimators} = 200$
- $\text{subforest_size} = 4$
- $\text{cv_folds} = 3$
- $\text{n_jobs} = -1$



COMPAS: seed-driven OOF instability

Emphasizes the implications for distributional fairness.

- Seed variability: because of model instantiation and folding.
- Uses the infamous Pro-Publica dataset on recidivism.
- Repeated 2-fold OOF across many seeds for RF, LR, and MLP.
- Summarise per-UID instability (variance / flip-rate).

COMPAS: seed-driven OOF instability

Emphasizes the implications for distributional fairness.

- Seed variability: because of model instantiation and folding.
- Uses the infamous Pro-Publica dataset on recidivism.
- Repeated 2-fold OOF across many seeds for RF, LR, and MLP.
- Summarise per-UID instability (variance / flip-rate).

Other obvious extensions:

COMPAS: seed-driven OOF instability

Emphasizes the implications for distributional fairness.

- Seed variability: because of model instantiation and folding.
- Uses the infamous Pro-Publica dataset on recidivism.
- Repeated 2-fold OOF across many seeds for RF, LR, and MLP.
- Summarise per-UID instability (variance / flip-rate).

Other obvious extensions:

1. Different domains (mortgage finance, health)

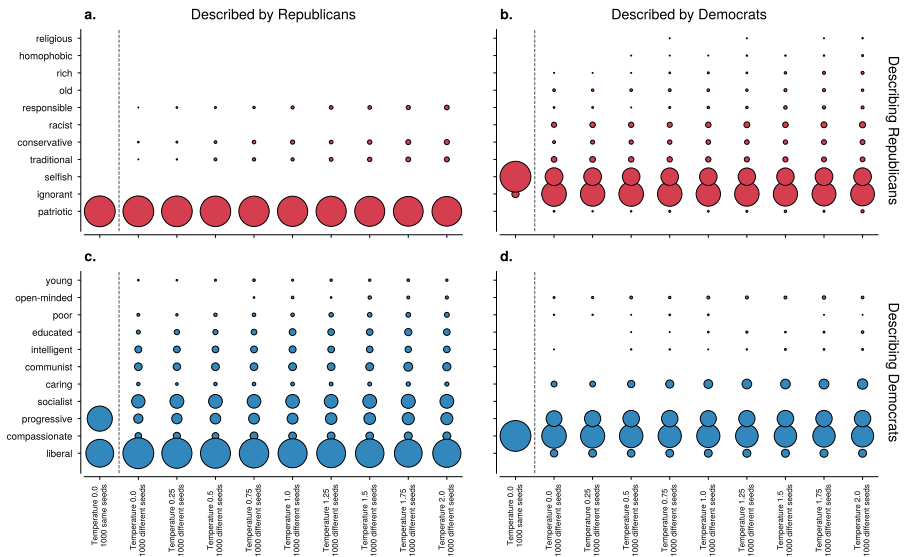
COMPAS: seed-driven OOF instability

Emphasizes the implications for distributional fairness.

- Seed variability: because of model instantiation and folding.
- Uses the infamous Pro-Publica dataset on recidivism.
- Repeated 2-fold OOF across many seeds for RF, LR, and MLP.
- Summarise per-UID instability (variance / flip-rate).

Other obvious extensions:

1. Different domains (mortgage finance, health)
2. Alternate metrics, e.g., statistical parity.



Seed variability in the experiment

- Outer loop varies temperature across nine fixed values.
- Inner loop iterates over up to 1000 seeds.
- Outputs depend jointly on temperature and seed.
- Fixed-seed experiment repeats 1000 calls at temperature zero.
- Temperature has little effect on variability compared to seed.
- The seed regulates the sampling step – the choice among logits at each token.
 - But the logits themselves come from a huge distributed system whose compute is not bit-deterministic.
- Even with *fixed seed* (and $\text{temp}=0$): non-determinism abounds!

Concluding Thoughts (Part One)

- Matt Salganik commented: 'That sounds great, but expensive!'
 - Expense no-longer prohibitive: everything here ran locally.
 - We live in an age of vast computational resources.

