

Soc306: Machine Learning with Social Data

Week 12: Interpretability

Brandon Stewart¹

Princeton

April 14–16, 2025

¹Huge thanks to Zenobia Chan for permission to include many of her slides here.

Where We've Been and Where We're Going...

Where We've Been and Where We're Going...

- Last Week
 - ▶ fairness and racial justice

Where We've Been and Where We're Going...

- Last Week
 - ▶ fairness and racial justice
- This Week
 - ▶ interpretability
 - ▶ computing topic of choice!

Where We've Been and Where We're Going...

- Last Week
 - ▶ fairness and racial justice
- This Week
 - ▶ interpretability
 - ▶ computing topic of choice!
- Next Week
 - ▶ machine learning in practice and policy

Where We've Been and Where We're Going...

- Last Week
 - ▶ fairness and racial justice
- This Week
 - ▶ interpretability
 - ▶ computing topic of choice!
- Next Week
 - ▶ machine learning in practice and policy
- Long Run
 - ▶ target tasks → data sources → guiding values

- 1 Interpretability
- 2 Interpretability vs. Explainability
- 3 Interpretable Algorithmic Forensics
- 4 Open vs. Closed LLMs
- 5 Preview: (Ir)reproducible machine learning

- 1 Interpretability
- 2 Interpretability vs. Explainability
- 3 Interpretable Algorithmic Forensics
- 4 Open vs. Closed LLMs
- 5 Preview: (Ir)reproducible machine learning

Three Broad Categories of Values

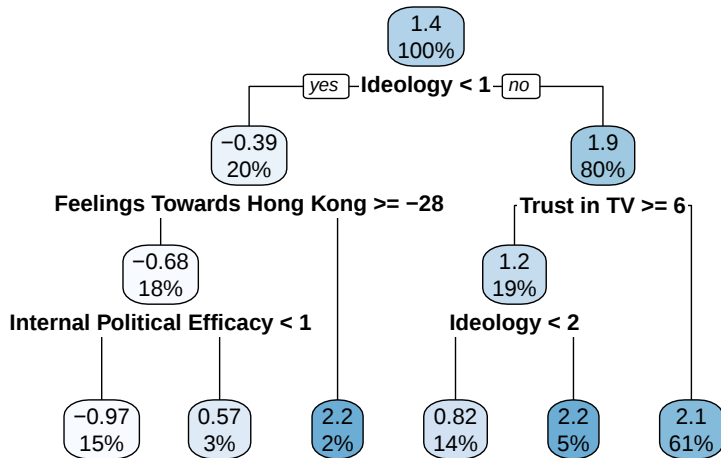
- **Privacy**
what information is exposed?
- **Fairness**
are gains and costs distributed equitably?
- **Interpretability**
do we know what the machine algorithm is doing?

What variables predict the ideology of Estonians' vote choice?



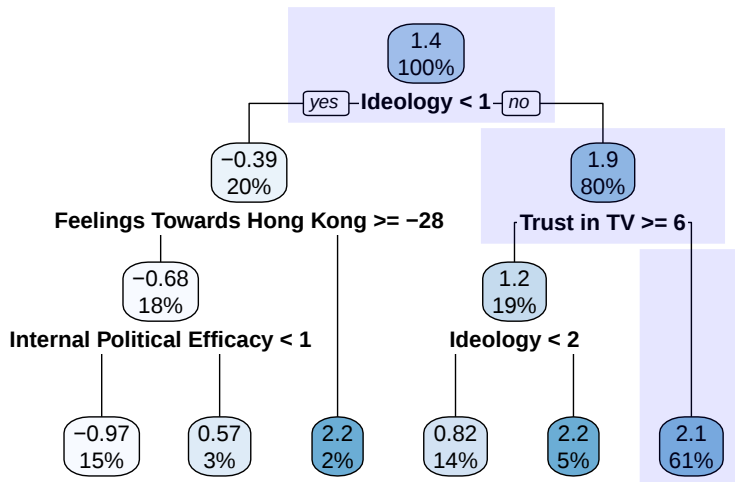
Slide Credit: Zenobia Chan

CART (Classification And Regression Trees)



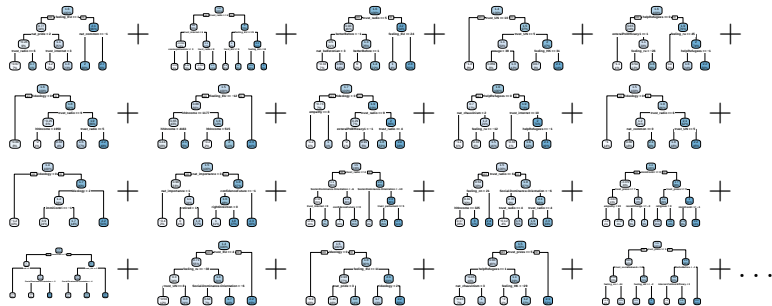
Slide Credit: Zenobia Chan

CART (Classification And Regression Trees)



