

– Seminar –

Cognitive Reasoning Seminar

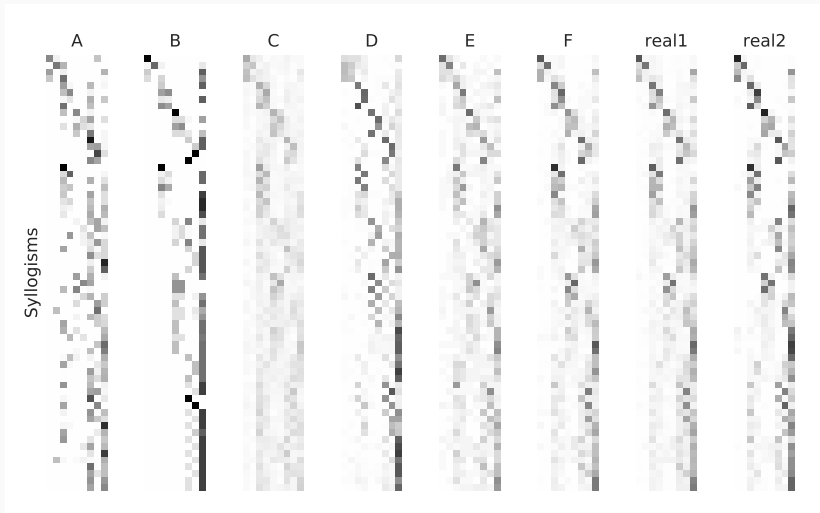
Midterm Meeting

apl. Prof. Dr. Dr. Marco Ragni
Nicolas Riesterer, Daniel Brand
November 19th, 2019

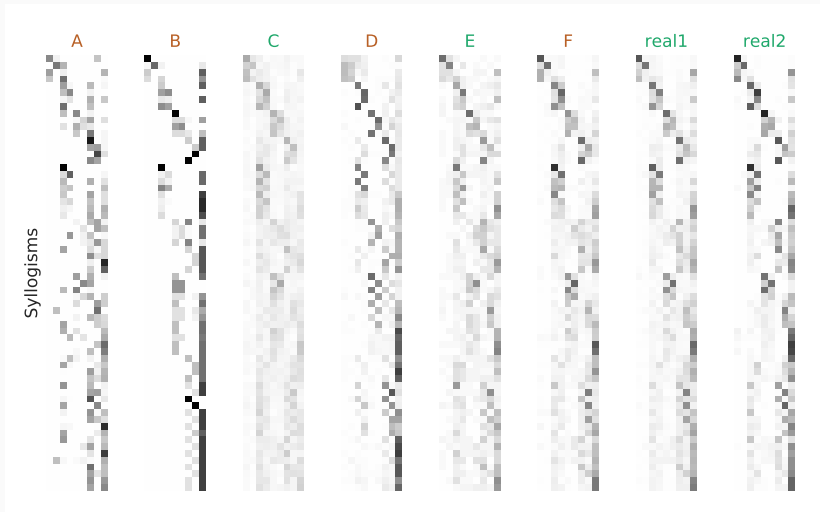
Cognitive Computation Lab,
Department of Computer Science,
University of Freiburg

Given known real data of human syllogistic reasoning, label other datasets as *real* or *artificial*.

Datasets



Datasets

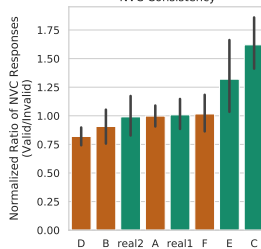
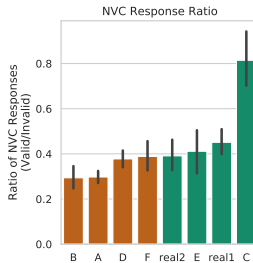
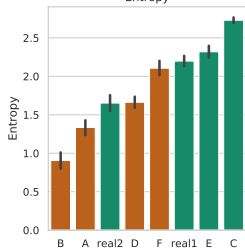
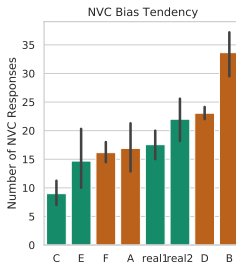
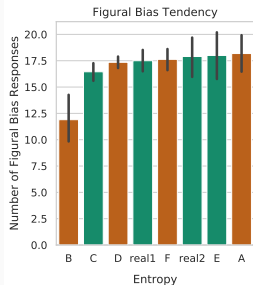


Data Generation Strategies

Dataset	Generation Strategy
A	Autoencoder supplied with random latent states
B	Cognitive model: <i>mReasoner</i>
D	Cognitive model: <i>Probability Heuristics Model</i>
F	Responses from votes by samples of individuals from real1

- Goal is to find discrepancies between the datasets
- Need to investigate different properties of the datasets
- Two directions to approach this task:
 1. Aggregate properties of the datasets
 - Distribution of responses
 - Variances of certain response biases
 - Comparison with theoretical insight into reasoning (figural bias, NVC responses, etc.)
 2. Individual response behavior
 - Individual response strategies in the datasets

