



Navigating Data Errors in Machine Learning Pipelines: Identify, Debug, and Learn

Bojan Karlaš (Harvard University), Babak Salimi (UC San Diego), Sebastian Schelter (BIFOLD & TU Berlin)



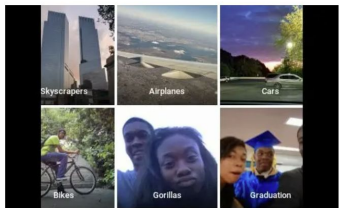
navigating-data-errors.github.io

Background: ML apps often behave in unintended ways

Wrong

Google apologises for Photos app's racist blunder

© 1 July 2015



diri noir avec banan
Google Photos, y'a
My friend's not a gorilla.
Mr Aline tweeted Google about the fact its app had misclassified his photo

Source: BBC

Biased

MIT
Technology
Review

Featured Topics Newsletters Events Audio

SIGN IN

SUBSCRIBE

SILICON VALLEY

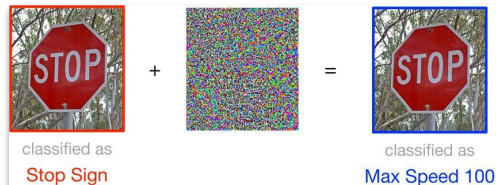
Amazon ditched AI recruitment software because it was biased against women

By Erin Winick

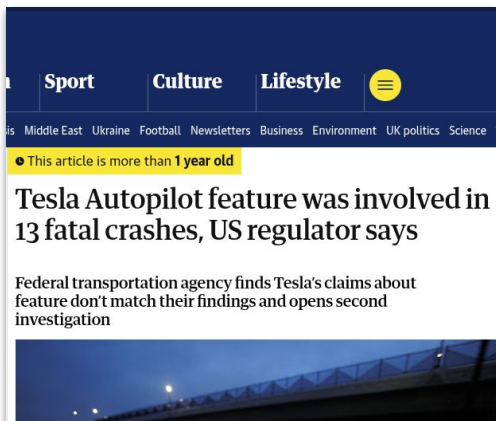
October 10, 2018

Source: MIT Technology Review

Unstable

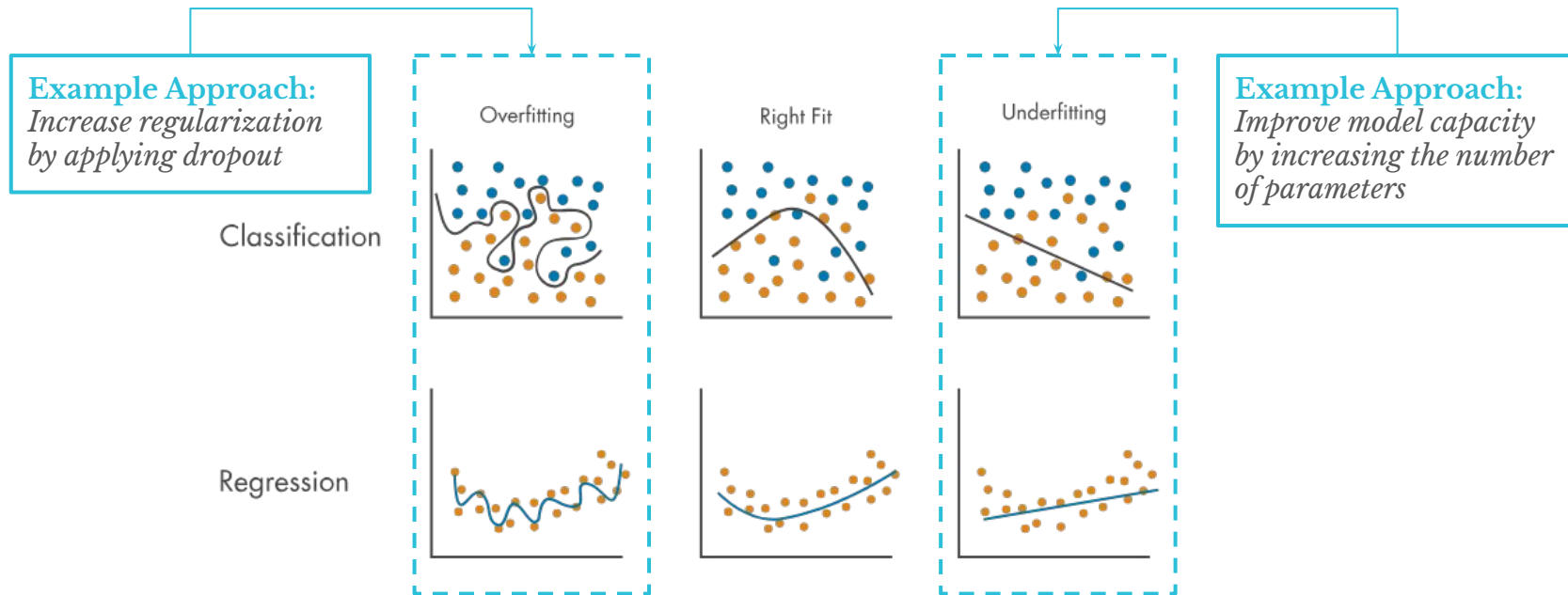


Source: Xiong et al. ACM Comput. Surv. 2023.



Source: The Guardian

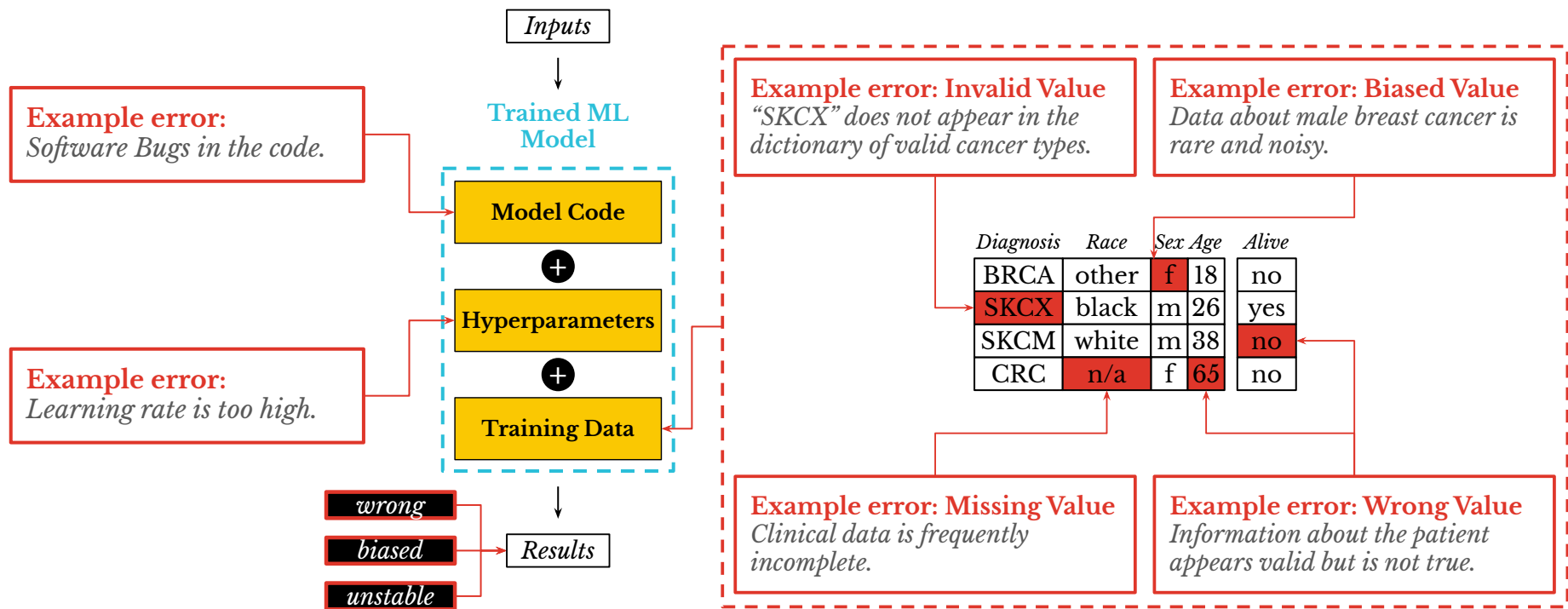
Primary approach: Focus on improving the model



Source: MathWorks

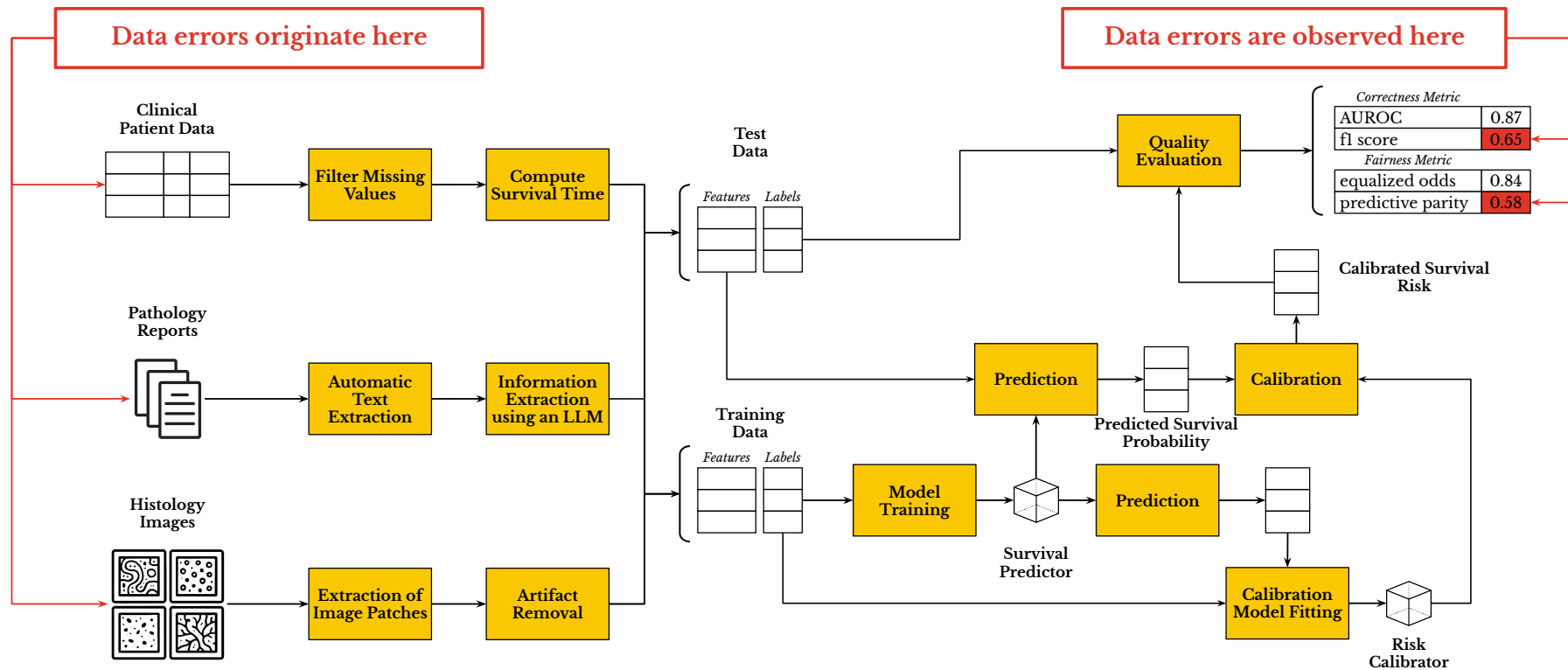
Problem: *This is only one piece of the puzzle!*

Observation 1: Data is a crucial piece of the puzzle



Challenge 1: Can we identify the most important data errors?

Observation 2: ML apps are built by complex pipelines



Challenge 2: *Can we trace data errors as they pass through the pipeline?*

Observation 3: Not all data errors are meant to be fixed

For each data error, we can choose to perform one of the following actions:

Discard



Remove the faulty data from the training set.

Repair



Perform manual quality control which might include repeating the data acquisition process.

Ignore



Let the faulty data remain in the training set.

Benefits:

Easy to Perform

Data Quality Improves

No Labor Required

Shortcomings:

Loss of Useful Data

Often Labor-intensive

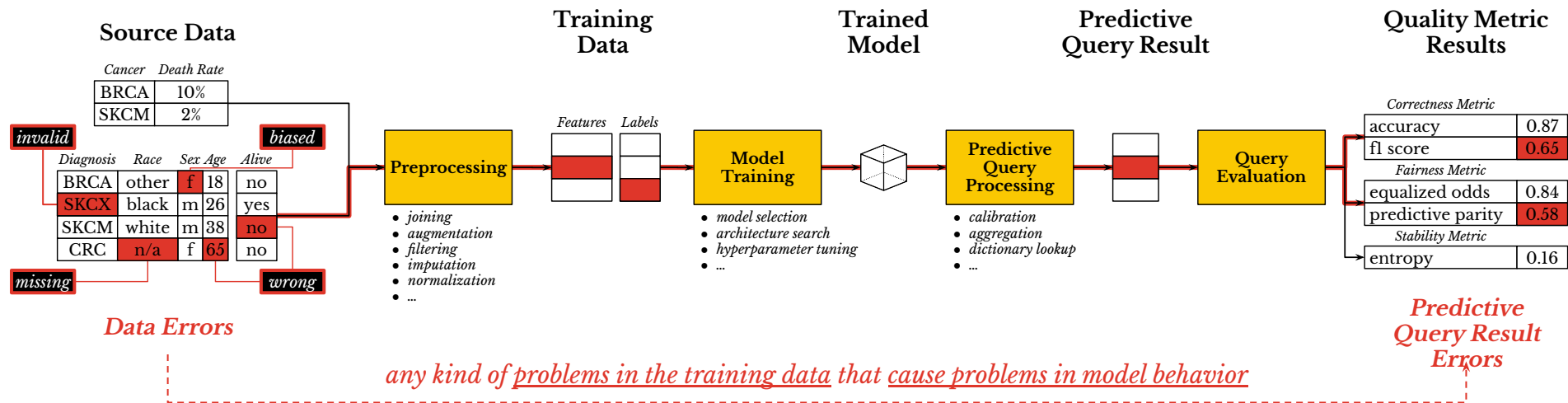
Risk Hurting Model Quality

Optimal trade-off:

Discard or Repair the Portion of Data that will Bring the Highest Model Quality Increase

Challenge 3: *Can we ensure reliable model performance after (partial) data repairs?*

Tutorial Overview: Data Errors in ML pipelines



Part I: Data Importance for Data Error Detection

What are good approaches for identifying data errors?

Part II: Data Debugging in ML Pipelines

What are practical challenges when debugging complex ML pipelines?

Part III: Learning from Uncertain and Incomplete Data

When we cannot repair all errors, can we still have reliable models?

Opportunities for the Data Management Community

- (1) Data quality is an established discipline in data management, but most practitioners still rely on **manual effort**.
- (2) ML pipelines are data processing pipelines. Models are learned data transformation operators. Many systems have been developed, but most practitioners still rely on **rudimentary scripts for crunching data**.
- (3) Many promising methods for handling data errors suffer from **scalability issues**.

Main Goal: *Present the current state of the art and inspire novel research.*

Part I: Data Importance for Data Error Detection

Bojan Karlaš



- 1) **Introducing the Concept of Data Importance**
- 2) **Examples of Data Attribution Functions**
- 3) **Case Study of Shapley Value as a Measure of Importance**
- 4) **Applications of Data Importance**

How can we identify data errors?

Trivial

Solution approach:

Apply a rule-based validation function that performs a dictionary lookup.

Solution approach:

Check if the value is marked as missing.

invalid

Diagnosis	Race	Sex	Age	Alive
BRCA	other	f	18	no
SKCX	black	m	26	yes
SKCM	white	m	38	no
CRC	n/a	f	65	no

missing

biased

wrong

Solution approach:

Measure the impact of the value on model quality.

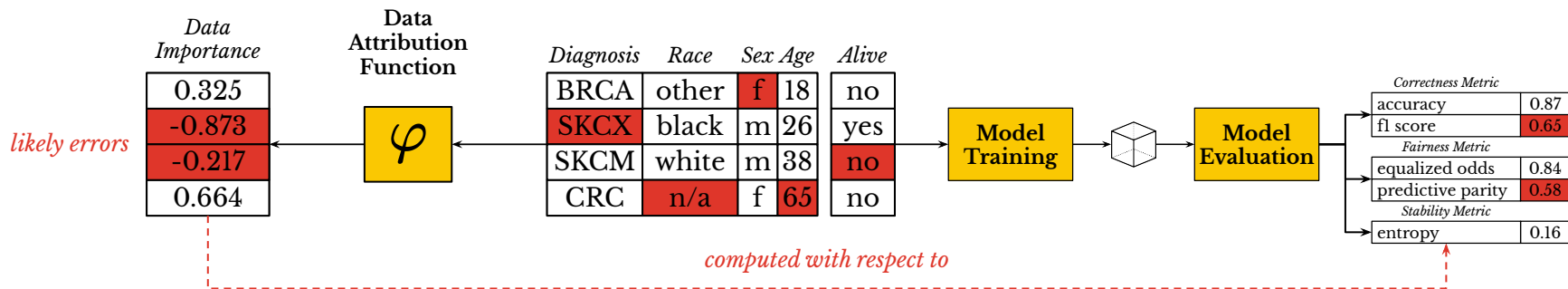
How do we measure this?

That is the main topic of this part of the tutorial.

Recall: Data errors are any kind of problem in the training data that cause problems in model behavior.