Slide Credit: Zenobia Chan

From Trees to Forests



Random Forest = Trees + Bagging + Randomizing Variables

Slide Credit: Zenobia Chan

From Trees to Boosted Trees

Initialize $\hat{Y}^0 = \bar{Y}$; λ a small number

$$\hat{Y}^1 = \hat{Y}^0 + \lambda \times$$


$$\hat{Y}^2 = \hat{Y}^1 + \lambda \times$$


$$\hat{Y}^3 = \hat{Y}^2 + \lambda \times$$


$$\hat{Y}^4 = \hat{Y}^3 + \lambda \times$$


$$\hat{Y}^5 = \hat{Y}^4 + \lambda \times$$


$$\hat{Y}^6 = \hat{Y}^5 + \lambda \times$$


$$\hat{Y}^7 = \hat{Y}^6 + \lambda \times$$


$$\hat{Y}^8 = \hat{Y}^7 + \lambda \times$$


$$\hat{Y}^9 = \hat{Y}^8 + \lambda \times$$

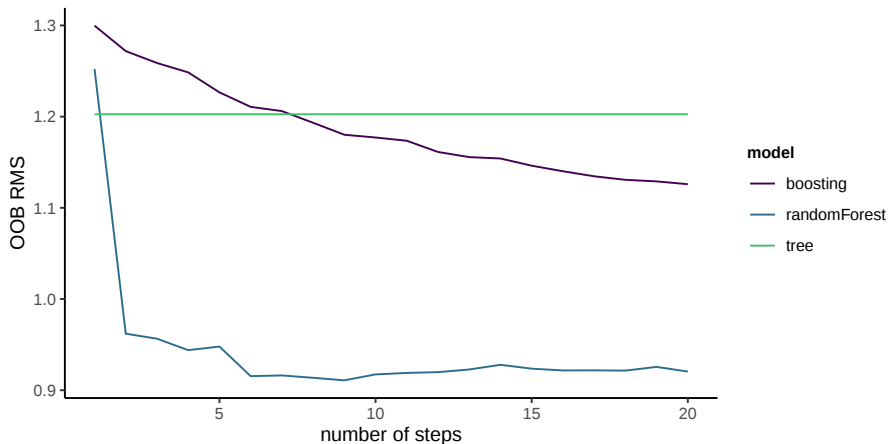

...

Slide Credit: Zenobia Chan

MSE Goes Down...

Performance Comparison

Dataset: Foster and Chan 2019; Outcome: Ideology of Vote Choice



...but so does interpretability!

Slide Credit: Zenobia Chan

Why do we need interpretability?

Why do we need interpretability?

① Legal requirements



Why do we need interpretability?

- 1 Legal requirements

Why do we need interpretability?

- ① Legal requirements
- ② Build trust for human decision makers

Why do we need interpretability?

- ① Legal requirements
- ② Build trust for human decision makers
- ③ Anticipate behavior in new situations

Why do we need interpretability?

- ① Legal requirements
- ② Build trust for human decision makers
- ③ Anticipate behavior in new situations
- ④ Help to monitor for other value problems: safety, ethics, privacy etc.

Why do we need interpretability?

- 1 Legal requirements
- 2 Build trust for human decision makers
- 3 Anticipate behavior in new situations
- 4 Help to monitor for other value problems: safety, ethics, privacy etc.

⇒ Incompleteness in the problem formalization

Interpretability can assist in qualitatively ascertaining other important desiderata of ML systems

Fairness, privacy, reliability, robustness, causality, usability, trust...

Interpretability can assist in qualitatively ascertaining other important desiderata of ML systems

Fairness, privacy, reliability, robustness, causality, usability, trust...

Use cases

- 1 Scientific understanding

Interpretability can assist in qualitatively ascertaining other important desiderata of ML systems

Fairness, privacy, reliability, robustness, causality, usability, trust...

Use cases

- 1 Scientific understanding
- 2 Safety

Interpretability can assist in qualitatively ascertaining other important desiderata of ML systems

Fairness, privacy, reliability, robustness, causality, usability, trust...

Use cases

- 1 Scientific understanding
- 2 Safety
- 3 Ethics

Interpretability can assist in qualitatively ascertaining other important desiderata of ML systems

Fairness, privacy, reliability, robustness, causality, usability, trust...

