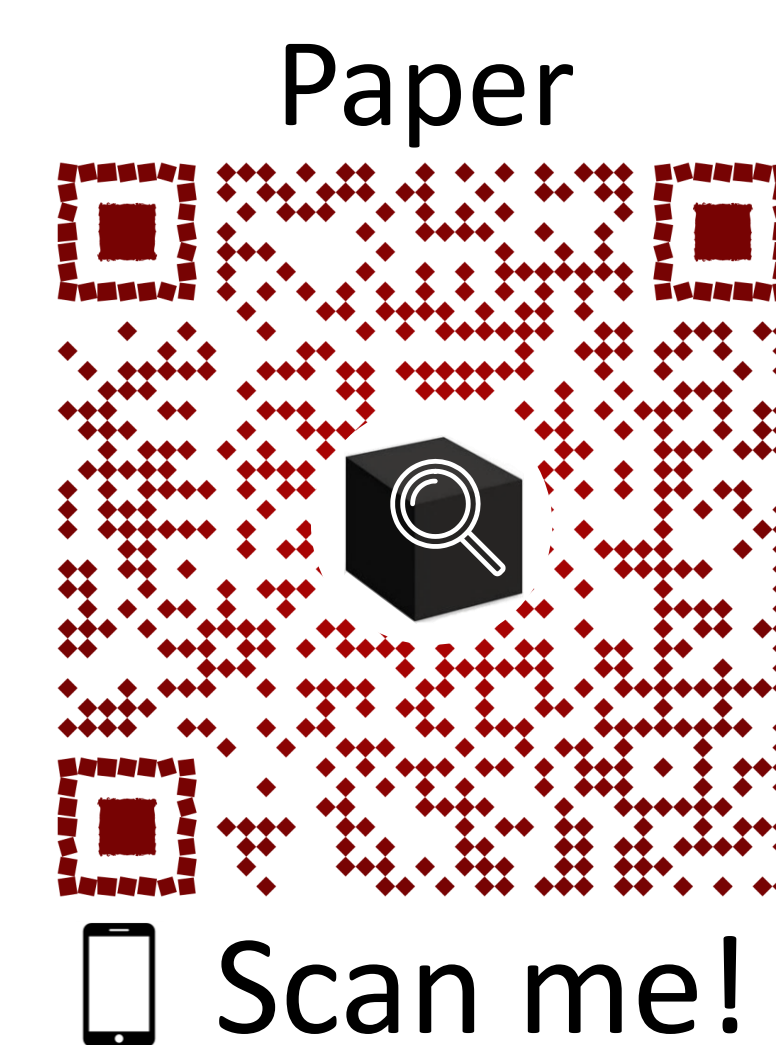


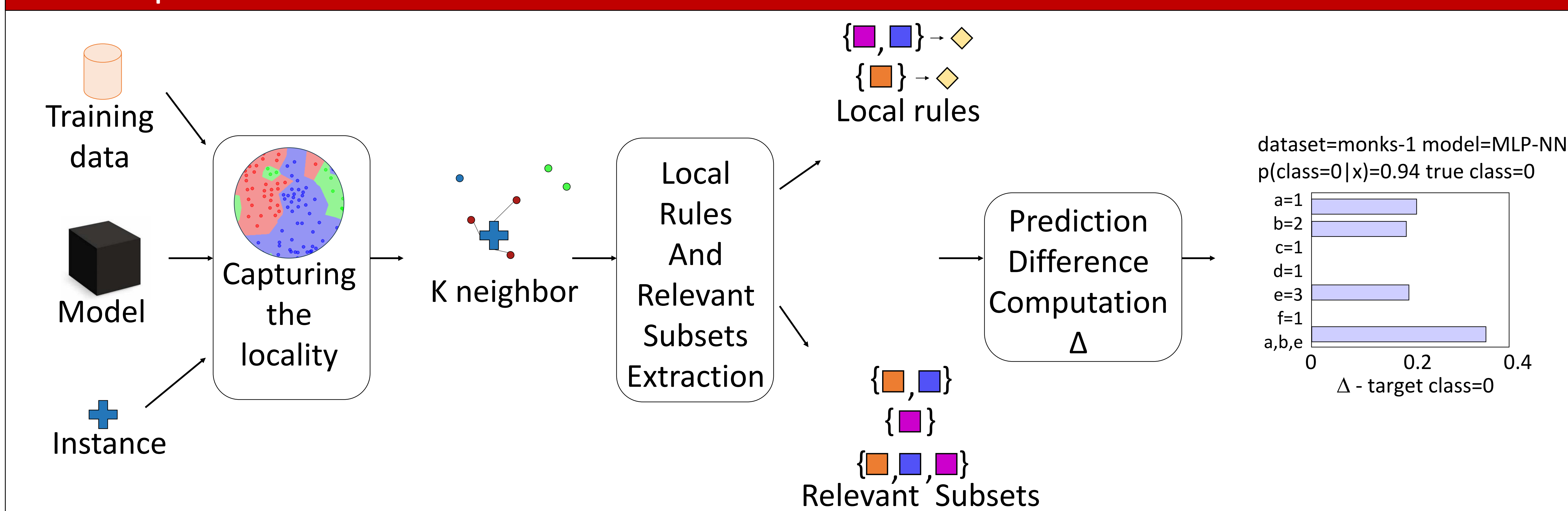
# Enhancing Interpretability of Black Box Models by means of Local Rules