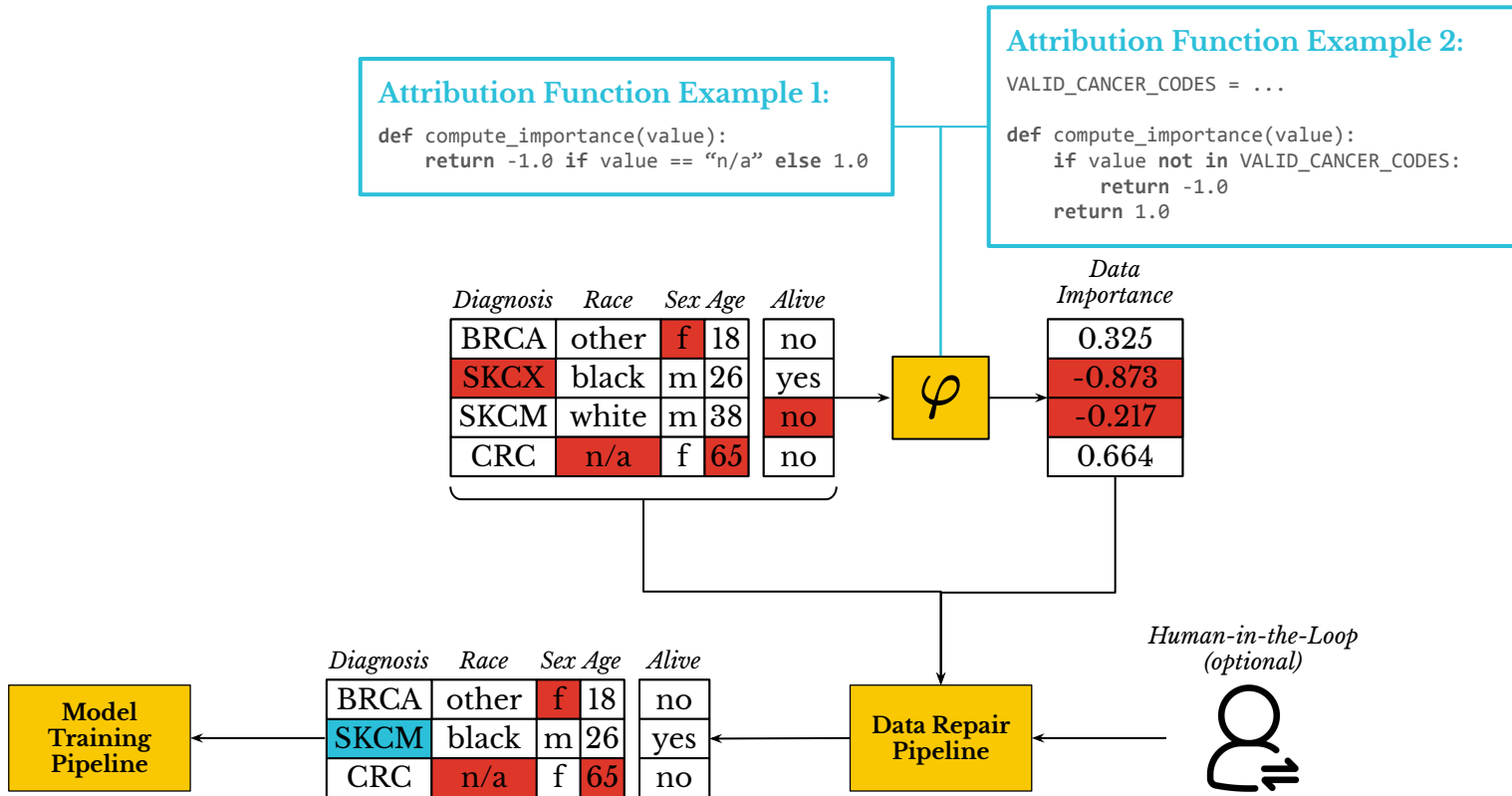
Challenge: Can we define a unified way to think about identifying data errors?

We can define a data attribution function



Recall: Data errors are any kind of problem in the training data that cause problems in model behavior.

How do we use importance to detect data errors?



What makes a good attribution function?

Design Consideration 1

Which model quality metric do we care about improving?

Correctness Metric

accuracy

f1 score

Fairness Metric

equalized odds

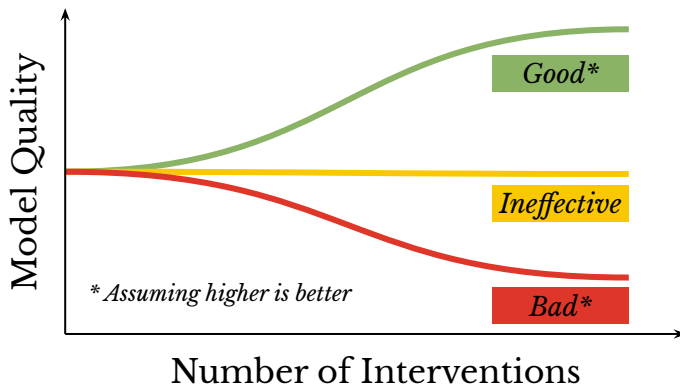
predictive parity

Stability Metric

entropy

Recall:

Data errors are any kind of problem in the training data that cause problems in model behavior.



Design Consideration 2

What kind of intervention do we intend to apply?



Discard



Repair



Something Else

Challenge: *How do we define an effective attribution function?*

1) Introducing the Concept of Data Importance

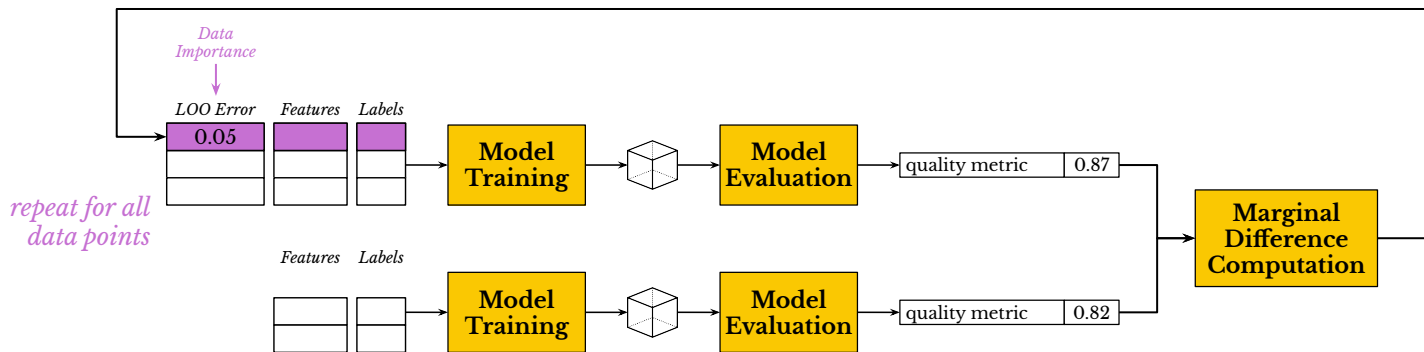
2) Examples of Data Attribution Functions

3) Case Study of Shapley Value as a Measure of Importance

4) Applications of Data Importance

Leave-one-Out Error

[Approach: Marginal Contribution]



Insights:

- Removing important data points affects model quality.

Approach:

- Remove a data point from the training set, train and evaluate the model again
- Interpret the difference in model quality as data importance.

Benefits:

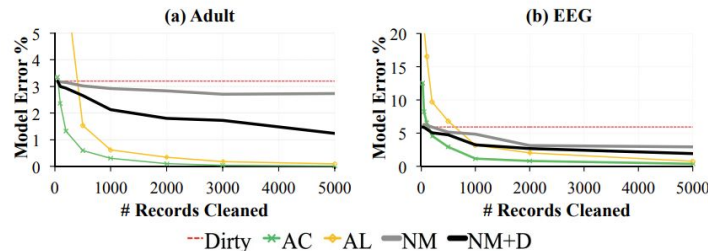
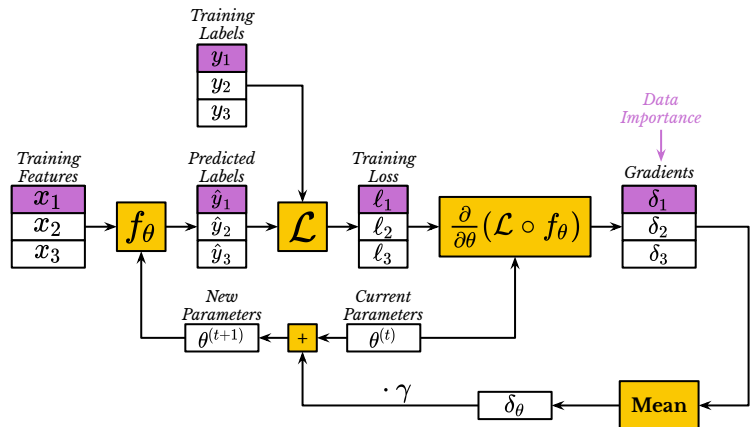
- Very simple to implement.

Shortcomings:

- Requires re-training the model once for each data point.
- Treats data points independently.

Error Gradient

[Approach: Gradient]



Insights:

- Data points vary in their contribution to the gradients that update the model.

Approach:

- Importance is proportional to the magnitude of the gradient.

Benefits:

- Simple to compute.

Shortcomings:

- Treats data points independently.

ActiveClean: Interactive Data Cleaning For Statistical Modeling

Sanjay Krishnan, Jiaman Wang¹, Eugene Wu², Michael J. Franklin, Ken Goldberg³
¹UC Berkeley, ²Simon Fraser University, ³Columbia University
 {sanjaykrishnan, franklin, goldberg}@berkeley.edu, jwang@sfu.ca, ewu@columbia.edu

ABSTRACT

Analysts often clean dirty data iteratively—cleaning some data, re-evaluating the analysis, and then cleaning more data based on the results. We explore the iterative cleaning process in the context of statistical model training, which is an increasingly popular form of data analysis. We propose ActiveClean, which tailors for properties of the analysis, cleaning an iterative model training while preserving convergence guarantees. ActiveClean supports an interactive class of models (all current tree models), linear models (generalized linear models), and probabilistic models (naïve Bayes, logistic regression, and SVMs), and iteratively cleaning these records likely to affect the results. We evaluate ActiveClean on the real-world datasets UCI Adult, UCI Credit, MNIST, IMDB, and Dailies For Dishes with both real and synthetic errors. The results show that our proposed algorithm can improve model accuracy by up to 2.6% for the same amount of data cleaned. Furthermore, it results in the highest and on all real dirty datasets, ActiveClean returns more accurate models than random sampling and Active Learning.

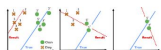


Figure 6: (a) Model error percentage vs # Records Cleaned for the Adult dataset. The dirty examples are initially 1.7 and the amount of records cleaned is 1000. (b) Model error percentage vs # Records Cleaned for the EEG dataset. The dirty examples are initially 1.7 and the amount of records cleaned is 1000. The results show that our proposed algorithm can improve model accuracy by up to 2.6% for the same amount of data cleaned. Furthermore, it results in the highest and on all real dirty datasets, ActiveClean returns more accurate models than random sampling and Active Learning.

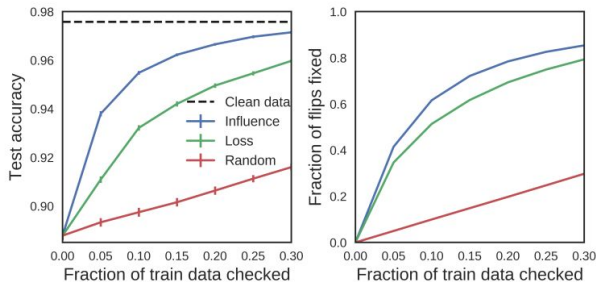
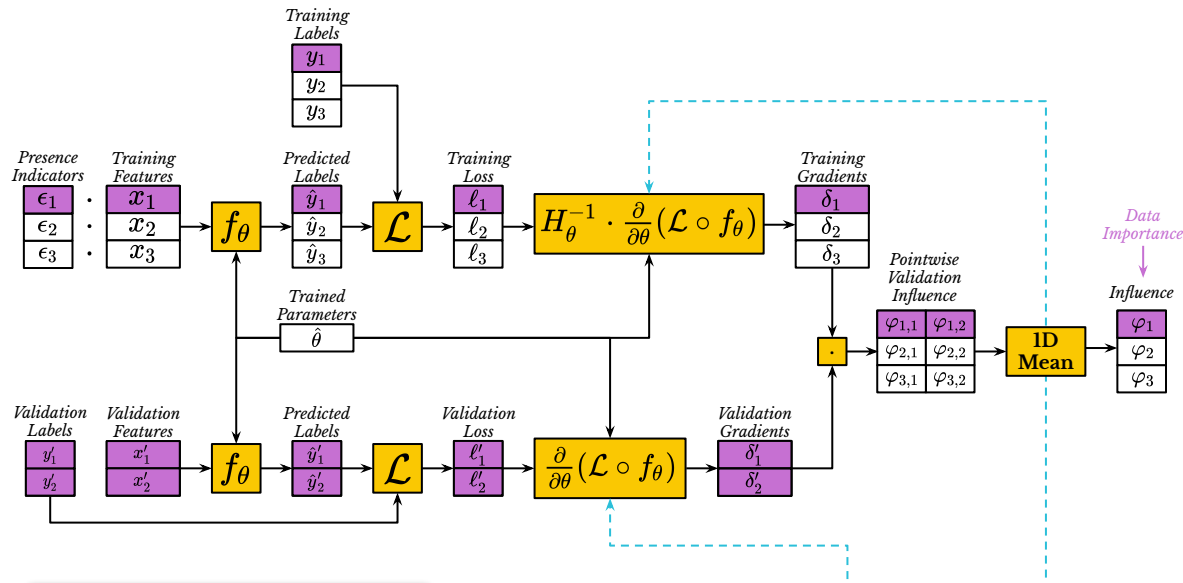
[Krishnan VLDB'16]

Krishnan, Sanjay, et al. "Activeclean: Interactive data cleaning for statistical modeling." Proceedings of the VLDB Endowment 9.12 (2016): 948-959.

[Paper] [Website]

Influence Function

[Approach: Marginal Contribution, Gradient]



Insights:

- The marginal contribution of a single data point can be approximated with gradients.

Approach:

- Introduce presence indicator variables ϵ for each data point and compute the gradient w.r.t. ϵ .

Benefits:

- Easily applicable to arbitrarily complex (twice) differentiable machine learning models.

Shortcomings:

- Treats data points independently.

[Koh ICML '17]

Koh, Pang Wei, and Percy Liang. "Understanding black-box predictions via influence functions." International conference on machine learning. PMLR, 2017. [\[Paper\]](#) [\[Code\]](#)

Understanding Black-box Predictions via Influence Functions

Pang Wei Koh¹, Percy Liang²

Abstract

How can we explain the predictions of a black-box model? In this paper, we use influence functions — a classic technique from robust statistics — to trace a model's prediction through the learning algorithm and back to its training data, thereby identifying training points most responsible for a given prediction. To scale to influence functions in modern machine learning settings, we develop a simple, efficient implementation that requires only access to gradients and Hessian-vector products. We show that even on non-convex and non-differentiable models where these black-box models, approximations to influence functions can still provide valuable insights. On linear models and convolutional neural networks, we demonstrate that influence functions are useful for multiple purposes: understanding model behavior, debugging models, de-

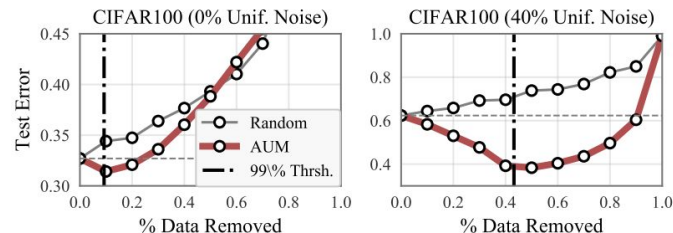
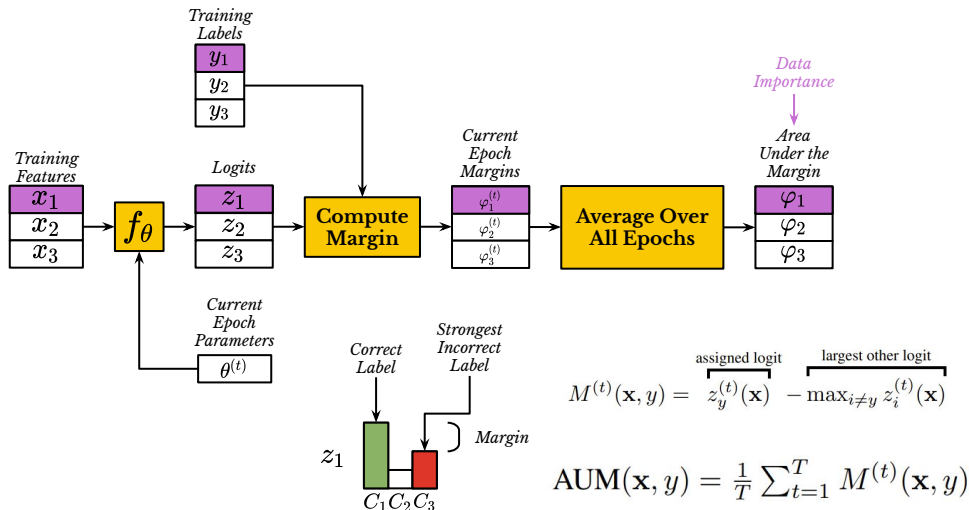
scribing (Bhattarai et al., 2015) or by perturbing the test point to see how the prediction changes (Gong et al., 2015; Li et al., 2016; Datta et al., 2016; Advers et al., 2015). These results explain the prediction in terms of the model, but how can we explain where the model came from?