Aggregate Data Analysis



Aggregate Data Analysis - Conclusions

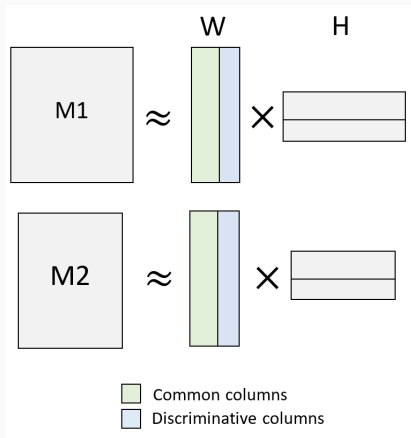
Dataset	Decision	Truth
A	Artificial	Artificial
B	Artificial	Artificial
C	Artificial	Real
D	–	Artificial
E	–	Real
F	–	Artificial

Data Analysis: Matrix Factorization

- Given a matrix $M = m \times n$, the goal is to find two matrices $W = m \times k$ and $H = k \times n$ (for $k \ll m, n$), so that $M \approx W \times H$
- Latent states have to be exploited to compress the data
- The principle is often used in topic extraction and recommender systems
- The approach can also be adapted to constrast data

Contrasting data using Matrix Factorization

- The main idea is to perform two matrix factorizations at the same time¹
- Regularization forces k_c columns to be similar and k_d columns to be dissimilar

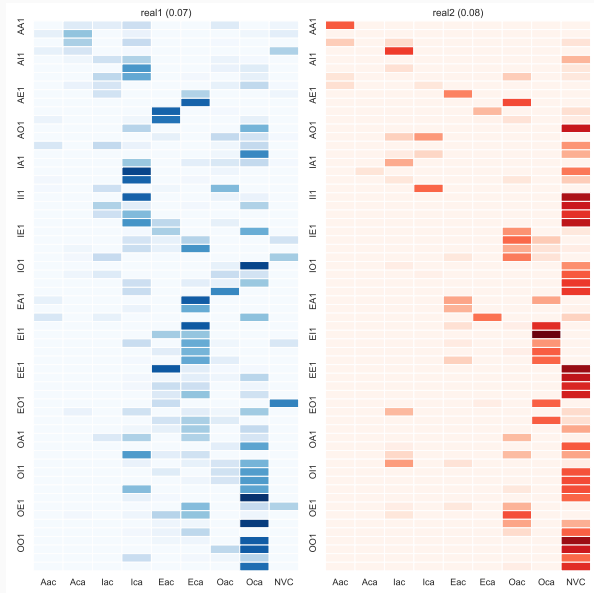


¹Kim, H., Choo, J., Kim, J., Reddy, C. K., & Park, H. (2015). Simultaneous discovery of common and discriminative topics via joint nonnegative matrix factorization

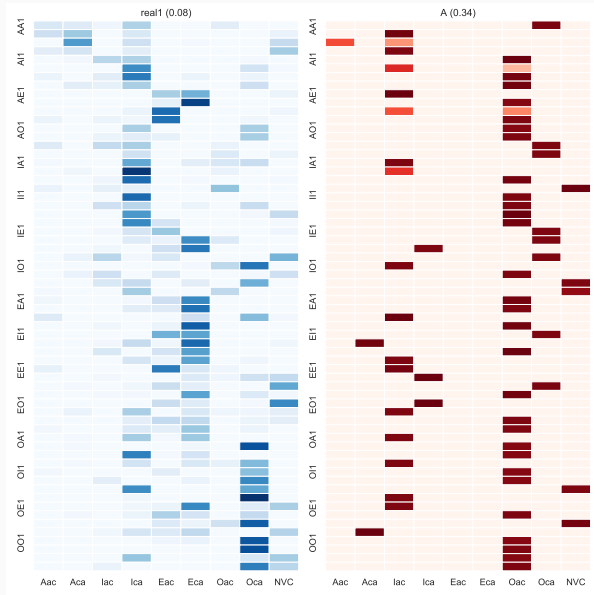
Contrasting data using Matrix Factorization

- We expect the datasets to contain mostly similar patterns, so we use $k_c = 9$ common and $k_d = 1$
- The different response patterns of the datasets can be obtained from the W-matrices
- The importance of the differences can be estimated from the H-matrices
- For this task, we used mainly the importance (is there a pattern that substantially discriminates between both datasets)