Eliana Pastor and Elena Baralis  
Politecnico di Torino, Italy

✉ eliana.pastor@polito.it



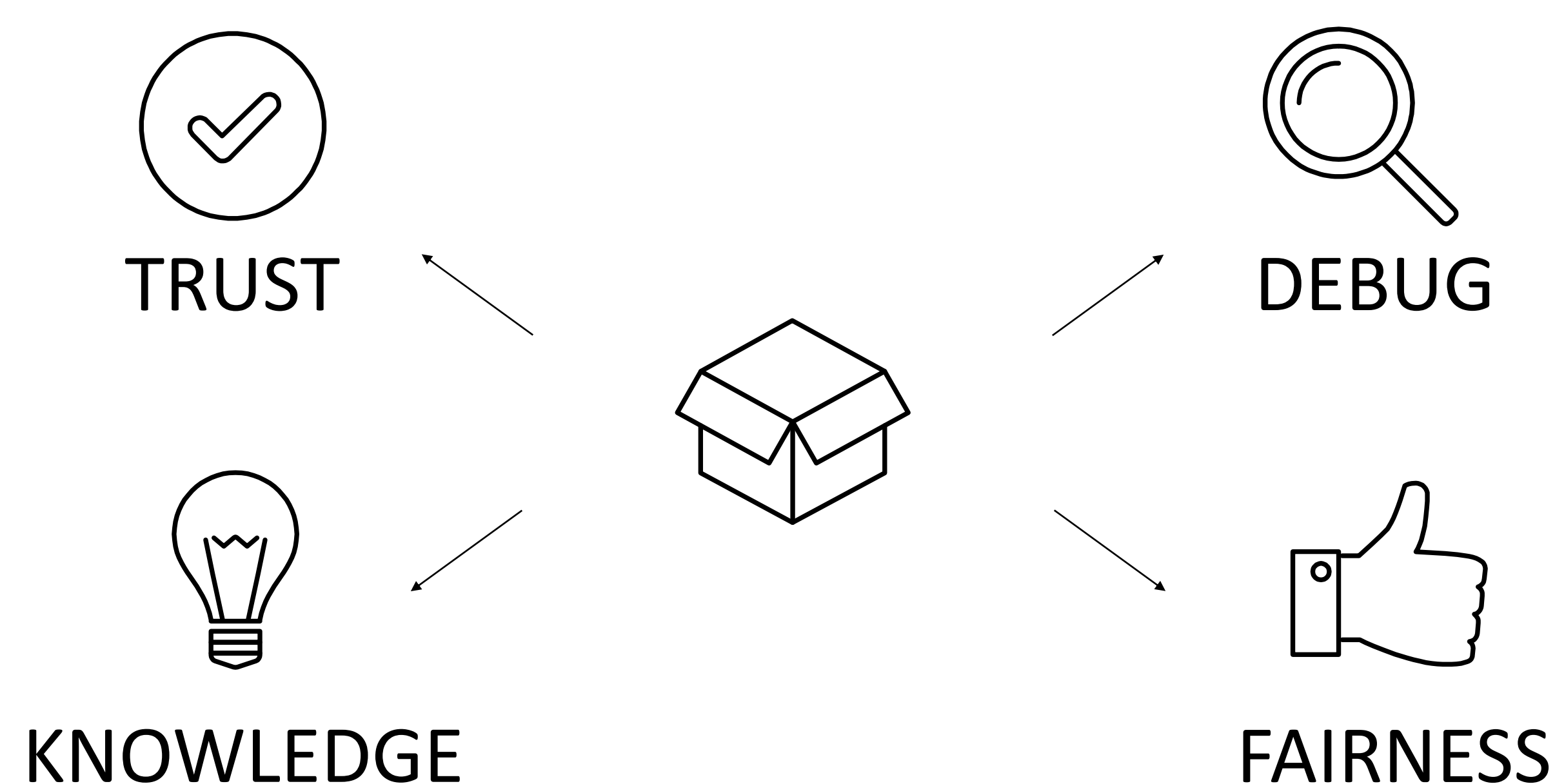
## LACE explanation method



## Enhancing interpretability

Black box models  $\rightarrow$  hide from users the reasons behind predictions

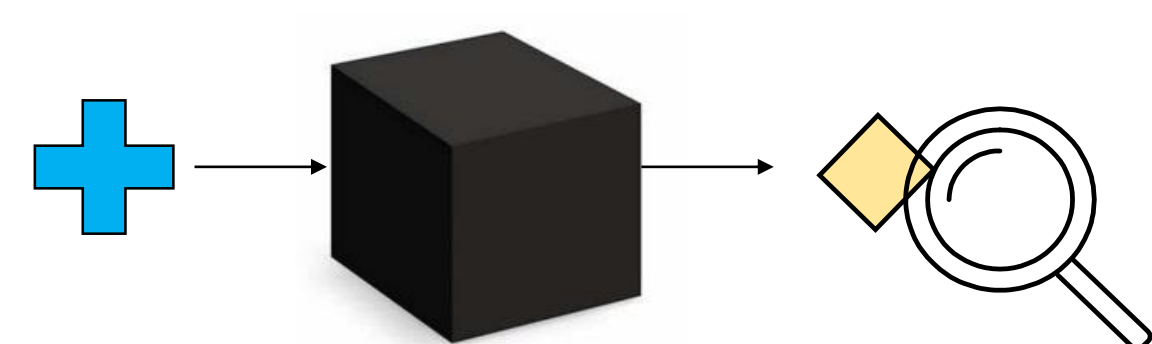
Why explainable AI?



Demand from institutions  $\rightarrow$  GDPR