In this paper, we tackle this question by tracing a model's prediction through its learning algorithm and back to the training data, where the model parameters ultimately derive from. To formalize the impact of a training point on a prediction, we ask the counterfactual: what would happen if we did not have this training point, or if the values of this training point were changed slightly?

Answering this question by perturbing the data and examining the model can be prohibitively expensive. To overcome this problem, we use influence functions, a classic technique from robust statistics (Cook & Weisberg, 1982) that tells us how the model parameters change as we sprinkle a training point by an infinitesimal amount. This allows us to "differentiate through the learner" to estimate by closed-

Area Under the Margin

[Approach: Uncertainty Analysis]



Insights:

- If similar samples have the same label, the model will learn to activate only the correct logit.
- In the presence of mislabeled samples, the model will learn to activate alternative logits.

Approach:

- The importance of a data point is proportional to its margin averaged across all training epochs.

Benefits:

- Very simple to implement in a wide array of models.
- Does not rely on a separate clean dataset.

Shortcomings:

- Focuses only on label noise.

Identifying Mislabeled Data using the Area Under the Margin Ranking

Geoff Pleiss¹
Columbia University
gpleiss@columbia.edu

Thang Luong²
Stanford University
tluong@stanford.edu

Ethan Haber³
AOL
ehaber@comcast.net

Kilian Q. Weinberger¹
Stanford University

Abstract

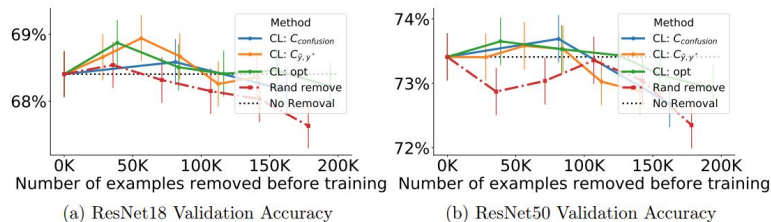
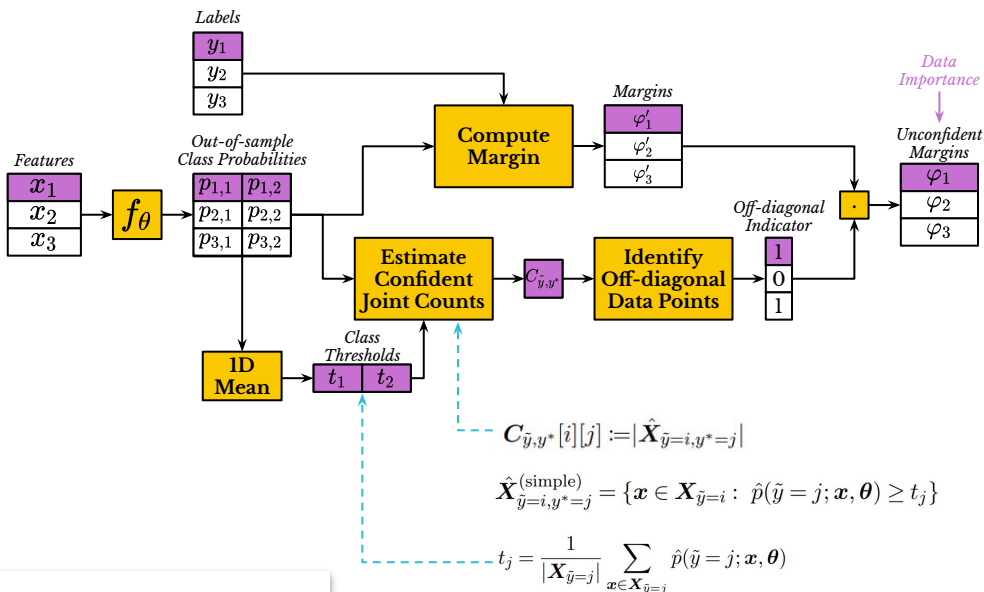
Not all data in a typical training set help with generalization; some samples can be overly challenging or simply mislabeled. This paper introduces a new method to identify such samples and mitigate their impact when training neural networks. At the heart of our algorithm is the Area Under the Margin (AUM) statistic, which captures differences in the training dynamics of clean and mislabeled samples. A simple procedure—adding an extra clean population with proportionally mislabeled derived samples—allows a AUM upper bound that isolates mislabeled data. This approach consistently improves upon prior work on synthetic and real-world

[Pleiss NeurIPS '20]

Pleiss, Geoff, et al. "Identifying mislabeled data using the area under the margin ranking." Advances in Neural Information Processing Systems 33 (2020): 17044-17056. [\[Paper\]](#) [\[Blog\]](#) [\[Code\]](#)

Unconfident Margins

[Approach: Uncertainty Analysis]



Insights:

- Given a data point, if a model assigns a higher than average probability to some specific class, it is likely because most similar data points have the same class label. This is likely to be the true label of that data point.

Approach:

- Identify likely mislabeled data points and assign negative importance using the margin. Remaining data points get zero importance.

Benefits:

- Very simple to implement in a wide array of models.
- Does not rely on a separate clean dataset.

Shortcomings:

- Focuses only on label noise.
- Relies on having an adequately powerful model.

Journal of Artificial Intelligence Research 70 (2021) 1373-1411

Submitted 06/2020; published 04/2021

Confident Learning: Estimating Uncertainty in Dataset Labels

Curtis G. Northcutt
Massachusetts Institute of Technology
Department of EECS, Cambridge, MA, USA
CU.NORTH@MIT.EDU

Lu Jiang
Google Research, Mountain View, CA, USA
LUJIAN@GOOGLE.COM

Isaac L. Chuang
Massachusetts Institute of Technology
Department of Physics, Cambridge, MA, USA
ICHUANG@MIT.EDU

Abstract

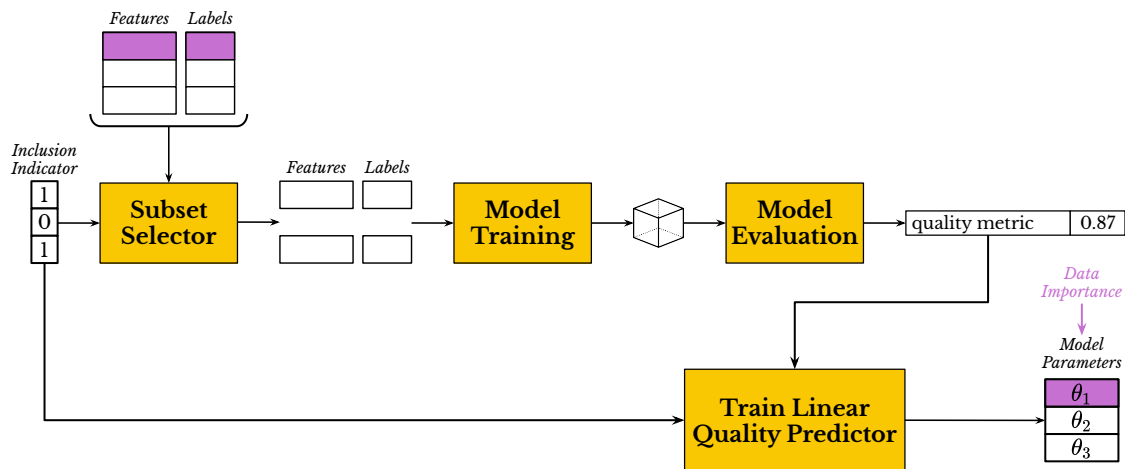
Learning exists in the context of data, yet notions of confidence typically focus on model predictions, not label quality. Confident Learning (CL) is an alternative approach which focuses instead on label quality by characterizing and identifying label errors in datasets, based on the principles of proving noisy data, counting with probabilistic thresholds to estimate noise, and ranking examples to train with confidence. Whereas numerous studies have developed these principles independently, here, we combine them, building on the assumption of a class-conditional noise process to directly estimate the joint distribution between noisy labels and uncorrupted unlabeled inputs. This results in a general and

[Northcutt JAIR '21]

Northcutt, Curtis, Lu Jiang, and Isaac Chuang. "Confident learning: Estimating uncertainty in dataset labels." *Journal of Artificial Intelligence Research* 70 (2021): 1373-1411. [\[Paper\]](#) [\[Blog\]](#) [\[Code\]](#)

Model Training Outcome

[Approach: Surrogate Data Model]



Insights:

- A linear model can be good at predicting the quality of a model trained on an arbitrary subset of the training data and tested on a single test example.

Approach:

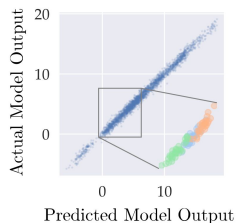
- Train a linear quality predictor and interpret its parameters as data importance.

Benefits:

- Conceptually simple yet powerful framework for analyzing datasets.

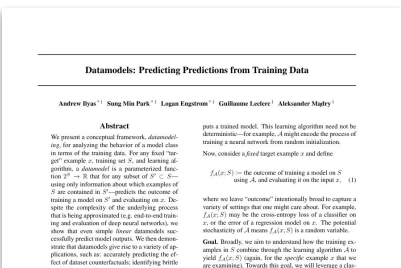
Shortcomings:

- The original method requires retraining the model many times.



[Ilyas ICML '22]

Ilyas, Andrew, et al. "Datamodels: Predicting Predictions from Training Data." Proceedings of the 39th International Conference on Machine Learning. 2022. [\[Paper\]](#) [\[Blog\]](#) [\[Code\]](#)



- 1) Introducing the Concept of Data Importance
- 2) Examples of Data Attribution Functions
- 3) Case Study of Shapley Value as a Measure of Importance**
- 4) Applications of Data Importance**

Improving Upon the Marginal Contribution Methods

Recall

Marginal contribution methods treat data points independently, ignoring any interactions that might exist.

Consequence

Let there be a data point that has high importance. If we make two copies of that data point, their individual marginal contribution to the dataset as a whole will be zero.

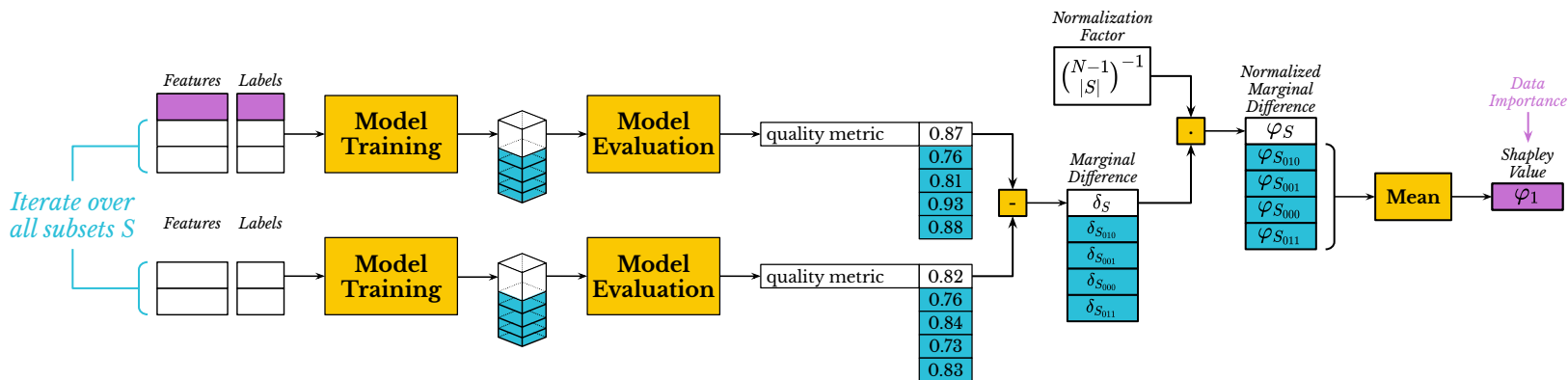
Approach

We should measure marginal contribution over all subsets.

Shapley value

A standard method from game theory for distributing surplus among a coalition of players.

$$\varphi_i = \frac{1}{N} \sum_{S \subseteq X \setminus \{i\}} \binom{N-1}{|S|}^{-1} (u(S \cup \{i\}) - u(S))$$



Effectiveness at Data Debugging

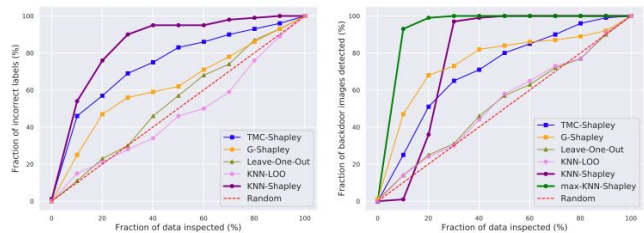


Figure 2: The experiment result of (a) noisy label detection on fashion-MNIST dataset; (b) instance-based watermark removal on MNIST dataset; (c) data summarization on UCI Adult Census dataset [15]; (d) data acquisition on MNIST dataset with injected noise. In (a)-(b) the “random” line shows the results of random guess; while in (c)-(d), the “random” line corresponds to the empirical results of the random baseline introduced in Section 4.1.

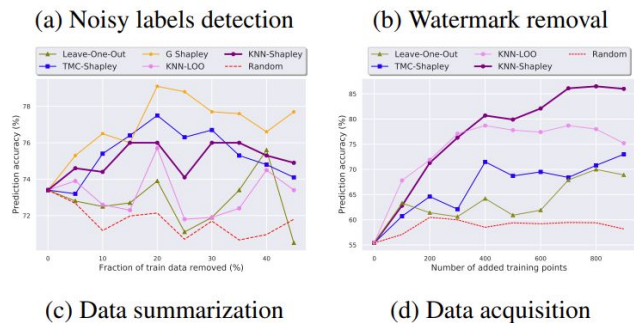


Table 2: Domain adaptation between MNIST and USPS.

Method	MNIST → USPS	USPS → MNIST
	  →  	  →  
KNN-Shapley	31.70% → 47.00%	23.35% → 29.80%
KNN-LOO	31.70% → 37.40%	23.35% → 24.50%
TMC-Shapley	31.70% → 44.90%	23.35% → 29.55%
LOO	31.70% → 29.40%	23.35% → 23.53%

This CVPR 2021 paper is the Open Access version, provided by the Computer Vision Foundation. Except for this watermark, it is identical to the accepted version; the final published version of the proceedings is available on IEEE Xplore.

Scalability vs. Utility: Do We Have to Sacrifice One for the Other in Data Importance Quantification?

Ruoxi Jia¹ Fan Wu^{2*} Xuehui Sun³ Jiaxin Xu⁴ David Dao⁵
 Bhavya Kulkarni⁶ Ce Zhang⁷ Bo Li⁸ Dim Song⁹
¹Virginia Tech ²UIUC ³Shanghai Jiaotong University ⁴UC Irvine ⁵ETH Zurich
⁶Lawrence Livermore National Laboratory ⁷UC Berkeley
 {ruoxi.jia@vt.edu, fanwu@uiuc.edu, xuehui.sun@sjtu.edu.cn, jiaxin.xu@uci.edu, ddaod@cs.berkeley.edu, ce.zhang@liit.edu.cn, kailixu@sjtu.edu.cn, dsong@cs.berkeley.edu}

Abstract

Quantifying the importance of each training point to a learning task is a fundamental problem in machine learning and the returned importance scores have been leveraged to guide a range of data-centric tasks such as data summarization and dataset adaptation. Our simple idea is to use the leave-one-out error of each training point to indicate its importance. Recent work has also proposed to use the Shapley value as a global importance measure for each data point. In this paper, we are driven by two questions around this fundamental problem. Our contribution is a novel theoretical

1. Introduction

[Jia CVPR '21]

Jia, Ruoxi, et al. "Scalability vs. utility: Do we have to sacrifice one for the other in data importance quantification?." Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. 2021. [\[Paper\]](#) [\[Code\]](#)

Benefits and Challenges

Beneficial Properties of the Shapley Value

Symmetry

If two data points have the same contribution to every subset, their value should be the same.

Efficiency

The sum of importances of all data points should equal the marginal contribution of the entire set over an empty set.

Linearity

If the utility function can be expressed as a sum of two other functions, then the importance of a data point using the combined function should equal the sum of importances computed using the individual functions.

Null Player

If a data point has a zero marginal contribution to every single subset, its importance should be zero.

Key Challenge

The number of subsets to enumerate is exponential, making it intractable to compute the exact Shapley value for an arbitrary model.

$$\varphi_i = \frac{1}{N} \sum_{S \subseteq X \setminus \{i\}} \binom{N-1}{|S|}^{-1} (u(S \cup \{i\}) - u(S))$$

Approximation: Monte Carlo Sampling

Challenge

Computing Shapley values is intractable.

Insight

Since Shapley value can be seen as a statistic over exponentially many subsets, we can estimate it using Monte Carlo sampling.

Approach

Use the permutation-based definition of the Shapley value and sample permutations.

$$\varphi_i(v) = \frac{1}{n!} \sum_R [v(P_i^R \cup \{i\}) - v(P_i^R)]$$

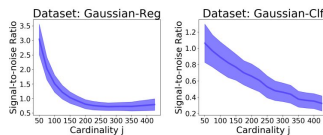
$$\phi_i = \mathbb{E}_{\pi \sim \Pi} [V(S_\pi^i \cup \{i\}) - V(S_\pi^i)]$$

Challenge

We need many Monte Carlo samples to produce good estimates.

Insight

When estimating the marginal contribution of a data point to a subset, we empirically observe that larger subsets incur a slower signal-to-noise ratio.