Use cases

- 1 Scientific understanding
- 2 Safety
- 3 Ethics
- 4 Mismatched objectives

Interpretability can assist in qualitatively ascertaining other important desiderata of ML systems

Fairness, privacy, reliability, robustness, causality, usability, trust...

Use cases

- 1 Scientific understanding
- 2 Safety
- 3 Ethics
- 4 Mismatched objectives
- 5 Multi-objective trade-offs

Slide Credit: Zenobia Chan

How should we evaluate interpretability?

(Doshi-Velez and Kim 2017)

		Evaluation	End Goal (s)	Real Human Involved	Task	
more	specific and costly	Application-grounded	Whether it results in better identification of errors, new facts, or less discrimination	Yes	Yes	Real
		Human-grounded	Whether humans and models yield similar results	No	Yes	Simple
		Functionality-grounded	Whether the models are known to be interpretable	No	No	Proxy
less						

Slide Credit: Zenobia Chan

How should we evaluate interpretability?

(Doshi-Velez and Kim 2017)

		Evaluation	End Goal (s)	Real Human Involved	Task	
more	specific and costly	Application-grounded	Whether it results in better identification of errors, new facts, or less discrimination	Yes	Yes	Real
		Human-grounded	Whether humans and models yield similar results	No	Yes	Simple
		Functionality-grounded	Whether the models are known to be interpretable	No	No	Proxy
less						

Slide Credit: Zenobia Chan

How should we evaluate interpretability?

(Doshi-Velez and Kim 2017)

		Evaluation	End Goal (s)	Real Human Involved	Task	
more	specific and costly	Application-grounded	Whether it results in better identification of errors, new facts, or less discrimination	Yes	Yes	Real
		Human-grounded	Whether humans and models yield similar results	No	Yes	Simple
		Functionality-grounded	Whether the models are known to be interpretable	No	No	Proxy
less						

Slide Credit: Zenobia Chan

How should we evaluate interpretability?

(Doshi-Velez and Kim 2017)

Evaluation		End Goal (s)	Real Human Involved	Task	
more specific and costly less	Application-grounded	Whether it results in better identification of errors, new facts, or less discrimination	Yes	Yes	Real
	Human-grounded	Whether humans and models yield similar results	No	Yes	Simple
	Functionality-grounded	Whether the models are known to be interpretable	No	No	Proxy

Slide Credit: Zenobia Chan

Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead

Cynthia Rudin 

Black box machine learning models are currently being used for high-stakes decision making throughout society, causing problems in healthcare, criminal justice and other domains. Some people hope that creating methods for explaining these black box models will alleviate some of the problems, but trying to explain black box models, rather than creating models that are interpretable in the first place, is likely to perpetuate bad practice and can potentially cause great harm to society. The way forward is to design models that are inherently interpretable. This Perspective clarifies the chasm between explaining black boxes and using inherently interpretable models, outlines several key reasons why explainable black boxes should be avoided in high-stakes decisions, identifies challenges to interpretable machine learning, and provides several example applications where interpretable models could potentially replace black box models in criminal justice, healthcare and computer vision.

Definitions

Black-box model: a function that is too complicated for human comprehension; and/or proprietary

Definitions

Black-box model: a function that is too complicated for human comprehension; and/or proprietary

Interpretability

- Domain-specific notion $\implies$ No all-purpose definition
- Common characteristics of interpretable models:
 - 1 Constrained in model form s.t. it is useful
 - 2 Obeys structural knowledge of the domain

Definitions

Black-box model: a function that is too complicated for human comprehension; and/or proprietary

Interpretability

- Domain-specific notion $\implies$ No all-purpose definition
- Common characteristics of interpretable models:
 - 1 Constrained in model form s.t. it is useful
 - 2 Obeys structural knowledge of the domain

Explainability

- Understanding how a model works
- NOT understanding how the world works

Slide Credit: Zenobia Chan

Key Issues with Explainable ML

Key Issues with Explainable ML

- 1 There's no necessary trade-off between accuracy

Key Issues with Explainable ML

- ① There's no necessary trade-off between accuracy
- ② Explanations from explainable ML models are not faithful to what the original model computes

Key Issues with Explainable ML

- 1 There's no necessary trade-off between accuracy
- 2 Explanations from explainable ML models are not faithful to what the original model computes
- 3 Explanations are sometimes not sensible and lack details for understanding the black box

Key Issues with Explainable ML

- 1 There's no necessary trade-off between accuracy
- 2 Explanations from explainable ML models are not faithful to what the original model computes
- 3 Explanations are sometimes not sensible and lack details for understanding the black box
- 4 Black-box models make it difficult to combine with information additional to what is in the dataset