- “meaningful information about the logic involved”
- “to ensure fair and transparent processing”

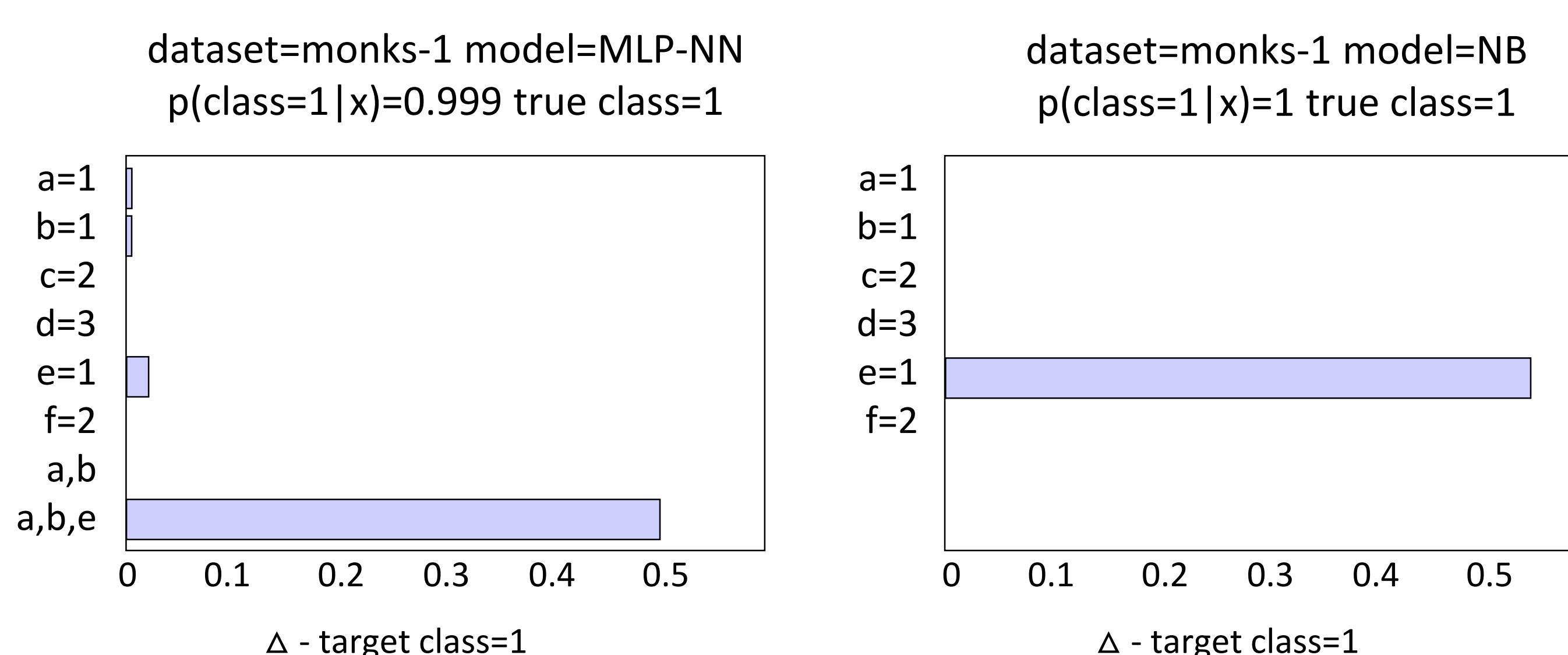
## LACE Explanations



- Qualitative insight  $\rightarrow$  **local rules**  
Rules learned in the locality of the prediction
- Quantitative insight  $\rightarrow$  **prediction difference**  
Prediction change when one or more attribute values are omitted.