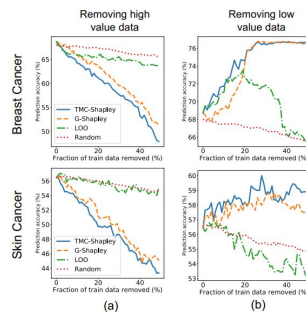


Approach

Leverage the importance sampling strategy and apply a larger weight to smaller subsets, based on the beta distribution.

Benefits

Estimating the Shapley value becomes tractable and is shown to be effective at identifying important data points.



Shortcomings

Each Monte Carlo sample relies on retraining the model from scratch, which is expensive for large models.

[Kwon AISTATS '22]

Kwon, Yongchan, and James Zou. "Beta Shapley: a Unified and Noise-reduced Data Valuation Framework for Machine Learning." International Conference on AI and Statistics. 2022. [\[Paper\]](#) [\[Code\]](#)

[Ghorbani ICML '19]

Ghorbani, Amirata, and James Zou. "Data shapley: Equitable valuation of data for machine learning." International conference on machine learning. PMLR, 2019. [\[Paper\]](#) [\[Code\]](#)



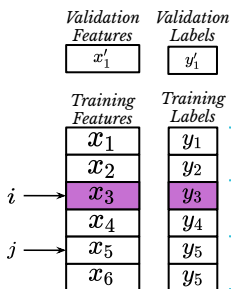
To get good Shapley value estimates, we need to retrain the model many times.

The simple KNN classifier can make it easy to design efficient and exact algorithms.

Use the KNN model as a proxy to develop an exact Shapley computation algorithm with polynomial time complexity.

- We are computing the Shapley value of data point i
- Data is sorted by similarity to the validation data point

Since $K=1$, for any subset S , the top-1 data point will determine the model prediction.



If the subset S contains these data points,
the data point i will not be the top-1.

If data point j is going to become the top-1 after i is removed, all data points above it cannot be included in S , while the ones below may or may not be included in S .

$$\varphi_i = \frac{1}{N} \sum_{S \subseteq X \setminus \{i\}} \underbrace{\binom{N-1}{|S|}^{-1} (u(S \cup \{i\}) - u(S))}_{\text{marginal contribution of } i \text{ to } S}$$

If data point i is not in the top-1, this term will be zero.

$$\varphi_i(t) = \frac{1}{N} \sum_{j=i+1}^N \sum_{a=1}^{n-j} \binom{N-1}{a}^{-1} (u(\{i\}) - u(\{j\})) \binom{N-j}{a}$$
$$\varphi_i(t) = \frac{1}{N} \sum_{j=i+1}^N (u(\{i\}) - u(\{j\})) \binom{N-j}{j+1}$$

After sorting the data, we can compute exact Shapley values in a single pass. Final computational complexity is $\mathcal{O}(N \log N)$

Approximation: Taylor Expansion

Challenge

If we are using a large and complex model, retraining will be extremely slow (preventing Monte Carlo approaches), and the KNN approximation will be biased.

Insight

Models trained with stochastic gradient descent (SGD) compute the loss function many times, over many random subsets of the training dataset. Furthermore, the changes in the model quality metric that are small enough to be effectively approximated with Taylor expansion.

Approach

Redefine the utility function to measure the cumulative impact of a training data point on the validation loss across gradient update steps.

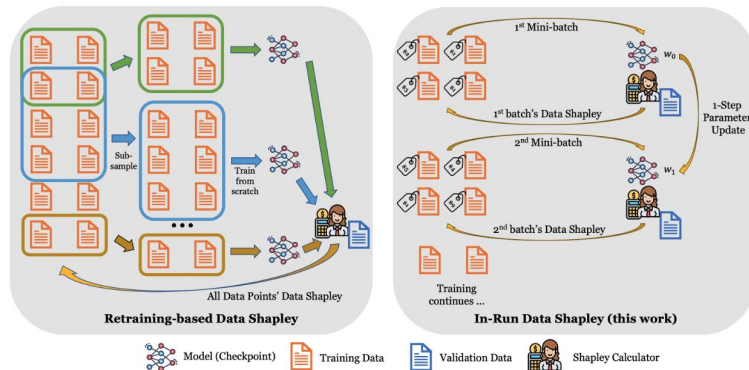
Redefined “local utility function” of subset S of a single SGD minibatch:

$$U^{(t)}(S; z^{(\text{val})}) := \underbrace{\ell(\tilde{w}_{t+1}(S), z^{(\text{val})})}_{\text{Model updated only using data from } S} - \underbrace{\ell(w_t, z^{(\text{val})})}_{\text{Model at SGD step } t}$$

$$\tilde{w}_{t+1}(S) := w_t - \eta_t \sum_{z \in S} \nabla \ell(w_t, z)$$

Redefined “global utility function” of subset S over the entire SGD run:

$$U(S) = \sum_{t=0}^{T-1} U^{(t)}(S)$$



Published as a conference paper at ICLR 2025

DATA SHAPLEY IN ONE TRAINING RUN

Jiachen T. Wang
Princeton University

Pranav Mittal
Princeton University

Dawn Song
UC Berkeley

Ravi Jia
Virginia Tech

ABSTRACT

Data Shapley offers a principled framework for attributing the contribution of data within machine learning contexts. However, the traditional notion of Data Shapley requires re-training models on various data subsets, which becomes computationally infeasible for large-scale models. Additionally, this retraining-based definition cannot evaluate the contribution of data for a specific model training run, which may often be of interest in practice. This paper introduces a novel concept, *In-Run Data Shapley*, which eliminates the need for model retraining and is specifically designed for assessing data contribution for a particular model of interest. In *In-Run Data Shapley*, we calculate the Shapley value for each gradient update iteration and accumulate these values throughout the training process. We present several techniques that allow the efficient scaling of *In-Run Data Shapley* to the size of federated models. In its most general implementation, our method adds negligible runtime overhead compared to standard model training. This framework efficiency improvement makes it possible to perform data attribution for the federated model pre-training stage. We present several case studies that offer fresh insights into pre-training data’s contribution and discuss their implications for copyright in generative AI and pre-training data curation.

[Wang ICLR ‘25]

Wang, Jiachen T., et al. "Data Shapley in One Training Run." The Thirteenth International Conference on Learning Representations. [Paper](#)

[Blog](#)

- 1) Introducing the Concept of Data Importance
- 2) Examples of Data Attribution Functions
- 3) Case Study of Shapley Value as a Measure of Importance
- 4) Applications of Data Importance**

Influence Function for Explaining Fairness Errors

Challenge

Data attribution gives us an ordered list of data points that impact model quality, but it does not explain what makes these data points impactful.

Insight

If we group important data points based on common predicates, we can derive more powerful conclusions about factors that cause models to underperform.

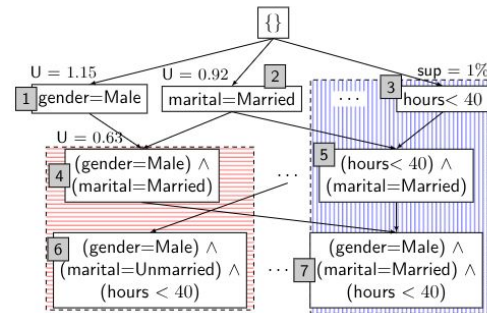
Approach

First, use influence functions to compute data importance with respect to fairness metrics. Second, use lattice-based search to identify combinations of predicates that define data subsets that are both small and impactful.

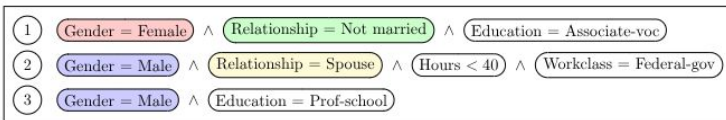
Data points ordered by importance

age	education	marital	...	gender	income
39	Bachelors	Never-married	...	Male	≤50K
53	11th	Never-married	...	Male	≤50K
28	Bachelors	Married-civ-spouse	...	Female	≤50K
37	Masters	Married-civ-spouse	...	Female	≤50K

Lattice-based search identifies predicates that select the most impactful training data subsets



Combinations of predicates that explain model behavior



[Zhu SIGMOD '22]

Pradhan, Romila, et al. "Interpretable data-based explanations for fairness debugging." Proceedings of the 2022 international conference on management of data. 2022. [Paper](#)

Debugging the LLM Retrieval Corpus

Challenge

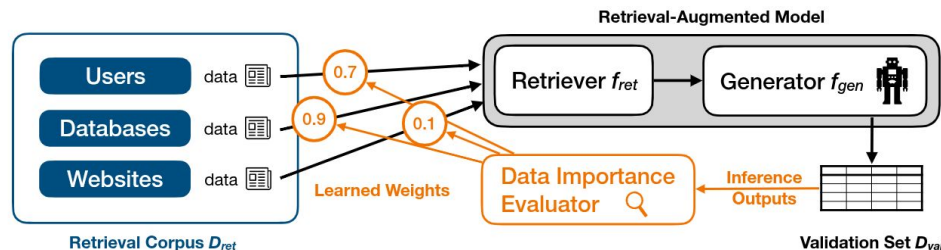
Retrieval augmented generation (RAG) is a widely used technique for providing pre-trained large language models (LLMs) with task-specific context. Data errors in the retrieval corpus have a negative impact on model quality.

Insight

The role of a retrieval corpus to an LLM is similar to the role of a training dataset to a classical ML model.

Approach

Define a data attribution function that will compute the importance of data points in the retrieval corpus. Use this to identify and debug data errors.



$$U(f_{gen}, f_{ret}, D_{val}, D_{ret}) := \sum_{x_i \subseteq D_{val}} U(f_{gen}(x_i, f_{ret}(x_i, D_{ret})))$$

$$\tilde{U}(w_1, \dots, w_M) := \sum_{S \subseteq D_{ret}} U(S) \underbrace{\prod_{d_i \in S} w_i \prod_{d_i \notin S} (1 - w_i)}_{P[S]}$$

DATASET	GPT-JT (6B)	GPT-JT (6B) W/ RETRIVAL				GPT-3.5 (175B)
		VANILLA	+LOO	+REWEIGHT	+PRUNE	
BUY	0.102	0.789	0.808	0.815	<u>0.813</u>	0.764
RESTAURANT	0.030	0.746	0.756	<u>0.760</u>	0.761	0.463

Improving Retrieval-Augmented Large Language Models via Data Importance Learning

Xiaozhong Lyu¹, Stefan Grathner², Samarth Singh³, Shuang Wu¹
 Hong Cao¹, Sebastian Schödl¹, Qi Zhang¹
¹ETH Zurich ²University of Amsterdam ³Apple

Abstract

Retrieval augmentation enables large language models to take advantage of external knowledge, for example on tasks like question answering and data imputation. However, the performance of such retrieval-augmented models is limited by the data quality of their underlying retrieval corpus. In this paper, we propose an algorithm based on additive importance learning for evaluating the data importance of retrieved data points. There are exponentially many terms in the additive extension, and one key contribution of this paper is a polynomial-time algorithm that computes, exactly, given a retrieval-augmented model with an additive utility function and a validation set, the data importance of data points in the retrieval corpus using the minimum extension of the model's utility function. We further propose an even more efficient (ϵ -) approximation algorithm. Our experimental results illustrate multiple strengths of the model's utility function.

[Lyu arXiv '23]

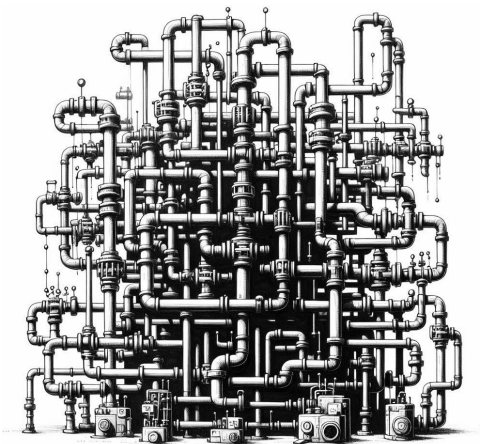
Lyu, Xiaozhong, et al. "Improving retrieval-augmented large language models via data importance learning." arXiv preprint arXiv:2307.03027 (2023). [\[Paper\]](#) [\[Code\]](#)

Key Takeaways of Part I

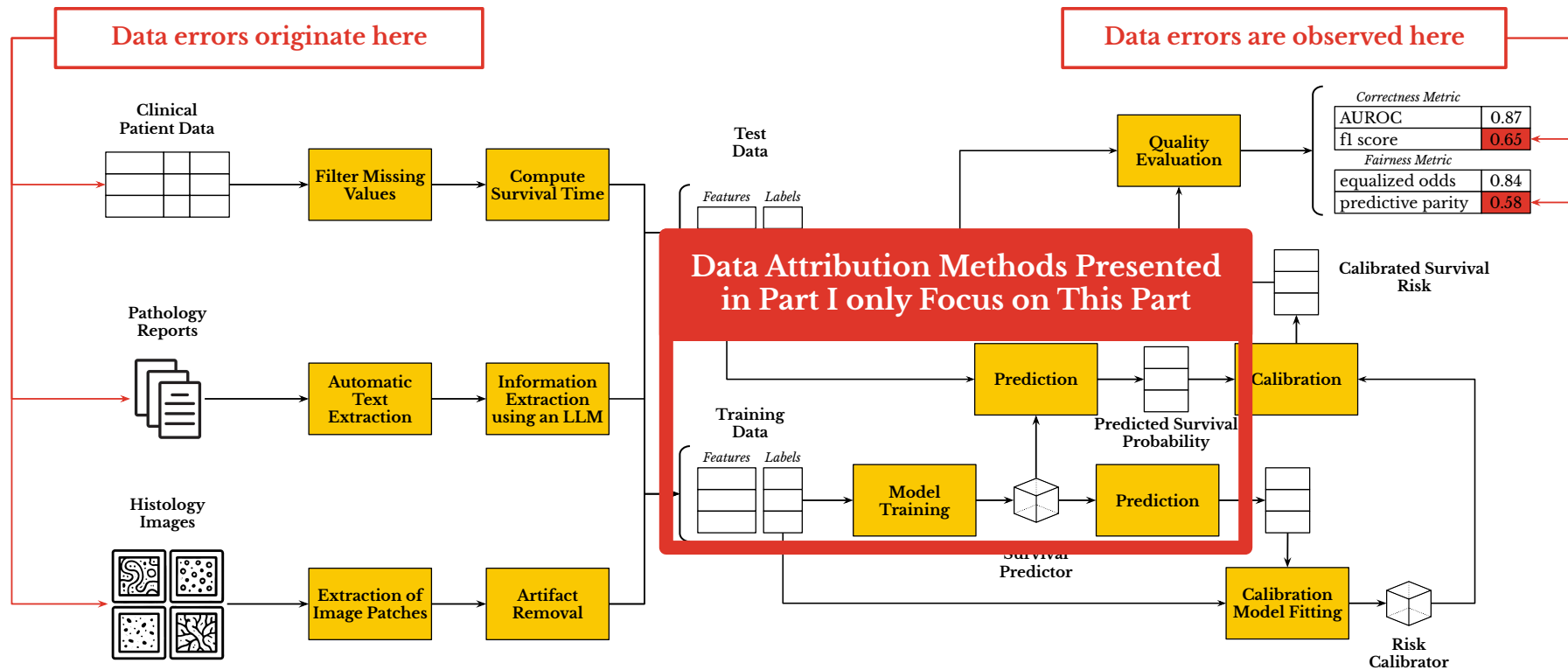
- Data attribution is a useful powerful framework for approaching the problem of data error detection.
- There are many existing data attribution methods with various strengths and shortcomings.
- The most powerful methods face scalability issues that have been tackled by existing research with many opportunities for future improvements.