Key Issues with Explainable ML

- 1 There's no necessary trade-off between accuracy
- 2 Explanations from explainable ML models are not faithful to what the original model computes
- 3 Explanations are sometimes not sensible and lack details for understanding the black box
- 4 Black-box models make it difficult to combine with information additional to what is in the dataset
- 5 Black-box models are prone to errors in training data

Slide Credit: Zenobia Chan



Fig. 2 | Saliency does not explain anything except where the network is looking. We have no idea why this image is labelled as either a dog or a musical instrument when considering only saliency. The explanations look essentially the same for both classes. Credit: Chaolei Chen, Duke University

Key Issues with Interpretable ML

- 1 Black-box models allow firms to profit

Key Issues with Interpretable ML

- ① Black-box models allow firms to profit
- ② Interpretable models create opportunities for gaming the system

Key Issues with Interpretable ML

- 1 Black-box models allow firms to profit
- 2 Interpretable models create opportunities for gaming the system
- 3 Interpretable models can be difficult and costly to create
 - ▶ High-stake decisions:
Costs of analysts + computation < Costs of wrong decision

Key Issues with Interpretable ML

- 1 Black-box models allow firms to profit
- 2 Interpretable models create opportunities for gaming the system
- 3 Interpretable models can be difficult and costly to create
 - ▶ High-stake decisions:
Costs of analysts + computation < Costs of wrong decision
- 4 Black-box models seem to uncover “hidden patterns”

Slide Credit: Zenobia Chan

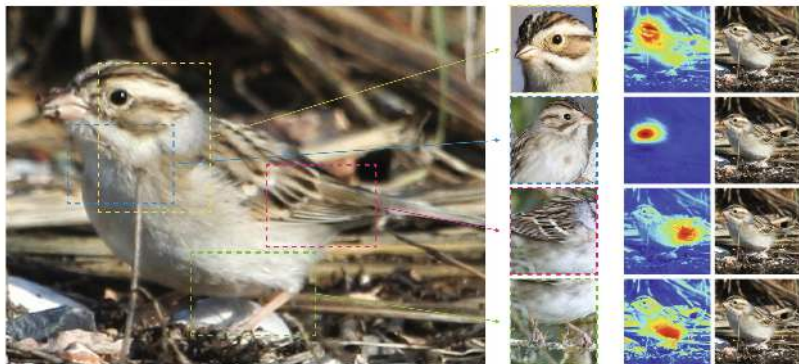


Fig. 3 | Image from the authors of ref. ¹⁶, indicating that parts of the test image on the left are similar to prototypical parts of training examples. The test image to be classified is on the left, the most similar prototypes are in the middle column, and the heatmaps that show which part of the test image is similar to the prototype are on the right. We included copies of the test image on the right so that it is easier to see to what part of the bird the heatmaps are referring. The similarities of the prototypes to the test image are what determine the predicted class label of the image. Here, the image is predicted to be a clay-coloured sparrow. The top prototype seems to be comparing the bird's head to a prototypical head of a clay-coloured sparrow, the second prototype considers the throat of the bird, the third looks at feathers, and the last seems to consider the abdomen and leg. Credit: Image constructed by Alina Barnett, Duke University

TECHNOLOGY

We Need Transparency in Algorithms, But Too Much Can Backfire

by [Kartik Hosanagar](#) and [Vivian Jair](#)

July 23, 2018 **UPDATED** July 25, 2018



A student revolt...



Too much transparency can indeed backfire

Enter Kizilcec (2016)

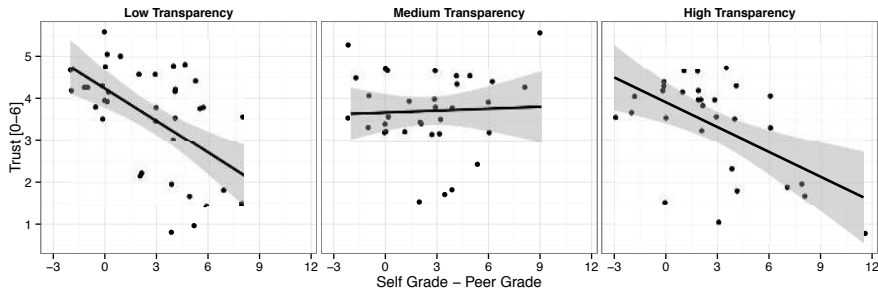


Figure 2. Trust examined as a function of the difference in expected (self) grade and received (peer) grade in each randomly assigned transparency condition. Points are jittered for presentation; OLS regression lines with 95% confidence bounds.

