

Automating validation of learning and decision making models using the CogniBench framework

Filip Melinscak

University of Zurich, Zurich, Switzerland

Eshref Yozdemir

University of Zurich, Zurich, Switzerland

Dominik R. Bach

University College London, London, United Kingdom

Abstract

Much of cognitive science is based on constructing, validating, and comparing formal models of the mind. Whereas coming up with new and useful models requires expertise and creativity, validating the proposed models and comparing them against the state-of-the-art mainly requires a systematic, rigorous approach. The task of model validation is therefore particularly well-suited for the types of automation that have propelled other research fields (cf. impact of bioinformatics on biology). Here we propose a model benchmarking framework implemented as an open-source Python package named CogniBench. Given a set of candidate models (which can be implemented in various languages), experimental observations, and scoring criteria, CogniBench automatically performs model benchmarks and reports the resulting matrix of scores. We demonstrate the potential of the proposed framework by applying it in the domain of learning and decision making, which poses unique requirements for model validation.