Concluding Thoughts (Part One)

- Matt Salganik commented: 'That sounds great, but expensive!'
 - Expense no-longer prohibitive: everything here ran locally.
 - We live in an age of vast computational resources.
- What are our suggestions?

Concluding Thoughts (Part One)

- Matt Salganik commented: 'That sounds great, but expensive!'
 - Expense no-longer prohibitive: everything here ran locally.
 - We live in an age of vast computational resources.
- What are our suggestions?
 - **Researchers:** Visualize your seed variability where affordable!
 - There are times when you can eliminate/reduce seed variability.
 - Pedagogically, lets teach students to parralelize code effectively.

Concluding Thoughts (Part One)

- Matt Salganik commented: 'That sounds great, but expensive!'
 - Expense no-longer prohibitive: everything here ran locally.
 - We live in an age of vast computational resources.
- What are our suggestions?
 - **Researchers:** Visualize your seed variability where affordable!
 - There are times when you can eliminate/reduce seed variability.
 - Pedagogically, lets teach students to parralelize code effectively.
 - **Reviewers:** Question papers you read – is there stochasticity?

Concluding Thoughts (Part One)

- Matt Salganik commented: 'That sounds great, but expensive!'
 - Expense no-longer prohibitive: everything here ran locally.
 - We live in an age of vast computational resources.
- What are our suggestions?
 - **Researchers:** Visualize your seed variability where affordable!
 - There are times when you can eliminate/reduce seed variability.
 - Pedagogically, lets teach students to parralelize code effectively.
 - **Reviewers:** Question papers you read – is there stochasticity?
 - Even well established results are susceptible.

Concluding Thoughts (Part One)

- Matt Salganik commented: 'That sounds great, but expensive!'
 - Expense no-longer prohibitive: everything here ran locally.
 - We live in an age of vast computational resources.
- What are our suggestions?
 - **Researchers:** Visualize your seed variability where affordable!
 - There are times when you can eliminate/reduce seed variability.
 - Pedagogically, lets teach students to parallelize code effectively.
 - **Reviewers:** Question papers you read – is there stochasticity?
 - Even well established results are susceptible.
 - **Editors:** Ubiquitous mandate for responsible use of PRNGs!

simple_ranges.py

```
1 import numpy as np # randomisation
2 import matplotlib.pyplot as plt # plotting
3
4 def analytical_function(func_input, seed):
5     '''Simple analytical function: can be anything'''
6     np.random.seed(seed)
7     return func_input*np.random.normal()
8
9 results = [] # store results
10 inputs = 42 # dataset, figure path, etc.
11
12 for seed in range(0, 100000): # simple loop; can be distributed
13     results.append(analytical_function(inputs, seed))
14
15 plt.hist(results) # plot results
```

better_seed_algorithm.py

```
1 import numpy as np # randomisation
2 import matplotlib.pyplot as plt # plotting
3
4 def analytical_function(func_input, seed):
5     '''Simple analytical function: can be anything'''
6     np.random.seed(seed)
7     return func_input*np.random.normal()
8
9 results = [] # store results
10 inputs = 42 # dataset, figure path, etc.
11
12 # Use predefined 'secret' seeds we provide
13 with open(os.path.join '..', 'assets', 'seed_list.txt')) as f:
14     seed_list = [int(line.rstrip('\n')) for line in f]
15
16 for seed in seed_list: # simple loop; can be distributed
17     results.append(analytical_function(inputs, seed))
18
19 plt.hist(results) # plot results
```

Concluding Thoughts (Part Two)

- Ideally, we want to move towards distributional reproducibility.
- Moving forward, PRNGs shouldn't exist (QRNGs the norm).
- Does anyone have ideas for other types of seed variability?
- Can this be corrective? Index historical seed variability?
- Prospectively: we make available a list of seeds (replication materials), encouraging their use as a pre-specified set.
- **TLDR**: when variation can't be eliminated, should be visualised.

Thanks to my Coauthors!



Mark Verhagen



Jiani Yan

Thanks to my Coauthors!



Mark Verhagen



Jiani Yan

And thanks to you for your attention and attendance!

Replication Materials



Naturally, containerised code available on GitHub at the QR code.