

`/home/philippe/Documents/courses/img/mba_tech_logo.jpg`

ETHICS IN DATA SCIENCE

DATA SCIENCE IN PRACTICE

DR. PHILIPPE J.S. DE BROUWER

ALL RIGHTS RESERVED BY THE AUTHOR
prepared for: AGH University of Krakow

SEPTEMBER 28, 2026

Table of Contents

Contents

Contents	2
1 Bias in Data and Mathematical Models	2
1.1 Basic Definitions	2
1.2 Algorithmic Fairness	3
1.3 The Confusion Matrix and Related Measures	3
1.4 Measures for Algorithmic Fairness	4
2 Origin and Types of Bias	5
3 Conclusions	6
4 Case Study	7

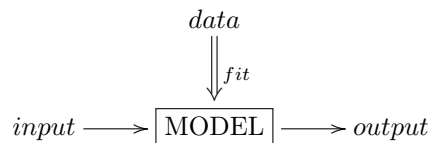
1 Bias in Data and Mathematical Models

1.1 Basic Definitions

What is a Model?

Definition 1 (Mathematical Model). A mathematical model is a description of a system using mathematical concepts and language.

Use of a model in banking is fact-based (data is used to fit the model):



What is Bias in Models?

Definition 2 (Bias in Mathematical Models). Model Bias (also Algorithmic Bias) refers to the systematic and repeatable error in a Model that creates outcomes that are statistically at odds with the reality of the population whose behaviour it is supposed to reflect.

1.2 Algorithmic Fairness

Fairness Assumptions

Assume S the protected feature with an under-privileged group S_u and a privileged group S_p , and predicted outcomes $\hat{Y}$.

- A. **Independence:** the probability of being in S_p or S_u has nothing to do with $\hat{Y}$ – $\hat{Y}$ is independent of S .
- B. **Separation:** The predicted probability of having a favourable outcome has nothing to do with membership of the protected feature – $\hat{Y}$ is independent of S , given Y .
- C. **Sufficiency:** the prediction should not depend on the protected group – Y is independent of S , given $\hat{Y}$.

1.3 The Confusion Matrix and Related Measures

A simple and clear way to show what the model does is a “confusion matrix”: simply show how much of the observations are classified correctly by the model.

Useful Concepts for the Confusion Matrix

The following are useful measures for how good a classification model fits its data:

- *Accuracy:* The proportion of predictions that were correctly identified.
- *Precision* (or positive predictive value): The proportion of positive cases that correct.
- *Negative predictive value:* The proportion of negative cases that were correctly identified.
- *Sensitivity* or Recall: The proportion of actual positive cases which are correctly identified.
- *Specificity:* The proportion of actual negative cases which are correctly identified.

Let us use the following definitions:

- Objective concepts (depends only on the data):
 - P : The number of positive observations ($y = 1$)
 - N : The number of negative observations ($y = 0$)
- Model dependent definitions:
 - True positive (TP) the positive observations ($y = 1$) that are by the model correctly classified as positive;

	Observed pos.	Observed neg.	
Pred. pos.	TP	FP	Pos.pred.val = $\frac{TP}{TP+FP}$
Pred. neg.	FN	TN	Neg.pred.val = $\frac{TN}{FN+TN}$
	Sensitivity	Specificity	Accuracy
	$= \frac{TP}{TP+FN}$	$= \frac{TN}{FP+TN}$	$= \frac{TP+TN}{TP+FN+FP+TN}$
	$= \frac{TP}{TP+FN}$	$= \frac{TN}{FP+TN}$	$= \frac{TP+TN}{TP+FN+FP+TN}$

Table 1: The confusion matrix, where “pred.” refers to the predictions made by the model, “pred.” stands for “predicted,” and the words “positive” and “negative” are shortened to three letters.

- False positive (FP) the negative observations ($y = 0$) that are by the model incorrectly classified as positive – this is a false alarm (Type I error);
- True negative (TN) the negative observations ($y = 0$) that are by the model correctly classified as negative;
- False negative (FN) the positive observations ($y = 1$) that are by the model incorrectly classified as negative – miss (Type II error).