Introduction to Interpretable AI in Forensics

- Topic: Importance of interpretable AI in criminal justice.

Introduction to Interpretable AI in Forensics

- Topic: Importance of interpretable AI in criminal justice.
- Context: Increasing use of black box AI which lacks transparency.

Introduction to Interpretable AI in Forensics

- Topic: Importance of interpretable AI in criminal justice.
- Context: Increasing use of black box AI which lacks transparency.
- Issue: Such approaches can obscure forensic evidence and violate constitutional rights.

Problems with Black Box AI

- Definition: Black box AI are systems whose workings are not understandable.

Problems with Black Box AI

- Definition: Black box AI are systems whose workings are not understandable.
- Examples: DNA mixture interpretation, facial recognition, risk assessments.

Problems with Black Box AI

- Definition: Black box AI are systems whose workings are not understandable.
- Examples: DNA mixture interpretation, facial recognition, risk assessments.
- Consequences: Risks to public safety and legal rights due to non-transparency.

Advantages of Glass Box AI

- Definition: Glass box AI is designed to be transparent and interpretable.

Advantages of Glass Box AI

- Definition: Glass box AI is designed to be transparent and interpretable.
- Research Findings: Shows higher accuracy than black box AI.

Advantages of Glass Box AI

- Definition: Glass box AI is designed to be transparent and interpretable.
- Research Findings: Shows higher accuracy than black box AI.
- Implications: Supports due process and safeguards defendant's rights.

Key Premise

The Black Box Performance Myth

LMSYS Chatbot Arena Leaderboard

[Home](#)
[Blog](#)
[About](#)
[FAQ](#)
[Privacy](#)
[Terms](#)
[Contact](#)

LMSYS [Chatbot Arena](#) is a crowdsourced open platform for LLM evals. We've collected over 500,000 human pairwise comparisons to rank LLMs with the [SeedEval](#) model and display the model ratings in Elo score. You can find more details in our [paper](#).

Arena
 [Full Leaderboard](#)

Total #models: 82. Total votes: 672,236. Last updated: April 13, 2024.

NEW! View leaderboards for different categories (e.g., coding, long-term query).

Code to regenerate leaderboard tables and plots in this [notebook](#). You can contribute your vote @ [chat.lmsys.org](#)!

Category

Overall

Overall Questions

#models: 82 (100%) #votes: 672,236 (100%)

Rank	Model	Arena Elo	95% CI	Votes	Organization	License	Knowledge Cutoff
1	GPT-4 Turbo 2024-04-09	1260	157-5	15751	OpenAI	Proprietary	2023/12
2	Claude 3.5 Sonnet	1255	+37/-4	36131	Anthropic	Proprietary	2023/9
3	GPT-4o-mini	1214	+37/-3	60199	OpenAI	Proprietary	2023/4
4	GPT-4o	1200	+37/-4	59923	OpenAI	Proprietary	2023/12
5	Bard (Gemini Pro)	1189	157/-5	12468	Google	Proprietary	On-line
6	Claude 3 Sonnet	1163	137/-3	62356	Anthropic	Proprietary	2023/8
7	Gemini Pro	1153	+47/-4	59637	Google	CC-BY-NC-SA 4.0	2022/3
8	GPT-4o-mini	1109	+47/-4	42925	OpenAI	Proprietary	2021/9
9	Claude 3 Haiku	1182	+37/-3	37727	Anthropic	Proprietary	2023/8
10	GPT-4o	1164	137/-3	61528	OpenAI	Proprietary	2021/9
11	Mistral Large 2/32	1158	137/-4	37658	Mistral	Proprietary	Unknown

Legislative and Judicial Response

- Current Trends: Absence of required transparency in AI use in courts.

Legislative and Judicial Response

- Current Trends: Absence of required transparency in AI use in courts.
- Proposed Changes: Legislation to mandate interpretable AI in investigations.

Legislative and Judicial Response

- Current Trends: Absence of required transparency in AI use in courts.
- Proposed Changes: Legislation to mandate interpretable AI in investigations.
- Judicial Rulings: Enforcing rights to interpretable forensic AI.

Case Studies and Legal Implications

- Case Studies: Instances where black box AI was legally challenged.

Case Studies and Legal Implications

- Case Studies: Instances where black box AI was legally challenged.
- Legal Considerations: Need for scrutinizable AI systems.

Case Studies and Legal Implications

- Case Studies: Instances where black box AI was legally challenged.
- Legal Considerations: Need for scrutinizable AI systems.
- Future Outlook: Potential for action to promote interpretable AI in forensics.