Part II: Data Debugging in ML Pipelines

Sebastian Schelter



Gap between Attribution Methods and ML Pipelines



1) Gap between Attribution Methods and ML Pipelines

2) Libraries and Systems for ML Pipelines

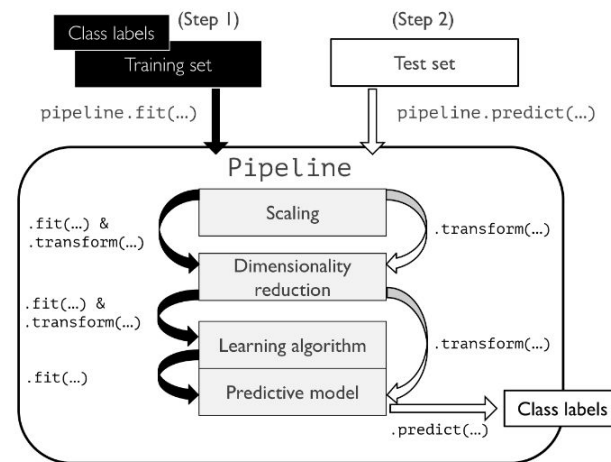
3) Characteristics of Real World ML Pipelines

4) Methods for Debugging ML Pipelines

Scikit-Learn

Highlights

- Among the most popular data science Python libraries
- Has implementations of many machine learning models, as well as feature encoding operators
- Introduced the **estimator/transformer abstraction** for composing complex, nested pipelines
 - **Transformer:** tuple-at-a-time transformation
 - **Estimator:** create a data-specific transformer via a global aggregation over the data



Source: <https://vitalflux.com/sklearn-machine-learning-pipeline-python-example/>

Journal of Machine Learning Research 12 (2011) 2825-2830

Submitted 11/11, Revised 11/11, Published 11/11

Scikit-learn: Machine Learning in Python

Fabian Pedregosa
Guillaume Varoquaux
Alexandre Gramfort
Vincent Michel
Bernard Thériault
Pedro de Amorim
Stéphane Mallat
Olivier Grisel
Nicolas
Thomas
Mathieu Blondel
J. F. B. de
Peter Prettenhofer
Bachman
Bachman

FABIAN.PEDREGOSA@INRIA.FR
GUILLAUME.VAROQUAUX@INRIA.FR
ALEXANDRE.GRAMFORT@INRIA.FR
VINCENT.MICHEL@INRIA.FR
BERNARD.THERIAULT@INRIA.FR
PEDRO.DE.AMORIM@INRIA.FR
STEPHANE.MALLAT@INRIA.FR
OLIVIER.GRISSEL@INRIA.FR
NICOLAS
THOMAS
MATHEU.BLODDEL@INRIA.FR
J.F.B.D.E
PETER.PRETTENHOFER@INRIA.FR
BACHMAN
BACHMAN

[Pedregosa] MLR '11

Pedregosa, Fabian, et al. "Scikit-learn: Machine learning in Python." the Journal of machine Learning research 12 (2011): 2825-2830. [\[Paper\]](#)

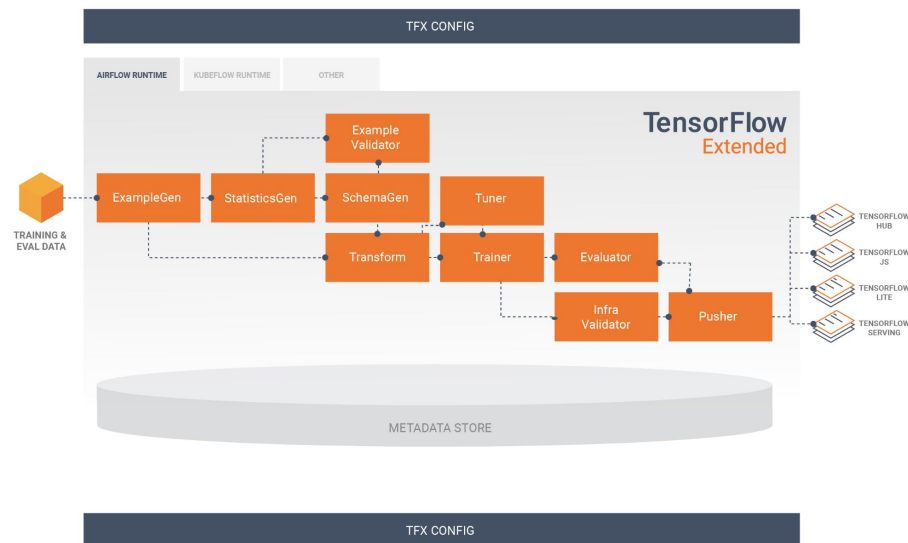
[\[Website\]](#) [\[Code\]](#)

Tensorflow Extended (TFX)



Highlights

- End-to-end platform for production ML pipelines
- Built on TensorFlow and optimized for scalability, strong emphasis on model validation and monitoring
- Includes reusable components for pipelines, inspired by estimator/transformer paradigm
- Apache Beam for dataflow operations, Tensorflow for numerical operations



Source: <https://www.tensorflow.org/tfx/guide>

KDD 2017 Applied Data Science Paper

KDD 17, August 13–17, 2017, Halifax, NS, Canada

TFX: A TensorFlow-Based Production-Scale Machine Learning Platform

Denis Boyko, Eric Breck, Hong-Tai Cheng, Noah Feid, Chun Yu Ho, Zakaria Hogg, Sahm Hyeon, Mustafa Isir, Vilan Jain, Levon Koc, Chih-Yen Kuo, Lukasz Lory, Clement Mowald, Akshay Narash Mahi, Nandini Pothuri, Sakshi Ramesh, Sully Roy, Steven Ruijter, Whang, Martin Wicke, Jack Wilkerson, Xin Zhang, Martin Zinkevich, Google Inc.*

ABSTRACT

Creating and maintaining a platform for reliable producing and deploying machine learning models requires careful orchestration of many components: a library for generating models based on training data, modules for monitoring and testing models, as well as services, and finally infrastructure for serving models in production. This becomes particularly challenging when data changes over time and fresh models need to be produced continuously. To address such challenges, we often draw on key using data and custom scripts developed by individual teams for specific use cases, leading to duplicated effort and fragile systems with high technical debt.

We present TensorFlow Extended (TFX), a TensorFlow-based production-scale machine learning platform implemented at Google. By integrating the abovementioned components into one platform, we were able to standardize the components and reduce the technical debt. In this paper, we describe the architecture of TFX and its components.

[Baylor SIGKDD '17]

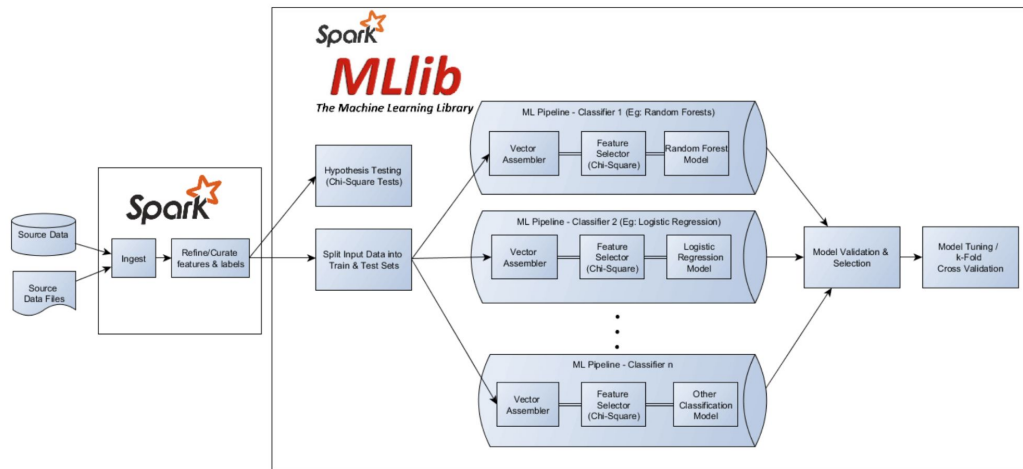
Baylor, Denis, et al. "Tfx: A tensorflow-based production-scale machine learning platform." Proceedings of the 23rd ACM SIGKDD international conference on knowledge discovery and data mining. 2017. [Paper](#) | [Website](#) | [Code](#)

Spark MLlib



Highlights

- Built on top of Apache Spark
- Includes implementations for classification, regression, clustering, collaborative filtering, and dimensionality reduction
- Works natively with Spark DataFrames, SQL, and streaming data
- Adoption of estimator/transformer paradigm from scikit-learn



Source: <https://www.qubole.com/developers/spark-getting-started-guide/workflow>

Journal of Machine Learning Research 17 (2016) 1-7

Submitted 5/15; Published 5/16

MLlib: Machine Learning in Apache Spark

Xiangrui Meng
Databricks, 100 Spear Street, 10th Floor, San Francisco, CA 94105

MENGX@DATABRICKS.COM

Joseph Bradley
Databricks, 100 Spear Street, 10th Floor, San Francisco, CA 94105

JOEB@DATABRICKS.COM

Burak Yavuz
Databricks, 100 Spear Street, 10th Floor, San Francisco, CA 94105

BURAK@DATABRICKS.COM

Evan Sparks
1700 Redwood City, CA 94060

SPARKS@DATABRICKS.COM

Shivaram Venkataratnam
1700 Redwood City, CA 94060

SHIVARAM@DATABRICKS.COM

Darwin Liu
Databricks, 100 Spear Street, 10th Floor, San Francisco, CA 94105

DAVID@DATABRICKS.COM

Jeremy Freeman
IBM Research, 5600 S. Hayes Ave., San Jose, CA 95128

FRJ@IBM.COM

DB Tadi
Netflix, 960 University Ave., Los Gatos, CA 95032

DT@NETFLIX.COM

Masahito Amaki
Google, 1600 Amphitheatre Parkway, Mountain View, CA 94035

AMAKI@GOOGLE.COM

[Meng] JMLR '16]

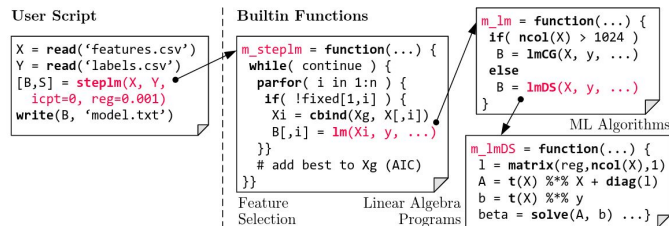
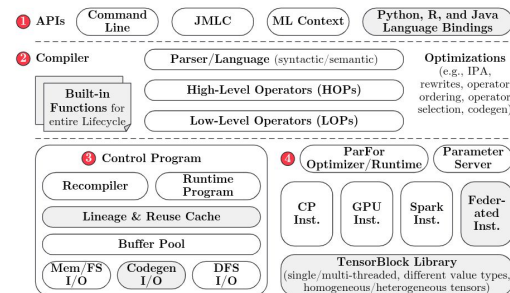
Meng, Xiangrui, et al. "Mllib: Machine learning in apache spark." Journal of Machine Learning Research 17.34 (2016): 1-7. [\[Paper\]](#) [\[Website\]](#) [\[Code\]](#)

Apache SystemDS



Highlights

- Designed for scalable and efficient execution on both single-node and distributed environments
- Offers a high-level scripting language for expressing ML algorithms and workflows with a declarative R-like language
- Performs cost-based optimization and automatic operator selection for efficient execution across different hardware endpoints
- Optimised feature encoders based on estimator/transformer paradigm



UPLIFT: Parallelization Strategies for Feature Transformations in Machine Learning Workloads

Arnab Phani
Graz University of Technology
phani@igi.tugraz.at

Lukas Erbacher
Graz University of Technology
luka.erbacher@student.tugraz.at

Matthias Boehm
Graz University of Technology
mboehm@igi.tugraz.at

ABSTRACT
Data science applications are typically exploratory. An integral task of such pipelines are feature transformations, which transform raw data into statistical features or features for training or scoring. They exist in a wide variety of transformations for different data

Journal of Machine Learning Research 17 (2016) 1-7

Submitted 5/15, Published 4/16

MLlib: Machine Learning in Apache Spark

Xiangrui Meng
Databricks, 400 Spear Street, 13th Floor, San Francisco, CA 94105
xmeng@databricks.com

Joseph Bradley
Databricks, 400 Spear Street, 13th Floor, San Francisco, CA 94105
jbradley@databricks.com

Burak Yigit
Databricks, 400 Spear Street, 13th Floor, San Francisco, CA 94105
byigit@databricks.com

Ryan Sprack
UC Berkeley, 405 Soda Hall, Berkeley, CA 94720
rsprack@cs.berkeley.edu

Shivaram Venkatasubramanian
UC Berkeley, 405 Soda Hall, Berkeley, CA 94720
shivaram@cs.berkeley.edu

[Boehm CIDR '20]

Boehm, Matthias, et al. "SystemDS: A Declarative Machine Learning System for the End-to-End Data Science Lifecycle." 10th Conference on Innovative Data Systems Research. 2020. [\[Paper\]](#) [\[Website\]](#) [\[Code\]](#)

[Phani VLDB '20]

Phani, Arnab, et al. "UPLIFT: parallelization strategies for feature transformations in machine learning workloads." Proceedings of the VLDB Endowment, Volume 15, Issue 11, 2020. [\[Paper\]](#)

ML Pipelines in the Cloud



Netflix Metaflow

[\[Website\]](#) [\[Documentation\]](#)

Highlights

- *Notebook based development environment*
- *Storing and tracking of code, data and models*
- *Scaling from local execution to the cloud*



Amazon SageMaker

Amazon SageMaker Pipelines

[\[Website\]](#) [\[Documentation\]](#)

Highlights

- *Define, automate, and manage end-to-end ML workflows*
- *Automatically tracks pipeline artifacts*
- *Leverages AWS Cloud infrastructure*



Azure Machine Learning

Azure Machine Learning Pipelines

[\[Website\]](#) [\[Documentation\]](#)

Highlights

- *Orchestration of ML workflows with reusable, modular pipeline components*
- *Versioning, monitoring, and CI/CD integration*



Vertex.ai

Vertex AI Pipelines

[\[Website\]](#) [\[Documentation\]](#)

Highlights

- *Connects with Vertex AI services*
- *Tracks pipeline steps, metadata, and artifacts*
- *Orchestrates ML workflows on Google Cloud*

- 1) Gap between Attribution Methods and ML Pipelines
- 2) Libraries and Systems for ML Pipelines
- 3) Characteristics of Real World ML Pipelines**
- 4) Methods for Debugging ML Pipelines

Study of Pipelines at Google

Highlights

- Study of 3000 production pipelines with over 450K models trained over a 4 month period
- About half the pipelines studied used data- and model-validation operators
- Input data typically has up to 100 features, but can have over 10K in extreme cases
- 53% of features were categorical, often with very large domains (averaging over 10M unique values)
- Training accounts for only 20% of the total runtime cost, over 30% is for model validation and 20% for data ingestion
- About 1/4 of model training runs results in model deployment
- Deep learning models account for 60% of pipelines

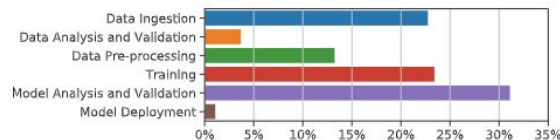


Figure 7: Compute cost of different operators.

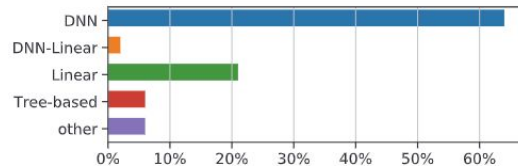
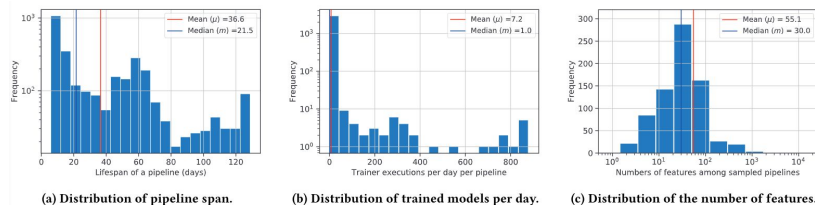
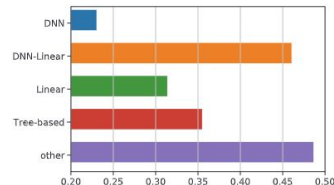
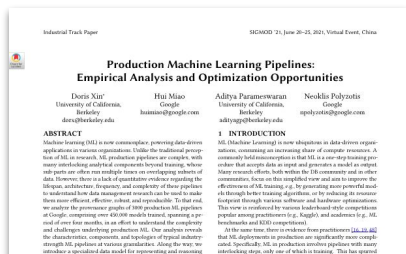


Figure 5: Percentage of Trainer runs with each model type



(f) Model type vs. likelihood of pushes.



[Xin SIGMOD '21]

Xin, Doris, et al. "Production machine learning pipelines: Empirical analysis and optimization opportunities." Proceedings of the 2021 international conference on management of data. 2021. [\[Paper\]](#)

Study of Pipelines at Microsoft

Highlights

- Study of over 8M public Jupyter notebooks on GitHub (from 2017, 2019, and 2020), and 2M enterprise pipelines developed with ML.NET
- Python is emerging as the de-facto standard language for data science (81% of notebooks in 2017 and 91% in 2020)
- Around 80% cells were linear (no conditional statements) and 76% were completely linear (no conditionals, classes, or functions)
- Libraries like numpy, matplotlib, pandas, and scikit-learn are used very frequently (e.g., numpy in >60% of notebooks)
- Few highly used libraries have significant coverage (e.g., top-10 cover ~40% of notebooks, top-100 cover ~75%), but there is a long tail
- Explicit ML pipelines (defined with sklearn.pipeline) are gaining traction but there are still 5 times more implicit pipelines in GitHub notebooks
- There is a large number of distinct operators, and a significant portion are user-defined (especially in ML.NET and implicit GitHub pipelines)

Data Science Through the Looking Glass: Analysis of Millions of GitHub Notebooks and ML.NET Pipelines

Pan Prallidas, Yuesha Zhu, Brian Katulic*, Jordan Holyst, Matteo Iervolino, Sahar Khatami, Brian Kish, Vitoria Eusebio, Wenxin Wu, Gu Zhang*, Marko Vucelja, Arvind Potharaju, Carlos Guestrin, Konstantinos Kamnitsas

ABSTRACT

The recent success of machine learning (ML) has led to an explosive growth of systems and applications built by an ever-growing community of system builders and data science (DS) practitioners. This quickly shifting landscape, however, is challenging for system builders and practitioners alike to follow. In this paper, we set out to capture this panorama through a wide-angle, performing the largest analysis of DS projects to date. In doing so, we question the ad-hoc nature of understanding the field and determine its trends. Specifically, we describe and analyze (a) over 8M public notebooks, available on GitHub (a) and (b) over 2M enterprise ML pipelines developed within Microsoft. Our analysis includes coarse-grained statistical characterizations, integrated analysis of libraries and pipelines, and coverage to, e.g., shall we use TensorFlow (1.14.0) or PyTorch (1.10.0)?, Discussing with experts in the field to further reconstruct views, too.