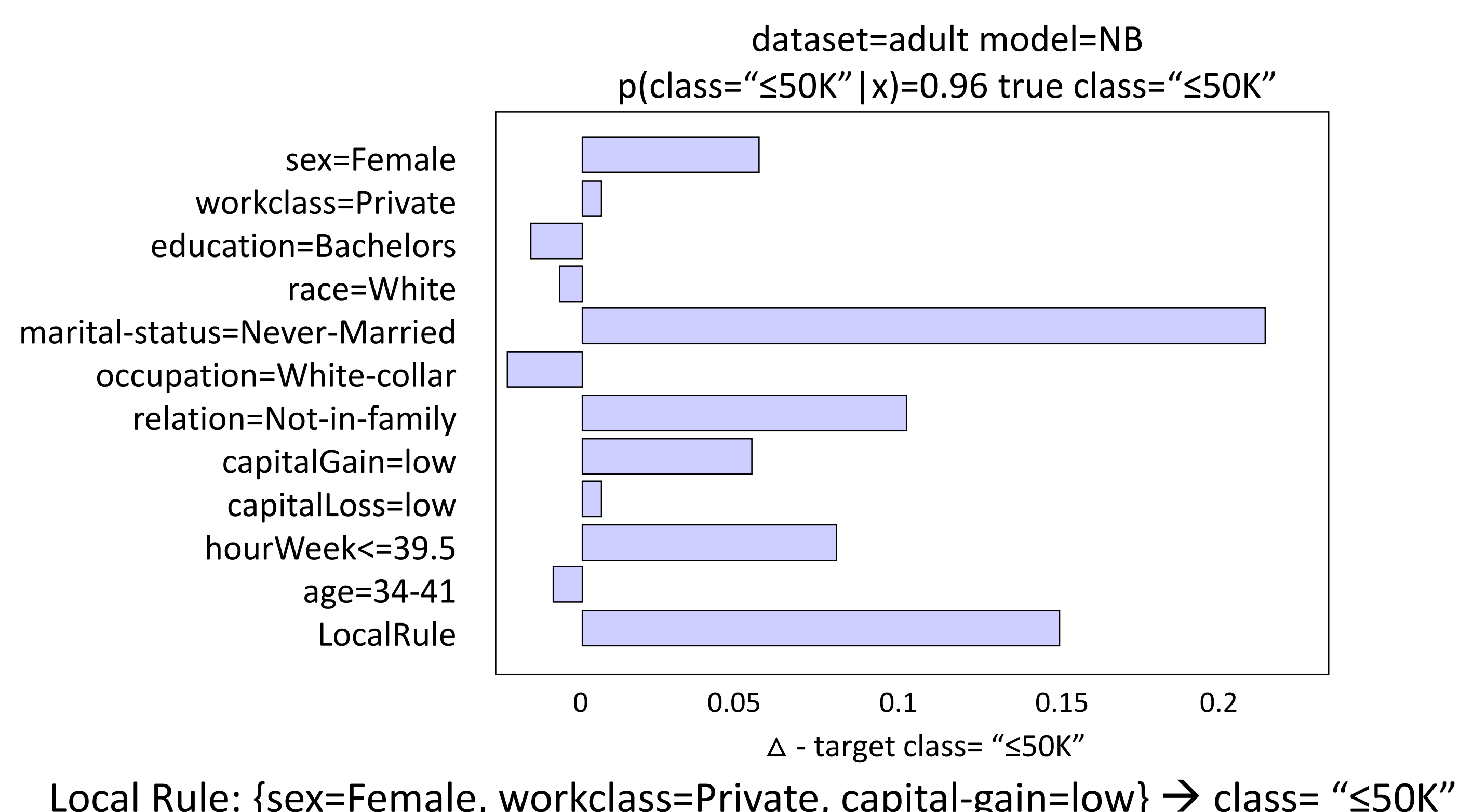
## LACE Explanation Results

- Explanation comparison



Local rules:  
 $\{e=1\} \rightarrow \text{class}=1$   
 $\{a=1, b=1\} \rightarrow \text{class}=1$

- Insight of reasons behind predictions



Local Rule:  $\{\text{sex}=\text{Female}, \text{workclass}=\text{Private}, \text{capital-gain}=\text{low}\} \rightarrow \text{class}=\leq 50K$

## Future work

- User studies  $\rightarrow$  to assess LACE ability to provide actionable insights on model behavior
- Explanation of Big Data model predictions