

Abstract

While we like to think of big data and algorithms as neutral, precise and accurate, there are multiple examples of not only biased, but also discriminating machine learning models in the world. Therefore fairness plays a crucial part in the development of any machine learning model. In person reId, fairness focuses mostly on the ability to re-identify all kinds of persons featuring all kinds of visible features. For a person reId model, which is currently in research, we explore, what kinds of biases generally need to be considered in this field. We analyze two commonly used public datasets for those types of bias, using information about how they were generated, as well as labelled attributes, which have been added for each identity in the sets. In addition, we compare the results of the model trained and tested on different datasets to see what effect biased training data has on the model, but also to see if the model itself amplifies or mitigates biases, that we already saw in the training data.

In the datasets we do see examples of bias, mainly caused by the fact, that, like many other public person reId datasets, they were generated within a short timespan in a specific location. The analysis of the model shows, that the biases of training data are not transferred directly onto the test results of the research-model. However, we can confirm that combining multiple datasets to a training set does lead to more accurate re-identification and our test results underline the relevance of background bias. Also we see that the general brightness of an image might have an impact on its signature, although it is not an attribute of the depicted person. It becomes clear, that biases have to be considered, when using public person reId datasets. How to mitigate all of them systematically, might be topic of further investigations. Also, with regard to our research-model, it ought to be evaluated, if the relation between the brightness of an image and its signature is a causal one, and if so, this bias should be mitigated in order to improve the model.