

Author: A. Erdelsky | Supervisor: S.R. Bongers | Responsible Professor: J.H. Krijthe

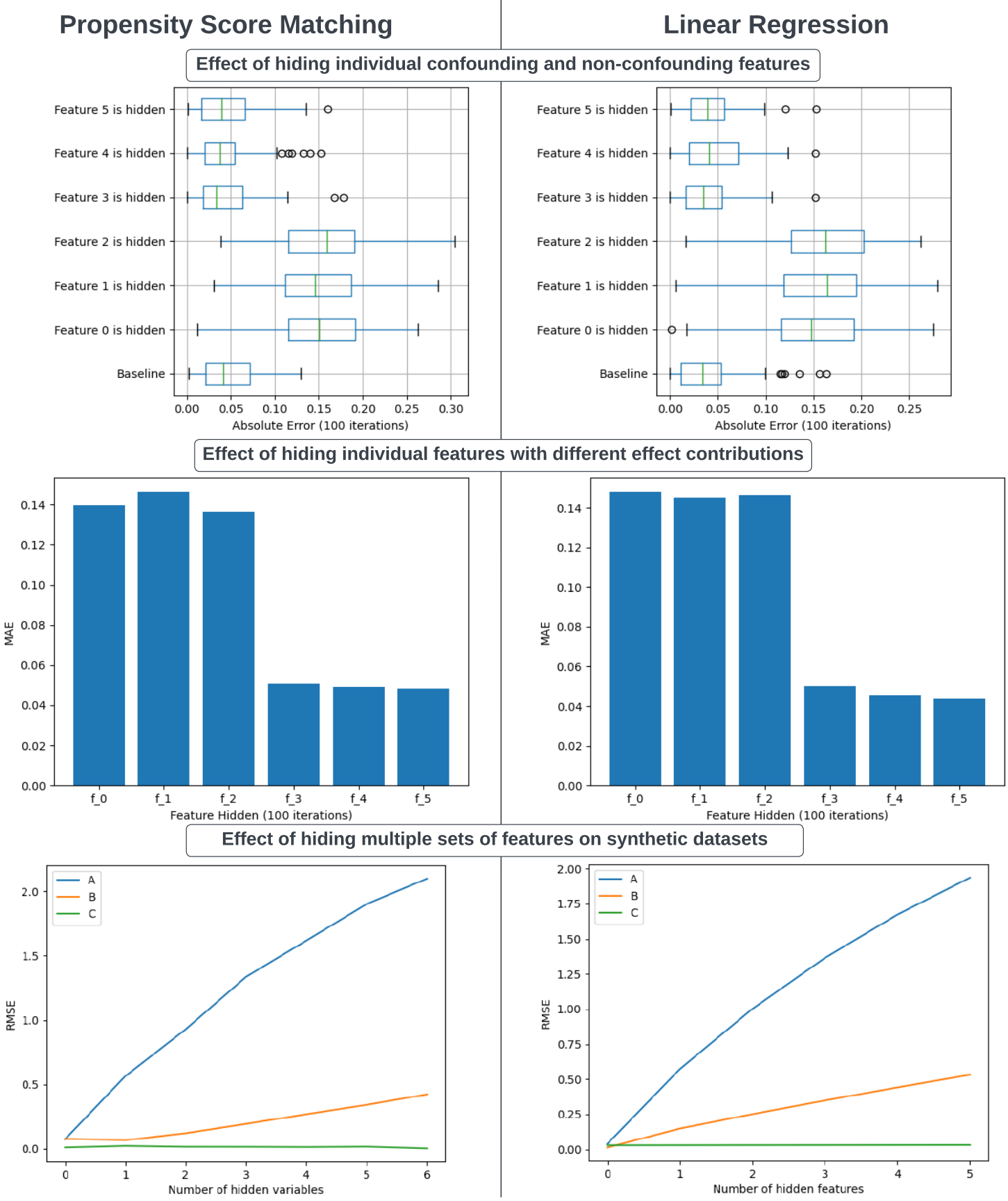
1 Background

- **Propensity Score Matching** is a causal machine learning algorithm that is used for the unbiased estimation of the **ATE**, the average treatment effect.
- PSM operates ideally only under specific conditions, the main assumption being **unconfoundedness** [1].
- Unconfoundedness of a dataset means that all variables that affect treatment and outcome have been measured. These variable are known as **confounding variables, covariates** or **features** [1].
- Propensity Score Matching entails **forming matched sets of treated and untreated subjects** who share a similar value of the **propensity score** [2].
- The **propensity score** is the probability of getting treatment based on observed confounding variables [2].

2 Methodology

- What is the effect of unconfoundedness on the performance of Propensity Score Matching ?
- Breaking the unconfoundedness assumption should negatively impact the ATE performance of PSM.
- Run the algorithm on data that **upholds the unconfoundedness condition**.
- Compare these results with measurements obtained from running the algorithm on data with **confounding features with varying contribution to other variable values and hiding these features individually or in progressively higher numbers**.
- These results are also then **compared to Linear Regression**, a generic machine learning algorithm, for the sake of comparison.

3 Results



4 Conclusions

- When running PSM with missing features, **the severity of the error in the output is dependent on how that feature affects all other variables present as well as how many of these features are missing**.
- If the **hidden feature does not influence any other variable**, the output of PSM will remain the **same as when every feature is observed** by the method.
- When hiding variables that **only contribute to the main effect, treatment effect or treatment propensity**, PSM performs with the **same error** no matter which of the three effects the hidden feature affects.
- In all experimental scenarios used in this work, **PSM performed nearly indently to Linear Regression** and did not seem to offer any advantages over the latter in these specific situations.

5 Limitations

- This research would benefit from:
- **Using real-world causal inference datasets** in all its experiments.
 - **Examining the ATT output of PSM** and see if the method behaves differently.
 - **Find experimental scenarios and data that demonstrates the strengths of PSM** over generic machine learning algorithms.

6 References

[1] Ruocheng Guo, Lu Cheng, Jundong Li, P Richard Hahn, and Huan Liu. A survey of learning causality with data: Problems and methods. *ACM Computing Surveys (CSUR)*, 53(4):1–37, 2020.

[2] Peter C Austin. An introduction to propensity score methods for reducing the effects of confounding in observational studies. *Multivariate behavioral research*, 46(3):399–424, 2011.