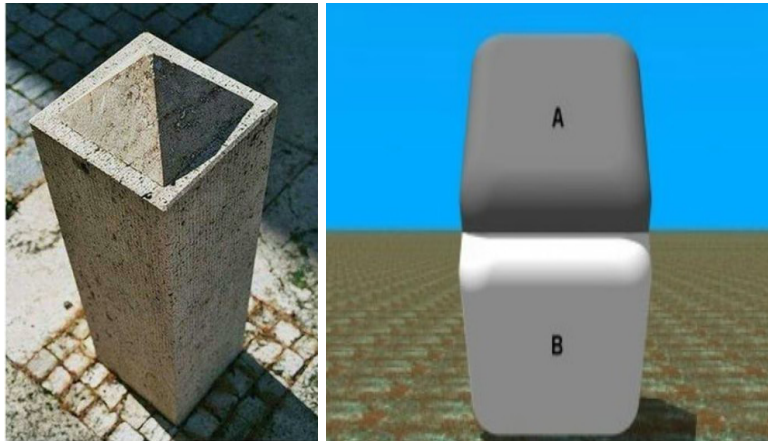
The Definition of the Confusion Matrix

1.4 Measures for Algorithmic Fairness

Fairness Metrics

- Equal Opportunity:** $FNR_p = FNR_u$ (equivalent with $TPR_p = TPR_u$)
- Predictive Equality:** $FPR_p = FPR_u$ (equivalent with $TNR_p = TNR_u$)
- Equalised Odds:** both previous
- Predictive Parity** or outcome test: $Prec_p = Prec_u$ (equal positive predictive value or precision ($Prec = \frac{TP}{TP+FP}$))
- Demographic Parity:** membership of S has no correlation with favourable outcome ($\hat{Y}$).

Visual Biases are systematic miss-interpretations



2 Origin and Types of Bias

General Causes of Bias in Data

- **Confirmation Bias:** bias governs what we search for and we find what we search for (e.g. belief preservation via social media)
- **Selection Bias:** sample is not representative
- **Historical Bias:** socio-cultural prejudices are reflected in data (e.g. Google image search)
- **Survival Bias:** the winner takes it all and losers are forgotten (e.g. hedge fund performance statistics)
- **Availability Bias:** violence has systematically decreased over the millennia, however we might think of modern times to be more violent. (focus on through the cycle data, different perspectives)
- **Outlier Bias:** thinking about a startup one things of Google, Amazon, Facebook, etc. These are the outliers. (use median instead of average, identify outliers, etc.)

Process Related Bias

- **Reporting bias:** selective reporting of some data: citation bias, language bias (ignore reports in other languages), duplicate publication bias (copied data found twice), location bias (some studies are hard to find), publication bias (secret or not popular data is hard to find), outcome reporting bias (e.g. company does not report with much bravo when results are bad), time lag bias
- **Automation bias:** we prefer automated systems to provide data

- **Selection bias:** data is not representative (e.g. sampling bias, convergence bias (only people who got a loan are in our database), participation bias, etc.)
- **Over-generalisation bias:** no black swans in the data
- **Group Attribution bias:** generalisation of stereotypes, in-group bias (preference for members in the group), out-group bias (stereotype for other groups), etc.
- **Implicit bias:** search for conforming information

What causes Bias in Models?

Bias can appear due to

- A. biases in the data
 - less representative group
 - the data is not a good representation of the reality (e.g. survivor bias, previous societal bias or limitations, etc.)
- B. the data processing pipeline (extraction, cleaning, transformation, binning)
- C. the model development process (including selection of algorithm, variables, etc.)
- D. the particulars of model implementation

3 Conclusions

Fairness is a matter of perspective

We do have an innate sense for what is fair and what not, however that sense is heavily biased towards ourselves. For example, we don't see it as a moral issue to build a house and destroy the home of a squirrel in the process. The squirrel, however, is only separated from us by 250 million years of separate evolution. From a society that is a million years ahead of it might seem obvious to simply make earth their home and get rid of us.

If that seems unfair, turn it around. Would it be fair to colonise a planet inhabited by squirrel like mammals and make it our home?



Figure 1: Fairness is a matter of perspective

Bias is learnt



4 Case Study

Case study on addressing bias: car insurance



References

- Agarwal, Sray and Shashin Mishra (2021). *Responsible AI: Implementing Ethical and Unbiased Algorithms*. Springer.
- Coeckelbergh, Mark (2020). *AI ethics*. Mit Press.
- De Brouwer, Philippe J. S. (2020). *The Big R-Book: From Data Science to Learning Machines and Big Data*. New York: John Wiley & Sons, Ltd. ISBN: 978-1-119-63272-6.
- Dubber, Markus Dirk, Frank Pasquale, and Sunit Das (2020). *The Oxford handbook of ethics of AI*. Oxford Handbooks.
- Hardt, Moritz, Solon Barocas, and Arvind Narayanan (2018). 'Fairness in Machine Learning: Limitations and Opportunities'. In: *Solon Barocas Moritz Hardt and Arvind Narayanan*.