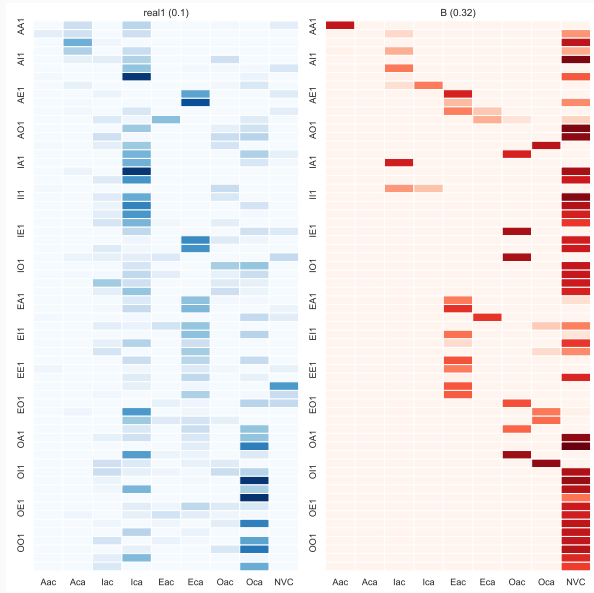
Contrasting data: Reference



Contrasting data: Dataset A



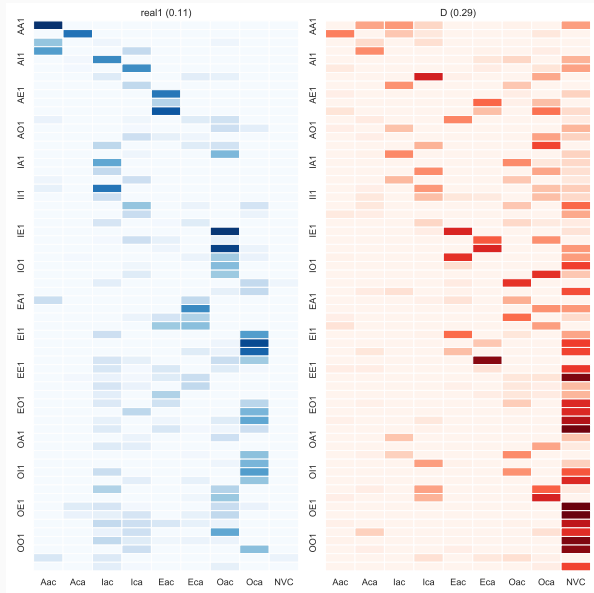
Contrasting data: Dataset B



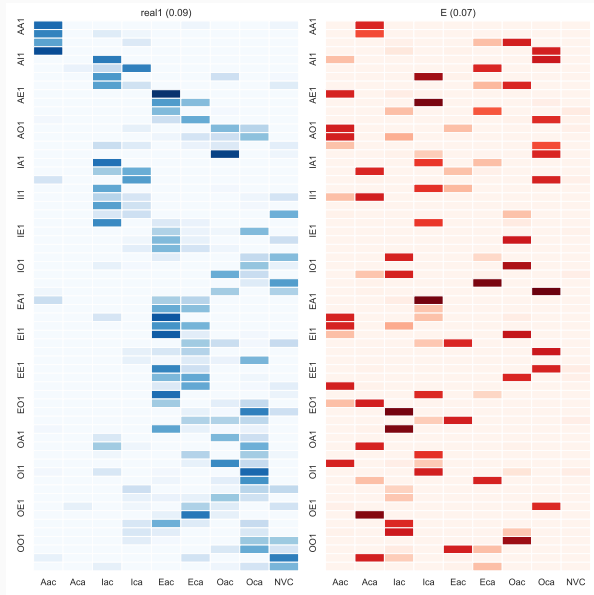
Contrasting data: Dataset C



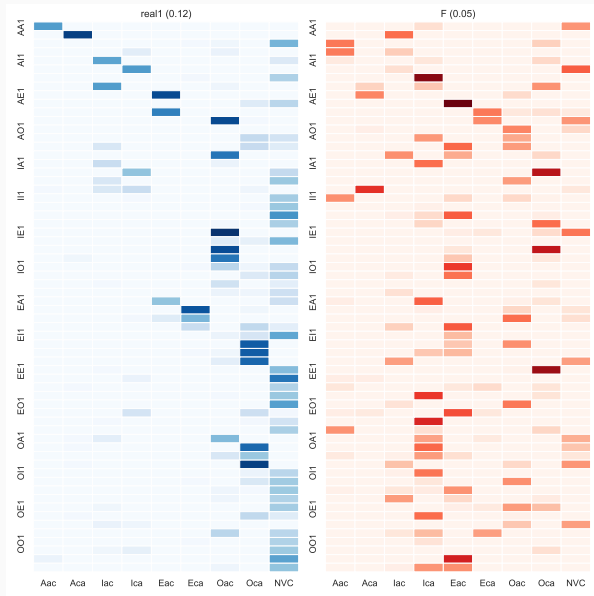
Contrasting data: Dataset D



Contrasting data: Dataset E



Contrasting data: Dataset F



Contrasting data - Conclusions

Dataset	Importance	Decision	Truth
A	0.22	Artificial	Artificial
B	0.17	Artificial	Artificial
C	0.12	Real	Real
D	0.18	Artificial	Artificial
E	0.09	Real	Real
F	0.08	Real	Artificial

- Use mean importance of both reference datasets
- Base rate is expected to be about 0.1
- Significantly higher values are suspicious

Next Steps

- Improve your classifiers by integrating knowledge about artificial datasets
- **Tip:** You can generate artificial data by yourself
- We will provide new unlabeled data for you to predict
- Keep in mind that this is a real-world problem and some fake datasets are really hard to detect

Upcoming Deadlines

- December 15th, 23:59: Deadline for final models & written report
- December 20th-21st, 09:00-17:00, room TBA: Blockseminary