[Psallidas SIGMOD Record '22]

Psallidas, Fotis, et al. "Data science through the looking glass: Analysis of millions of github notebooks and ml. net pipelines." ACM SIGMOD Record 51.2 (2022): 30-37. [\[Paper\]](#)

Dimension	Metric	GH17	GH19	GH20
Notebooks	Total	1.23M	4.6M	8.7M
	Deduped	66.0%	65.5%	65.7%
	Linear	26.4%	29.1%	30.3%
	Completely Linear	21.2%	23.3%	24.6%
Languages	Python	81.7%	91.7%	91.1%
	Other	18.3%	8.3%	8.9%
Cells	Total	34.6M	143.1M	261.2M
Code Cells	Total	64.5%	66.4%	66.9%
	Deduped	41.0%	38.6%	38.5%
	Linear	72.1%	80.2%	79.3%
	Completely Linear	68.3%	76.1%	75.6%
Users	Total	100K	400K	697K

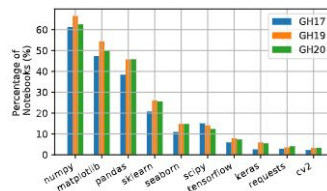


Figure 2: Top-10 used libraries.

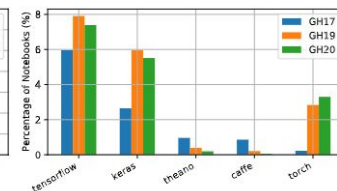
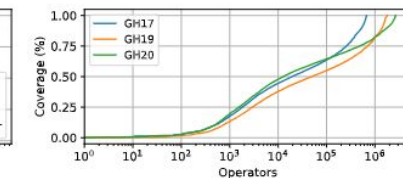
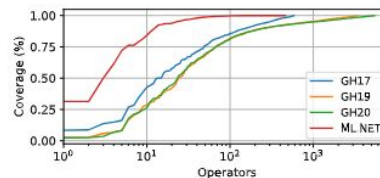


Figure 3: DL libraries usage percentages.

		GH17	GH19	GH20	ML.NET
#Pipelines	Implicit	164K	415K	1.4M	N/A
	Explicit	10K	129K	252K	29.7M
#Distinct Ops	Implicit	668K	1.8M	2.6M	N/A
	Explicit	584	3.4K	5.5K	23.5K

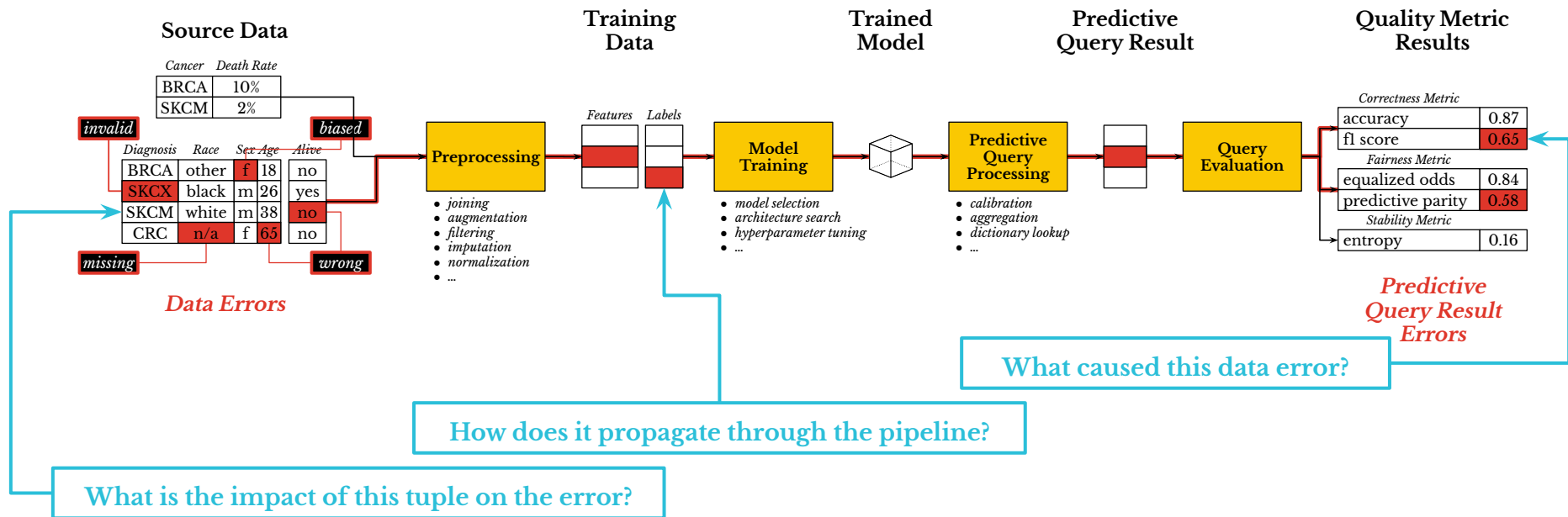


Observations

- No universal way to express ML pipelines, design often prioritises flexibility and ease-of-use
- Many pipelines combine relational / dataflow operators with ML-specific operators based on estimator/transformer abstraction
- Pipelines often executed via multiple runtimes
- Lack of algebraic operator semantics
- Lack of fine-grained data provenance

- 1) Gap between Attribution Methods and ML Pipelines
- 2) Libraries and Systems for ML Pipelines
- 3) Characteristics of Real World ML Pipelines
- 4) Methods for Debugging ML Pipelines**

How should we reason about pipelines?



Modeling ML Pipelines with “Logical Query Plans”

Challenge

Understanding of the semantics of operations and the flow of data required to reason about ML pipelines

Insight

Many common pipeline abstractions offer declarative operations, enables the extraction and definition of “logical query plans” modeling their operations

Approach

Instrument functions of Python data science libraries to extract query plan, enable annotation propagation through operators. Apply rule-based approaches to determine if an error has occurred.

Potential issues in preprocessing pipeline:

- 1 Join might change proportions of groups in data
- 2 Column 'age_group' projected out, but required for fairness
- 3 Selection might change proportions of groups in data
- 4 Imputation might change proportions of groups in data
- 5 'race' as a feature might be illegal!
- 6 Embedding vectors may not be available for rare names!

Python script for preprocessing, written exclusively with native pandas and sklearn constructs

```
# load input data sources, join to single table
patients = pandas.read_csv(...)
histories = pandas.read_csv(...)
data = pandas.merge([patients, histories], on=['ssn'])

# compute mean complications per age group, append as column
complications = data.groupby('age_group')
    .agg(mean_complications=('complications', 'mean'))
data = data.merge(complications, on='age_group')

# Target variable: people with frequent complications
data['label'] = data['complications'] >
    1.2 * data['mean_complications']

# Project data to subset of attributes, filter by counties
data = data[['smoker', 'last_name', 'county',
            'num_children', 'race', 'income', 'label']]
data = data[data['county'].isin(counties_of_interest)]

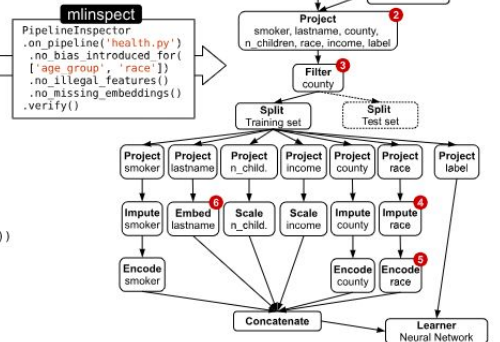
# Define a nested feature encoding pipeline for the data
impute_and_encode = sklearn.Pipeline([
    (sklearn.SimpleImputer(strategy='most_frequent')),
    (sklearn.OneHotEncoder())])
featurisation = sklearn.ColumnTransformer(transformers=[
    (impute_and_encode, ['smoker', 'county', 'race']),
    (Word2VecTransformer(), 'last_name')],
    (sklearn.StandardScaler(), ['num_children', 'income']))

# Define the training pipeline for the model
neural_net = sklearn.KerasClassifier(build_fn=create_model())
pipeline = sklearn.Pipeline([
    ('features', featurisation),
    ('learning_algorithm', neural_net)])

# Train-test split, model training and evaluation
train_data, test_data = train_test_split(data)
model = pipeline.fit(train_data, train_data.label)
print(model.score(test_data, test_data.label))
```

Corresponding dataflow DAG for instrumentation, extracted by *mlinspect*

Declarative inspection of preprocessing pipeline



The VLDB Journal 33.1 (2022): 1103–1126.
https://doi.org/10.1007/978-3-031-00029-4

SPECIAL ISSUE PAPER

Data distribution debugging in machine learning pipelines

Stefan Grafberger¹ · Paul Groth¹ · Julia Stoyanovich² · Sebastian Schuster¹

Received: 27 February 2021 / Revised: 9 September 2021 / Accepted: 2 December 2021 / Published online: 31 January 2022
© The Author(s), under exclusive license to Springer Nature GmbH Germany, part of Springer Nature 2022

Abstract

Machine learning (ML) is increasingly used to automate important decisions, and the risks arising from this widespread use are growing attention from policy makers, scientists, and the media. ML applications are often brittle with respect to data input data, which leads to concerns about their correctness, reliability, and fairness. In this paper, we describe *mlinspect*, a library that helps diagnose and mitigate technical bias that may arise during preprocessing steps in an ML pipeline. We refer to these problems collectively as *data distribution bugs*. The key idea is to extract a directed acyclic graph representation of the dataflow from a preprocessing pipeline and use this representation to automatically instrument the code with predicted operations. These inspections are based on a lightweight annotation propagation approach to propagate metadata such as change information throughout a pipeline. In contrast to existing work, *mlinspect* operates on declarative abstractions of popular data science libraries like estimator/transformer pipelines and does not require manual code instrumentation. We discuss the design and implementation of the *mlinspect* library and give a comprehensive and to-end example that illustrates its functionality.

Keywords Data debugging · Machine learning pipelines · Data preparation for machine learning

[Grafberger VLDB] ‘22]

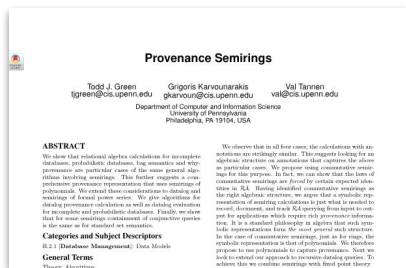
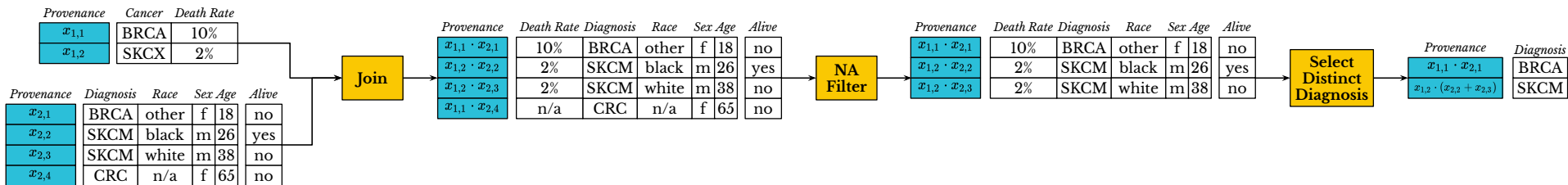
Grafberger, Stefan, et al. "Data distribution debugging in machine learning pipelines." The VLDB Journal 31.5 (2022): 1103–1126. [\[Paper\]](#) [\[Code\]](#)

Leveraging the Provenance Semiring Framework

Highlights

- *Theoretical framework analyzing the relationship between input and output tuples of relational queries*
- *Allows us to determine the presence of an output tuple as a function of the presence of the input tuples*
- *Easy to adapt for ML pipelines once logical query plan with “relational-like” operations is known*

Application to an Example Pipeline



[Green SIGMOD '07]

Green, Todd J., Grigoris Karvounarakis, and Val Tannen. "Provenance semirings." Proceedings of the twenty-sixth ACM SIGMOD-SIGACT-SIGART symposium on Principles of database systems. 2007. [Paper](#)

Debugging Preprocessing Pipelines with Datascope

[Attribution Function: Shapley Value]

Challenge

KNN Proxy methods not directly applicable to arbitrary pipelines. Presence of a single source data point does not map directly to a single data point fed to the model.

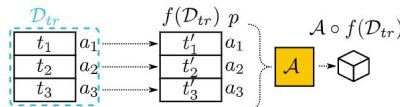
Insight

*Observe **three canonical types of pipelines** based on shape of produced provenance polynomials. Possible to develop efficient PTIME algorithms for computing the Shapley value for them.*

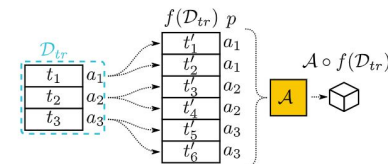
Approach

Compile provenance polynomials to Additive Decision Diagrams and use them to compute Shapley values in PTIME.

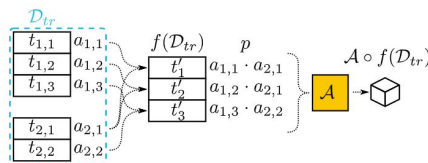
(a) Map pipeline



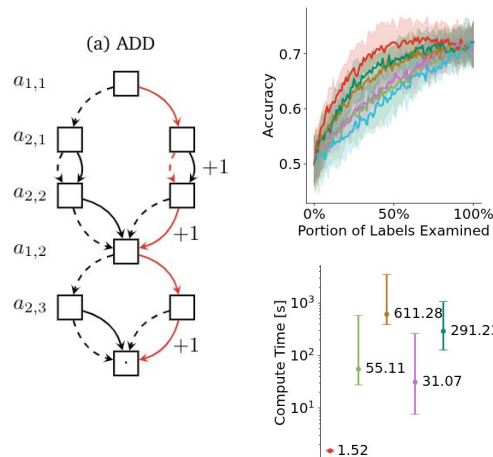
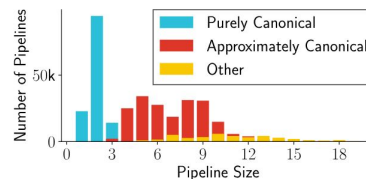
(b) Fork pipeline



(c) One-to-many join pipeline



(d) Distribution of canonical pipelines



Published as a conference paper at ICLR 2024

DATA DEBUGGING WITH SHAPLEY IMPORTANCE OVER MACHINE LEARNING PIPELINES

Bojan Karlaš*, David Dvor*, Matteo Interlandi*, Sebastian Scheller*, Wentao Wu*, Co Zhang*
 *Research University, †ETH Zurich, ‡Microsoft, §University of Amsterdam, ¶University of Chicago
 †bojan.karlas@ethz.ch, ‡david.dvor@ethz.ch

ABSTRACT

When a machine learning (ML) model exhibits poor quality (e.g., poor accuracy or fairness), the problem can often be traced back to errors in the training data. Being able to discover the data examples that are the most likely culprits is a fundamental concern that has received a lot of attention recently. One prominent way to measure “data importance” with respect to model quality is the Shapley value. Unfortunately, existing methods only focus on the ML model in isolation, without considering the broader ML pipeline for data preparation and feature extraction, which appears to be the majority of real-world ML code. This presents a major limitation to applying existing methods in practical settings. In this paper, we propose Datascope, a method for efficiently computing Shapley-based data importance over ML pipelines. We introduce several approximations that lead to dramatic improvements in terms of computational speed. Finally, our experimental evaluation demonstrates that our method is capable of data error discovery that is as effective as existing Shapley-based methods, and in some cases even outperforms them. We release our code as an open-source data debugging library available at <https://github.com/ethz-asl/datascope>.

[Karlaš ICLR '24]

Karlaš, Bojan, et al. "Data Debugging with Shapley Importance over Machine Learning Pipelines." The Twelfth International Conference on Learning Representations. 2024. [\[Paper\]](#) [\[Website\]](#) [\[Code\]](#)

Debugging Predictive Queries with Rain

[Attribution Function: Influence]

Challenge

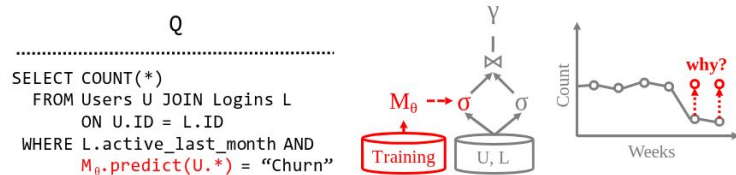
*Model inference often part of a larger predictive query.
Influence-based attribution methods must account for structure of query.*

Insight

Provenance polynomials for tracking lineage starting from training tuples all the way to predictive query outputs allows us to make the entire expression differentiable.

Approach

*User complaints on query outputs (e.g. what-if-queries) used to identify errors. Make the **entire query differentiable** using provenance polynomials and run the influence framework to identify errors in the training dataset.*



Research 15: Machine Learning for Clustering, Integration, and Search

SIGMOD '20, June 14–19, 2020, Portland, OR, USA

Complaint-driven Training Data Debugging for Query 2.0

Weiyuan Wu
Simon Fraser University
Burnaby, BC, Canada
youngw@sfu.ca

Lampros Flokas
Columbia University
New York, NY
luflokas@cs.columbia.edu

Eugene Wu
Columbia University
New York, NY
ewu@cs.columbia.edu

Jiannan Wang
Simon Fraser University
Burnaby, BC, Canada
jnwang@sfu.ca

ABSTRACT

As the need for machine learning (ML) increases rapidly across all industry sectors, there is a significant interest among commercial database providers to support "Query 2.0" which integrates model inference into SQL queries. Debugging Query 2.0 is very challenging since an incorrect query result may be caused by the bugs in training data (e.g., wrong labels, corrupted features). In response, we propose Rain, a complaint-driven training data debugging system. Rain allows users to specify complaints over the query's

Proceedings of the 2020 ACM SIGMOD International Conference on Management of Data (SIGMOD '20), June 14–19, 2020, Portland, OR, USA. ACM, New York, NY, USA, 14 pages. <https://doi.org/10.1145/339864.339866>

1 INTRODUCTION

Database researchers have long advocated the value of integrating model inference within the DBMS: data and for model inference is already in the DBMS, it brings the code

[Wu SIGMOD '20]

Wu, Weiyuan, et al. "Complaint-driven training data debugging for query 2.0." Proceedings of the 2020 ACM SIGMOD International Conference on Management of Data. 2020. [Paper](#)

ArgusEyes - Continuous Integration for ML Pipelines

Challenge

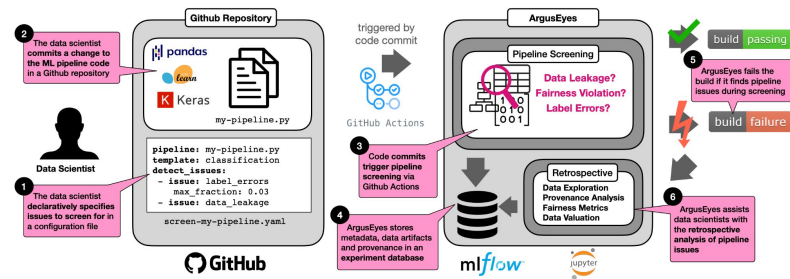
ML systems lack sophisticated testing infrastructure developed for classical software engineering. Many data-related problems only become apparent in production.