Conclusion

- Summary: Importance of shifting from black box to glass box AI in justice.

Conclusion

- Summary: Importance of shifting from black box to glass box AI in justice.
- Call to Action: Advocacy for policies that prioritize transparency.

Conclusion

- Summary: Importance of shifting from black box to glass box AI in justice.
- Call to Action: Advocacy for policies that prioritize transparency.
- Goal: Ensuring AI aligns with constitutional protections and public safety.

Abstract

The use of machine learning (ML) methods for prediction and forecasting has become widespread across many scientific sciences. However, there are many known methodological pitfalls, including data leakage, in ML-based science. In this paper, we systematically investigate reproducibility issues in ML-based science. We show that data leakage is indeed a widespread problem and has led to severe reproducibility failures. Specifically, through a survey of literature in research communities that adopted ML methods, we find 17 fields where errors have been found, collectively affecting 339 papers and in some cases leading to wildly overoptimistic conclusions. Based on our survey, we present a fine-grained taxonomy of 8 types of leakage that range from textbook errors to open research problems.

We argue for fundamental methodological changes in ML-based science so that cases of leakage can be caught before publication. To that end, we propose model info sheets for reporting scientific claims based on ML models that would address all types of leakage identified in our survey. To investigate the impact of reproducibility errors and the efficacy of model info sheets, we undertake a reproducibility study in a field where complex ML models are believed to vastly outperform older statistical models such as Logistic Regression (LR): civil war prediction. We find that all papers claiming the superior performance of complex ML models compared to LR models fail to reproduce due to data leakage, and complex ML models don't perform substantially better than decades-old LR models. While none of these errors could have been caught by reading the papers, model info sheets would enable the detection of leakage in each case.

1. Overview

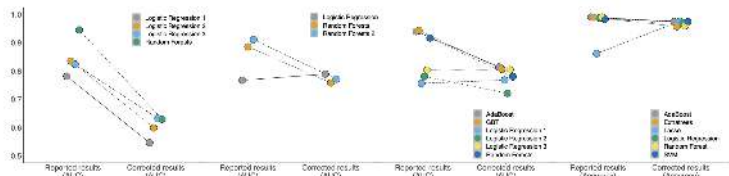
There has been a marked shift towards the paradigm of predictive modeling across quantitative science fields. This shift has been facilitated by the widespread use of machine learning (ML) methods. However, pitfalls in using ML methods have led to exaggerated claims about their performance. Such errors can lead to a feedback loop of over-optimism about the paradigm of prediction—especially as non-replicable publications tend to be cited more often than replicable ones (Serra-García & Górriz, 2021). It is therefore important to examine the reproducibility of findings in communities adopting ML methods.

Scope. We focus on reproducibility issues in ML-based science, which involves making a scientific claim using the performance of the ML model as evidence. There is a much better known reproducibility crisis in research that uses traditional statistical methods (Open Science Collaboration, 2015). We also situate our work in contrast to other ML domains, such as methods research (creating and improving widely-applicable ML methods), ethics research (studying the ethical implications of ML methods), engineering applications (building or improving a product or service), and model criticism (improving predictive performance on a fixed dataset created by an independent third party). Investigating the validity of claims in all of these areas is important, and there is ongoing work to address reproducibility issues in these domains (Hollman et al., 2022; Pincois et al., 2020; Erik Cundesen, 2021; Dell & Kampman, 2021).

We define a research finding as reproducible if the code and data used to obtain the finding are available and the data is correctly analyzed (Hohmann et al., 2021a; Leek & Peng, 2015; Pincois et al., 2020). This is a broader definition than computational reproducibility—when the results in a paper can be replicated using the exact code and dataset provided by the authors (see Appendix A).

Leakage. Data leakage has long been recognized as a leading cause of errors in ML applications (Nabae et al., 2009). In formative work on leakage, Kaufman et al. (2013) provide an overview of different types of errors and give several recommendations for mitigating these errors. Since this paper was published, the ML community has investigated leakage in several engineering applications and modeling competitions (Traser, 2016; Ghani et al., 2020; Decker, 2018; Urmowicz, 2015; Collins-Thompson). However, leakage occurring in ML-based science has not been comprehensively

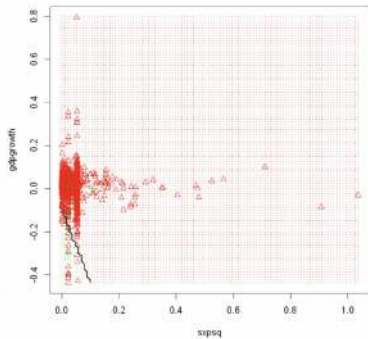
¹Department of Computer Science and Center for Information Technology Policy, Princeton University. Correspondence to: Sreyash.Kapoor@princeton.edu.