Insight

Logical query plans for ML pipelines combined with data debugging techniques enable ML-specific CI infrastructure.

Approach

Instrument, execute and screen ML pipelines for declaratively specified pipeline issues, and analyze data artifacts and their provenance to catch potential problems early before deployment to production.



[Schelter SIGMOD Demo '23]

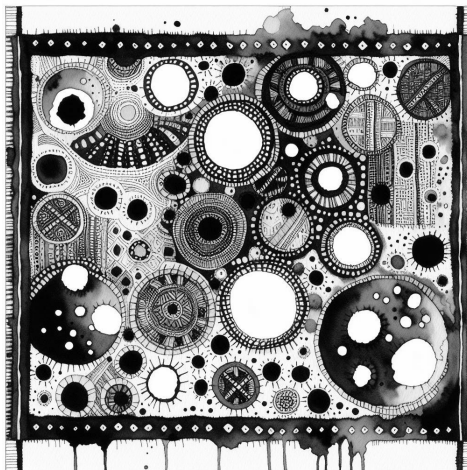
Schelter, et al.: "Proactively Screening Machine Learning Pipelines with ArgusEyes." Proceedings of the 2023 ACM SIGMOD International Conference on Management of Data (demo). 2023. [\[Paper\]](#)

Key Takeaways of Part II

- Attribution methods presented in Part I assume models are trained with source data.
- ML pipelines are complex and present many opportunities for methods development.
- Logical query plans combined with data provenance offer a powerful framework for analyzing ML pipelines.

Part III: Learning from Uncertain and Incomplete Data

Babak Salimi



The Standard ML Pipeline

Input Data



Data Cleaning

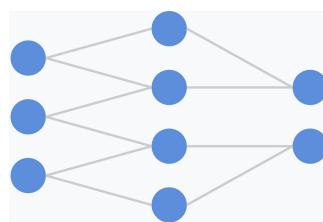


Model



Inference

ID	Age	Income	...	Loan
1	25	50K	...	5K
2	NULL	60K	...	8K
3	35	NULL	...	[10K, 12K]
...	...	...	...	...



Loan Denied:
High Risk

Loan Approved:
Low Risk

Alice



Bob



⚠ **Common Assumption:** once we “clean” the data, the pipeline consumes accurate and unbiased inputs.

✗ **Reality:** cleaning/pre-processing yields one reconstruction, driven by heuristic choices & domain assumptions → it can embed hidden bias and hide genuine uncertainty.

➡ **Key insight for Part III:** even after best-effort cleaning, *real-world data remains incomplete and uncertain*. Our models—and the theory behind them—must make that uncertainty explicit rather than ignore it.

Why “Fixing” Data Errors Is Impossible in Principle

Missing values (/)

Irrecoverable uncertainty: any imputation is just a guess; the true value is unobservable.

Unverifiable assumption: “missing at random,” parametric model of the data, etc.

[Pearl & Mohan, AAAI 2014], [Mohan, Pearl & Tian, NeurIPS 2013]

Measurement / annotation bias (sentiment, diagnoses)

Systematic distortion: recorded values can be consistently wrong.

Unverifiable assumption: symmetric, independent label-noise model.

[Pearl, UAI 2010], [Zhang & Yu, IJCAI 2015]

Why “Fixing” Data Errors Is Impossible in Principle

Selection bias & missing counterfactuals (⚠ rejected-loan applicants, excluded patients)

Unknown outcomes: whole sub-populations are never seen.

Finite-sample limits: re-weighting needs the true selection mechanism—which we can't test.

[Bareinboim, Tian & Pearl, AAAI 2014] [Cortes et al., ALT2008],
[Heckman, Econometrica 1979]

Schema / integration mismatch (⚠ inconsistent units, 🚫 fuzzy entity resolution)

Ambiguous merges: no ground-truth correspondences.

Pre-processing bias: heuristics distort original distributions; matching is probabilistic.

[Dong, Halevy & Madhavan, VLDB 2009],
[Getoor & Machanavajjhala, ACM 2012]

Challenges with Traditional Data Pipelines

Input Data



Data Cleaning



Model



Inference

Loan Denied: **High Risk** Loan Approved: **Low Risk**

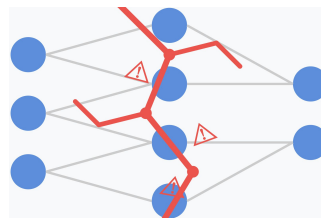
Alice



Bob



ID	Age	Income	...	Loan
1	25	50K	...	5K
2	NULL	60K	...	8K
3	35	NULL	...	[10K, 12K]
...	...	...	...	...



Generalization Failure – Models trained on “repaired” data collapse under real-world shifts.

✗ High-Stakes Mis-decisions – Hidden bias drives flawed credit, medical, and justice outcomes.

⚠ Broken Uncertainty – Bayesian & conformal intervals lose calibration when data are incomplete.

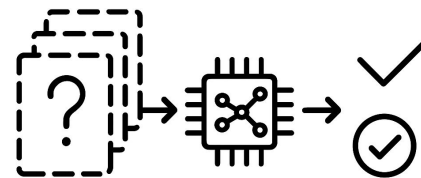
Learning from Incomplete Databases

Perfect cleaning is a myth. Even with best-effort repairs, many plausible datasets remain

Hidden uncertainty $\Rightarrow$ hidden risk. A model trained on one arbitrary repair can look accurate yet flip decisions on another equally valid repair.

Needed: an explicit uncertainty framework.

- capture what is *unknown* in the data,
- propagate that uncertainty through training,
- surface it at inference time.



Practical pay-off.

- **Robustness check:** see when all admissible models agree (safe to act).
 - **Guardrail:** abstain or seek more data when predictions diverge.
- Targeted cleaning:** focus effort on the cells that actually shrink uncertainty.

Incomplete Databases

Formalism from databases & AI to handle uncertainty by modeling all plausible data interpretations. (*Rooted in modal logic & philosophy*)

Dataset with Quality Issues

ID	Age	Income	...	Loan
1	25	50K	...	5K
2	NULL	60K	...	8K
3	35	NULL	...	[10K, 12K]
...	...	...	...	...

Q : What is the total income?

Possible Worlds Semantics

Inference:

- All repairs agree → **Certain answer**
- $\text{Range} \leq \tau \rightarrow$ **Robust interval** (e.g., [5 k – 6 k])
- $\text{Range} > \tau \rightarrow$ **Uncertain** → warn / seek more cleaning

Dataset with Quality Issues

ID	Age	Income	...	Loan
1	25	50K	...	5K
2	NULL	60K	...	8K
3	35	NULL	...	[10K, 12K]
...	...	...	...	...

Q : What is the total income?

ID	Age	Income	...	Loan
1	25	50K	...	5K
2	30	60K	...	8K
3	35	55K	...	7K
...	...	...	...	...

$$\rightarrow Q(D_1) = 6k$$

ID	Age	Income	...	Loan
1	25	50K	...	5K
2	35	60K	...	8K
3	35	60K	...	8K
...	...	...	...	...

$$\rightarrow Q(D_2) = 9k$$

ID	Age	Income	...	Loan
1	25	50K	...	5K
2	35	60K	...	8K
3	35	60K	...	8K
...	...	...	...	...

$$\rightarrow Q(D_3) = 5k$$

Range consistent
answers:
[0.5 - 0.3]

Min/Max query result across all possible database repairs.

Representing Uncertainty in Databases

C-Tables/M-Tables: Compactly represent multiple possible worlds using variables and conditions.

[Imieliński & Lipski, JACM 1984], [Sundarmurthy et al., ICDT 2017]

Probabilistic Databases: Assign probabilities to possible worlds, quantifying their likelihood.

[Suciu, Olteanu, Ré & Koch, Book 2022]

Answering queries across possible worlds is computationally expensive, often NP-hard or exponential.



ML from Possible Repairs

Inference

- All models ($h_{D_i}^*$) concur $\rightarrow$ **Certain** prediction (e.g., payout = 3 K)
- disagree $\rightarrow$ **Range** prediction (e.g., payout $\in [2\text{ K}, 4\text{ K}]$)

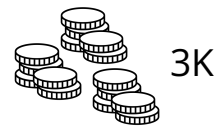
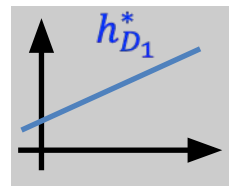


Dataset with Quality Issues

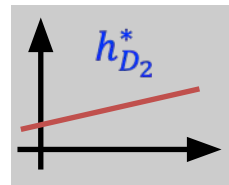
ID	Age	Income	...	Loan
1	25	50K	...	5K
2	NULL	60K	...	8K
3	35	NULL	...	[10K, 12K]
...	...	...	...	...

Machine-learning analogue of
Consistent Query Answering:
 swap the SQL query Q for a training
 routine T —e.g., gradient descent,
 decision-tree induction, SVM fitting.

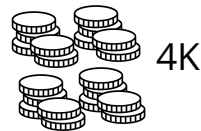
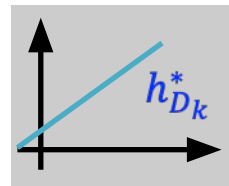
ID	Age	Income	...	Loan
1	25	50K	...	5K
2	30	60K	...	8K
3	35	55K	...	7K
...	...	...	...	...



ID	Age	Income	...	Loan
1	25	50K	...	5K
2	35	60K	...	8K
3	35	60K	...	8K
...	...	...	...	...



ID	Age	Income	...	Loan
1	25	50K	...	5K
2	35	60K	...	8K
3	35	60K	...	8K
...	...	...	...	...



KNN Classifiers over Incomplete Information

[Approach: “Certain-kNN” → returns a label only when it is guaranteed across all completions of the missing values]

Insights:

- Missing attributes can flip k-NN labels; intersecting votes across all imputations yields a *guaranteed* label.

Approach:

- Model each incomplete record as a value set (hyper-rectangle).
- Two polynomial-time tests (SS, MM) decide if a test point is “certain” without enumerating possible worlds.

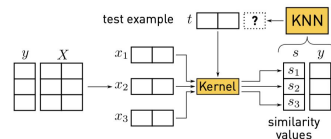
Benefits:

- 100 % precision on “certain” points – i.e., points whose prediction is certain across every imputation.
- CPClean add-on ranks the missing cells whose repair would turn “uncertain” points into certain ones, guiding targeted data cleaning.

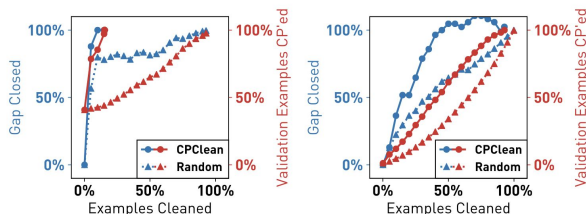
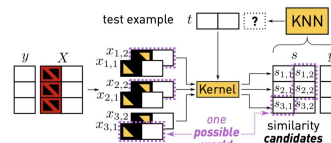
Shortcomings:

- Guarantees apply only to numeric-feature k-NN

a KNN classification over a regular training dataset



b KNN classification over a training dataset with incomplete information



Nearest Neighbor Classifiers over Incomplete Information: From Certain Answers to Certain Predictions

Bojan Karlaš¹, Peng Li², Renchi Wu¹, Nermin Merve Gürel³, Xu Chu⁴, Wentian Wu⁵, Ce Zhang⁶
¹ETH Zurich, ²Georgia Institute of Technology, ³University of Birmingham, ⁴University of Cambridge, ⁵University of Oxford, ⁶University of Warwick

ABSTRACT

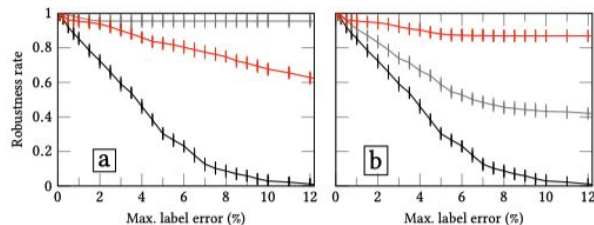
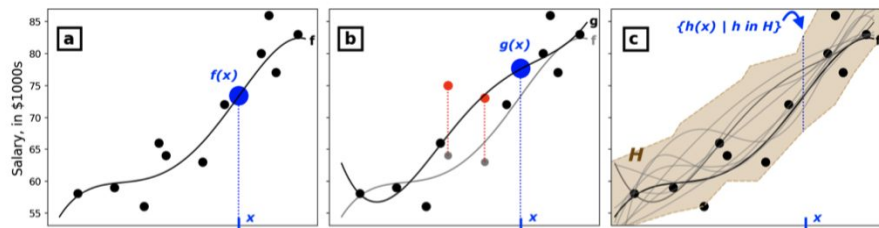
Machine learning (ML) applications have been thriving recently, largely attributed to the increasing availability of data. However, incompleteness and uncertainty information are ubiquitous in real-world datasets, and their presence in ML applications remains a challenge. In this paper, we present a formal study of this impact by extending the notion of Certain Answers (CA) to ML, which has been explored by the database research community for decades, into the field of machine learning. Specifically, we focus on classification problems and propose the notion of Certain Predictions (CP). A test data example can be *certainly predicted* (CPed) if all possible worlds induced by the imputation of data would yield the same prediction. We study two fundamental CP queries (CPQ) checking queries that determine whether a data example can be CPed, and CPQ covering queries that compute the number of classifiers that support a test prediction (i.e., label). Given that general solutions to CP queries are not computationally tractable without assumptions on the

[Karlaš VLDB '20]

Karlaš, Bojan, et al. "Nearest neighbor classifiers over incomplete information: from certain answers to certain predictions." Proceedings of the VLDB Endowment 14.3 (2020): 255-267. [\[Paper\]](#)

The Dataset Multiplicity Problem

[Approach: bound model risk across every dataset consistent with the errors]



Insights:

- Introduces a risk interval: the tightest possible lower/upper bound on test error that any admissible dataset can induce for a fixed linear model.

Approach:

- Derive closed-form formulas for the worst- and best-case hinge / logistic loss of any linear classifier under those rules, avoiding enumeration.

Benefits:

- Gives practitioners a numeric certificate of how much reported accuracy can deteriorate.

Shortcomings:

- Theory currently limited to linear models and label-noise rules; deep nets need looser convex relaxations.

The Dataset Multiplicity Problem: How Unreliable Data Impacts Predictions

Anna Meyer
ameyer@wisc.edu
University of Wisconsin - Madison
Madison, USA

Awa Albarghouthi
awalg@wisc.edu
University of Wisconsin - Madison
Madison, USA

Loris D'Antoni
loris@wisc.edu
University of Wisconsin - Madison
Madison, USA

ABSTRACT
We introduce dataset multiplicity, a way to study how inaccuracies, uncertainty, and social bias in training datasets impact test-time predictions. The dataset multiplicity framework asks a counterfactual question of what the set of real-world models (and associated test-time predictions) would be if we could somehow access all hypothetical, unlabeled versions of the dataset. We discuss how to use this framework to encapsulate various sources of uncertainty in datasets' faithfulness, including systemic social bias, data collection practices, and noisy labels or features. We show how to exactly analyze the impact of dataset multiplicity for a specific model architecture and type of uncertainty: linear models with label errors. Our empirical analysis shows that real-world datasets, under reasonable assumptions, contain many test samples whose predictions are affected by dataset multiplicity. Furthermore, the choice of domain-specific dataset multiplicity definitions determines what samples are affected, and whether different demographic groups are disparately impacted. Finally, we discuss implications of dataset multiplicity for machine learning practice and research, including considerations for when model outcomes should not be trusted.

1 INTRODUCTION
Datasets that power machine learning algorithms are supposed to be accurate and fully representative of the world, but in practice, this level of precision and representativeness is impossible [27, 44]. Dataset quality issues— which we use as a catch-all term for both errors and misrepresentations— due to sampling bias [10], human errors in label or feature transcription [29, 43], and sometimes deliberate poisoning attacks [3, 52]. Datasets can also reflect undesirable societal inequalities. But more broadly, datasets never reflect objective truths because the worldview of their creators is infused in the data collection and preprocessing [27, 44, 46]. Additionally, seemingly-trivial decisions in the data collection or annotation process influence exactly what data is included, or not [44, 45]. In psychology, these minute decisions have been termed “researcher degrees of freedom,” i.e., choices that can inadvertently influence conclusions that one ultimately draws from the data analysis [5]. In this paper, we study how unreliable data of all kinds impacts the predictions of the models trained on such data and frame this analysis as a “multiplicity problem.”