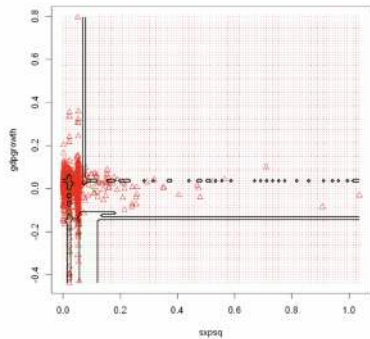
Paper	Muchlinski et al.	Colaresi and Mahmood	Wang	Kaufman et al.
Claim	Random Forests model drastically outperforms Logistic regression models	Random Forests models drastically outperform Logistic regression model	Adaboost and Gradient Boosted Trees (GBT) drastically outperform other models	Adaboost outperforms other models
Error	[L1.2] Pre-proc. on train-test (Incorrect imputation)	[L1.2] Pre-proc. on train-test (Incorrect reuse of an imputed dataset)	[L1.2] Pre-proc. on train-test. (Incorrect reuse of an imputed dataset) [L3.1] Temporal leakage (k-fold cross validation with temporal data)	[L2] Illegitimate features (Data leakage due to proxy variables) [L3.1] Temporal leakage (k-fold cross validation with temporal data)
Impact	Random Forests perform no better than Logistic Regression	Random Forests perform no better than Logistic Regression	Difference in AUC between Adaboost and Logistic Regression drops from 0.14 to 0.01	Adaboost no longer outperforms Logistic Regression. None of the models outperform a baseline model that predicts the outcome of the previous year
Discussion	Impact of the incorrect imputation is severe since 86% of the out-of-sample dataset is missing and is filled in using the incorrect imputation method	Re-use the dataset provided by Muchlinski et al., which uses an incorrect imputation method	Re-use the dataset provided by Muchlinski et al., which uses an incorrect imputation method	Use several proxy variables for the outcome as predictors (e.g., <i>colwars</i> , <i>covwars</i> , <i>swars</i> , all proxies for civil war), leading to near perfect accuracy

Figure 1. A comparison of reported and corrected results in civil war prediction papers published in top political science journals. The main findings of each of these papers are invalid due to various forms of data leakage: Muchlinski et al. (2016) impute the training and test data together, Colaresi & Mahmood (2017) and Wang (2019) incorrectly reuse an imputed dataset, and Kaufman et al. (2019) use proxies for the target variable which causes data leakage. The use of model info sheets (Section 3) would detect leakage in every paper. When we correct these errors, complex ML models (such as Adaboost and Random Forests) do not perform substantively better than decades-old Logistic Regression models for civil war prediction in each case. Each column in the table outlines the impact of leakage on the results of a paper. The figure above each column shows the difference in performance that results from fixing leakage issues.

Logistic Regression



RF



Field	Paper	Number of papers reviewed Number of papers with pitfalls [L1.1] No test set [L1.2] Pre-proc. on train-test [L1.3] Feature set. on train-test [L2.1] Duplicates [L2.2] Illegitimate features [L2.3] Temporal leakage [L3.1] Non-ind. b/w train-test [L3.2] Sampling bias Comput. reproducibility issues Data quality issues Metric choice issues Standard dataset used?										
Medicine	Bouwmeester et al. (2012)	71	27	o							o	
Neuroimaging	Whelan & Garavan (2014)	–	14	o		o						
Autism Diagnostics	Bone et al. (2015)	–	3				o					
Bioinformatics	Blagus & Lusa (2015)	–	6		o				o	o	o	
Nutrition Research	Ivanescu et al. (2016)	–	4	o								
Software Eng.	Tu et al. (2018)	58	11					o		o	o	
Toxicology	Alves et al. (2019)	–	1			o				o		
Satellite Imaging	Nalepa et al. (2019)	17	17					o		o	o	
Tractography	Poulin et al. (2019)	4	2	o						o	o	
Clinical Epidem.	Christodoulou et al. (2019)	71	48			o				o		
Brain-computer Int.	Nakanishi et al. (2020)	–	1	o							o	
Histopathology	Oner et al. (2020)	–	1					o				
Neuropsychiatry	Poldrack et al. (2020)	100	53	o	o					o	o	
Medicine	Vandewiele et al. (2021)	24	21					o	o	o	o	
Radiology	Roberts et al. (2021)	62	62	o		o			o	o	o	
IT Operations	Iyu et al. (2021)	9	3					o			o	
Medicine	Filho et al. (2021)	–	1				o					
Neuropsychiatry	Shim et al. (2021)	–	1							o		
Genomics	Barnett et al. (2022)	41	23		o					o		
Computer Security	Arp et al. (2022)	30	30	o	o	o	o	o	o	o	o	

Table 1. Survey of 20 papers that identify pitfalls in the adoption of ML methods across 17 fields, collectively affecting 329 papers. In each field, papers adopting ML methods suffer from data leakage. The column headings for types of data leakage, shown in bold, are based on our taxonomy of data leakage. We also highlight other issues that are reported in the papers, including issues with computational reproducibility (the availability of code, data, and computing environment to reproduce the exact results reported in the paper), data quality (for example, small size or large amounts of missing data), metric choice (using incorrect metrics for the task at hand, for example, using accuracy for measuring model performance in the presence of heavy class imbalance), and standard dataset use, where issues are found despite the use of standard datasets in a field.

Model Card for Avoiding Leakage in ML-Based Science

Paper Handle. Detail where the scientific paper that motivated the model:

- Authors and description of the model
- Purpose and use
- Data sources
- License
- Where to ask questions or make requests about the model

Scientific Claims. Details of the scientific claims made by the team or the manager of the model:

- Population or distribution that exhibits the relevant model
- Spatial and temporal characteristics of the claims

Model. List all models created and used in the paper, from best to worst:

- Performance metrics
- Training and testing
- Feature importance

Include model details not in the paper, from best to worst:

- Training and use data source/segment, and model for why the model and use data do not have the same distribution as the training data

Conclusion. Interpretation of the model, stating whether there are limitations in the data and the model are limited:

- Performance metrics for the best model, from best to worst
- Model training data
- Statistical significance
- Statistical significance

Model for application. How to use the model in the real world, including steps for deployment:

- Feature selection
- Feature scaling
- Model training
- Model validation

Validation. Details of the validation process, including the data used and the model used:

- Validation type: separate validation or bootstrap
- Test set size
- Test results
- Model used for validation
- Validation data source and model
- Validation results and model

Model for application. How to use the model in the real world, including steps for deployment:

- Feature selection
- Feature scaling
- Model training
- Model validation

Model for application. How to use the model in the real world, including steps for deployment:

- Feature selection
- Feature scaling
- Model training
- Model validation

Model for application. How to use the model in the real world, including steps for deployment:

- Feature selection
- Feature scaling
- Model training
- Model validation

Model for application. How to use the model in the real world, including steps for deployment:

- Feature selection
- Feature scaling
- Model training
- Model validation

Model for application. How to use the model in the real world, including steps for deployment:

- Feature selection
- Feature scaling
- Model training
- Model validation

Model for application. How to use the model in the real world, including steps for deployment:

- Feature selection
- Feature scaling
- Model training
- Model validation

Model for application. How to use the model in the real world, including steps for deployment:

- Feature selection
- Feature scaling
- Model training
- Model validation

Model Cards Template

1. Train and test sets never interact
2. Features used in the model are legitimate to answer the scientific question of interest
3. Test set is drawn from the distribution of scientific interest

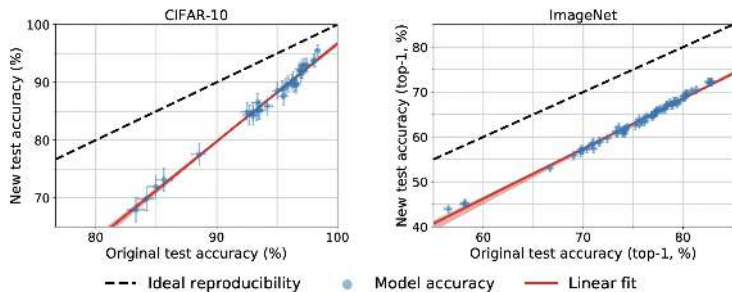


image from: "Do ImageNet Classifiers Generalize to ImageNet?", Recht et al., ICML 2019

This Week in Review

This Week in Review

- Interpretability

This Week in Review

- Interpretability

Next Week: Machine Learning in Practice and Policy

- Sayash Kapoor and Arvind Narayanan. 2023. “Leakage and the reproducibility crisis in machine-learning-based science” *Patterns*
- Johnson and Zhang. 2022. “What is the Bureaucratic Counterfactual? Categorical versus Algorithmic Prioritization in U.S. Social Policy” *FAccT '22*.
- Wang, Kapoor, Barocas, and Narayanan. 2024. “Against Predictive Optimization: On the Legitimacy of Decision-Making Algorithms that Optimize Predictive Accuracy” *ACM Journal of Responsible Computing*