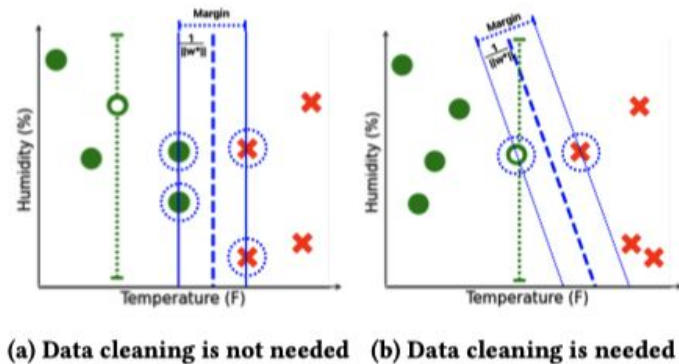
Multiplicity occurs when there are multiple combinations for

[Meyer FAccT'23]

Meyer, A. P.; Albarghouthi, A.; D'Antoni, L. “The Dataset Multiplicity Problem: How Unreliable Data Impacts Predictions. [\[Paper\]](#)

Certain & Approximately Certain Models for Statistical Learning

[Approach: Fast “certainty test” that lets you skip imputation whenever the missing cells don’t affect the optimum]



Certain and Approximately Certain Models for Statistical Learning

Cheng Zhen
Oregon State University
Corvallis, Oregon
zhen@oregonstate.edu

Arash Termechi
Oregon State University
Corvallis, Oregon
termechi@oregonstate.edu

Nischal Aryal
Oregon State University
Corvallis, Oregon
aryal@oregonstate.edu

Amandeep Singh Chabada
Oregon State University
Corvallis, Oregon
chabada@oregonstate.edu

ABSTRACT

Real-world data is often incomplete and contains missing values. To train accurate models over real-world datasets, users need to spend a substantial amount of time and resources imputing and finding proper values for missing data items. In this paper, we demonstrate that it is possible to learn accurate models directly from data with missing values for certain training data and target models. We propose a unified approach for checking the necessity of data imputation to learn accurate models across various widely-used machine learning paradigms. We build efficient algorithms with theoretical guarantees to check this necessity and return accurate models in cases where imputation is unnecessary. Our extensive experiments indicate that our proposed algorithms significantly reduce the amount of time and effort needed for data imputation without incurring considerable computational overhead.

To address the problem of training over incomplete data, users usually replace each missing data item with a value, i.e., data imputation, and train their models over the resulting repaired data. To repair incomplete data, users must figure out the mechanisms and causes of data missingness, e.g., completely at random or based on observed values of some features [26]. Based on this mechanism, they build a (statistical) model for missing data and replace the missing values with some measurements derived from the model, e.g., mean. Users may also leverage a variety of ML models to repair missing values, e.g., the least or linear regression [21]. Researchers have shown that the desired imputation method may vary depending on the downstream ML task [24]. Hence, it is often challenging to find a model of data missingness that results in an accurate ML model for a downstream task [26]. The aforementioned steps of finding a missingness mechanism, constructing an accurate

Insights:

- Not every example with missing values requires cleaning.
- If the missing cells lie in directions that do **not** change the model’s optimum, we can train directly on the incomplete data—with full guarantee.

Approach:

- Provide **fast algebraic tests** (no world enumeration) that decide certainty for linear regression, linear SVM, and two kernel SVMs. When tests pass → output the **certain model** (exactly optimal).
- When tests fail → compute an ϵ -**certain model** whose loss is within ϵ of the global optimum.

Benefits:

- **Skips imputation** for datasets that pass the test, saving cleaning effort and avoiding imputation bias.
- Same code works across several common model families.

Shortcomings:

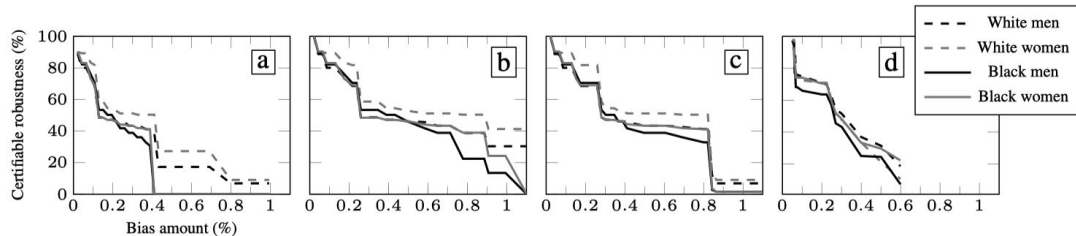
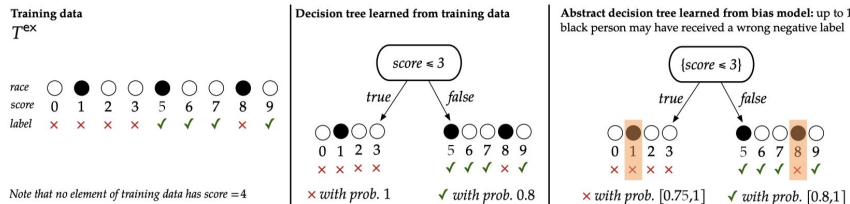
- Certainty rarely holds under heavy missingness. Guarantees limited to the studied linear & kernel models; deep nets need other methods.

[Zhen SIGMOD’24]

Zhen, C. et al. “Certain and Approximately Certain Models for Statistical Learning. [\[Paper\]](#)

Certifying Robustness to Programmable Data Bias in Decision Trees

[Approach — ProgBiasCert: encode “tree + bias program” in SMT to prove the label never flips]



Insights:

- Treat data bias as a **user-written program** (e.g., $age \pm 2$, $race\ swap$, $income \times 0.9-1.1$).
- A tree is *robust* if its prediction is invariant under **all** transformations allowed by that program.

Approach:

- Translate each path of the decision tree **and** the bias constraints into a single SMT formula.

Benefits:

- Exact guarantees—no sampling; works with real & categorical features and generates independently checkable proofs

Shortcomings:

- Does not yet handle ensembles or probabilistic bias distributions.

[Meyer NeurIPS'21]

Zhen, C.; Aryal, N.; Termehchy, A.; Chabada, A. S. “Certifying Robustness to Programmable Data Bias in Decision Trees.” [\[Paper\]](#)

Certifying Robustness to Programmable Data Bias in Decision Trees

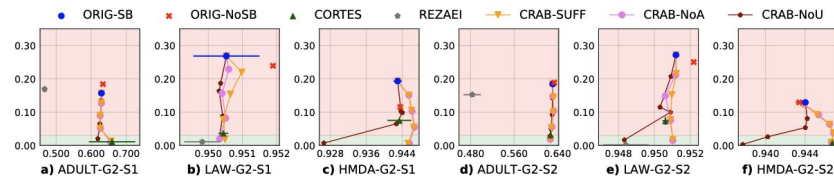
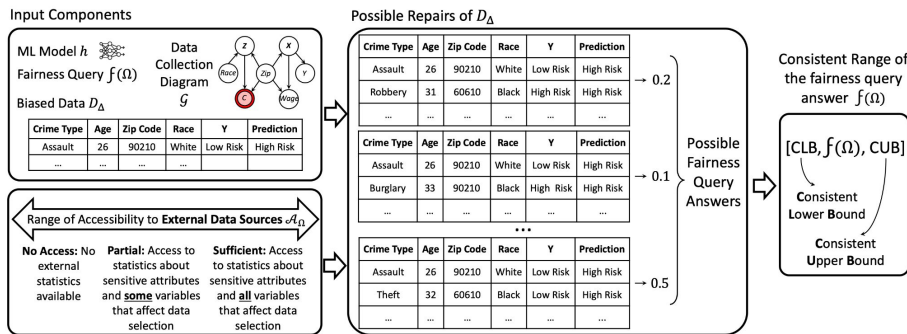
Anna P. Meyer, Aws Albarghouti, and Loris D'Antoni
Department of Computer Sciences
University of Wisconsin-Madison
Madison, WI 53706
(annameyer, aws, loris)@cs.wisc.edu

Abstract

Datasets can be biased due to societal inequalities, human biases, under-representation of minorities, etc. Our goal is to certify that models produced by a learning algorithm are *pointwise-robust* to potential dataset biases. This is a challenging problem: It entails learning models for a large, or even infinite, number of datasets, ensuring that they all produce the same prediction. We focus on decision-tree learning due to the interpretable nature of the models. Our approach allows programmatically specifying bias models across a variety of dimensions (e.g., missing data for minorities), composing types of bias, and targeting bias towards a specific group. To certify robustness, we use a novel symbolic technique to evaluate a decision-tree learner on a large, or infinite, number of datasets, certifying that each and every dataset produces the same prediction for a specific test point. We evaluate our approach on datasets that are commonly used in the fairness literature, and demonstrate our approach's viability on a range of bias models.

Consistent Range Approximation for Fair Predictive Modeling

[Approach: Fair-aware prediction ranges: bound each score so it stays fair under every repair of noisy / missing sensitive attributes]



Insights:

- With selection bias we **don't know** the target-population fairness.
- Treat fairness evaluation as a **query over incomplete data**; answer with a *range* that is guaranteed to contain the truth.

Approach:

- Derive a closed-form range for fairness aggregates.
- Train a classifier that minimises risk while keeping the worst-case value inside the acceptable fairness range.

Benefits:

- Certifies fairness without unbiased samples; needs only the biased data + background knowledge.

Shortcomings:

- Relies on correct causal diagram; ranges may be wide if knowledge is weak.

Consistent Range Approximation for Fair Predictive Modeling

Jiongli Zhu
University of California,
San Diego
jjz143@ucsd.edu

Sainyam Gaihotra
Cornell University
sgc@cs.cornell.edu

Nazanin Sabri
University of California,
San Diego
nsabri@ucsd.edu

Babak Salimi
University of California,
San Diego
bsalimi@ucsd.edu

ABSTRACT

This paper proposes a novel framework for certifying the fairness of predictive models trained on biased data. It starts from query answering for incomplete and inconsistent databases to formulate the problem of consistent range approximation (CRA) of fairness queries for a predictive model on a target population. The framework employs background knowledge of the data collection process and biased data, working with or without limited statistics about the target population, to compute a range of answers for fairness queries. Using CRA, the framework builds predictive models that are certifiably fair on the target population, regardless of the availability of external data during training. The framework's efficacy is demonstrated through evaluations on real data, showing substantial improvement over existing state-of-the-art methods.

results, deploying these models in the target population may lead to unfair and inaccurate predictions [6, 31, 33, 37, 48].

A significant issue in predictive modeling is **selection bias**, resulting from training data selection based on specific attributes, which creates unrepresentative datasets. This problem is prevalent in sensitive areas like predictive policing, healthcare, and finance, attributed to data collection costs, historical discrimination, and biases [13, 20, 32, 40]. For example, in predictive policing, the data is biased as it is gathered exclusively from police interactions, which are influenced by the sociocultural traits of the officers [28, 43]. Similarly, in healthcare, selection bias occurs when data is relied upon from individuals who are hospitalized or have tested positive, leading to disproportionate effects on racial, ethnic, and gender minorities due to barriers in healthcare access [2, 16, 45, 85].

Example 1.1. Consider the dataset in Table 1, which represents

[Zhu VLDB '23]

Consistent Range Approximation for Fair Predictive Modeling. [Paper]

Learning from Uncertain Data: From Possible Worlds to Possible Models

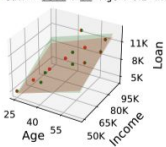
[Approach: Abstract interpretation +
zonotopes: train once on a single convex
polytope that encodes every possible repair

Training Data

Age	Income	Savings
25	50K	5K
35	60K	6K
45	90K	12K
50	NULL	[8K,9K]
55	75K	9K
60	85K	10K
65	80K	13K

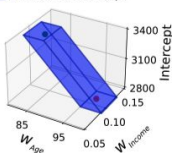
(a)

● Possible World 1 ● Possible World 2
 ■ Sav. = -2855 + 98 × Age + 0.1 × Inc.
 ■ Sav. = -3361 + 84 × Age + 0.1 × Inc.



(b)

● Possible Model 1 ● Possible Model 2
 ■ Abstract Model Zonotope



(c)

Possible Model 1 Predictions

Age	Income	Savings
20	90K	8.1K
70	50K	9K

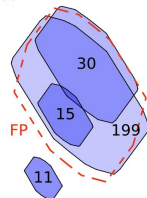
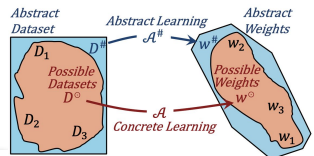
Possible Model 2 Predictions

Age	Income	Savings
20	90K	7.3K
70	50K	11.5K

Zorro Prediction Ranges

Age	Income	Savings
20	90K	[7K, 8.2K]
70	50K	[8.5K, 12K]

(d)



Learning from Uncertain Data: From Possible Worlds to Possible Models

Jiongli Zhu¹ Su Feng² Boris Glavic³ Babak Salimi¹

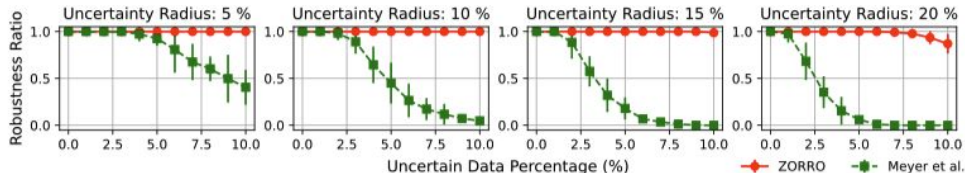
¹University of California, San Diego ²Nanjing Tech University ³University of Illinois, Chicago

Abstract

We introduce an efficient method for learning linear models from uncertain data, where uncertainty is represented as a set of possible variations in the data, leading to predictive multiplicity. Our approach leverages abstract interpretation and zonotopes, a type of convex polytope, to compactly represent these dataset variations, enabling the symbolic execution of gradient descent on all possible worlds simultaneously. We develop techniques to ensure that this process converges to a fixed point and derive closed-form solutions for this fixed point. Our method provides sound over-approximations of all possible optimal models and viable prediction ranges. We demonstrate the effectiveness of our approach through theoretical and empirical analysis, highlighting its potential to reason about model and prediction uncertainty due to data quality issues in training data.

[Zhu NeurIPS'24]

Zhu, J.; Feng, S.; Glavic, B.; Salimi, B. "Learning from Uncertain Data: From Possible Worlds to Possible Models. [\[Paper\]](#)



(a) Robustness verification with uncertain labels (MPG data).

Insights:

- Zonotope = all repairs in a compact affine form.
- Training on the zonotope gives one weight-box that subsumes every per-repair model.

Approach:

- Map each uncertain record to an affine form; the full dataset becomes **one zonotope**. Run gradient descent **symbolically**. Output is a convex box of model weights; any concrete repair yields weights inside this box.

Benefits:

- **Guaranteed intervals for weights & predictions**—true model always inside.

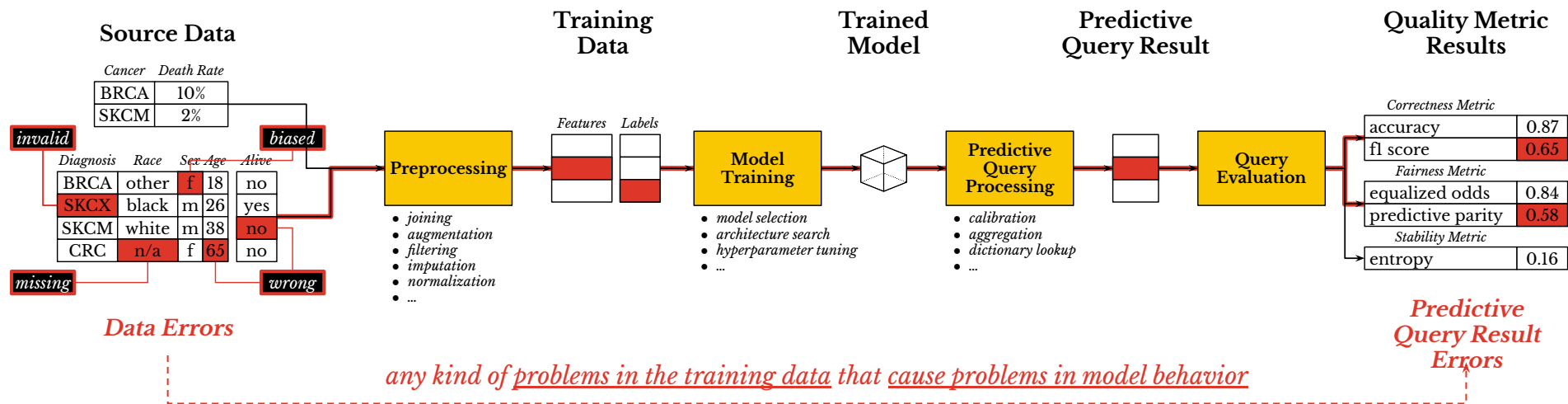
Shortcomings:

- Supports **linear models only**.

Key Takeaways of Part III

- Data cleaning can not only be costly, but it can also introduce the risk of reducing data quality.
- Possible worlds semantics is a fundamental concept for analyzing uncertainty over possible data repairs.
- Since the number of possible worlds is exponential, approximations and algorithmic tricks are needed to make reasoning about uncertainty tractable.

Conclusion: How should we navigate data errors?



Error Detection:
Compute Data Importance

ML Pipeline Debugging:
Leverage Data Provenance

Learning from Uncertain Data:
Apply Possible Worlds Semantics