

Individually Fair Clustering with Outliers

Machine learning based applications have become prevalent in various domains of life, be it in e-commerce, governance, banking, or healthcare. As an example, take the case of customer segmentation for bank loans, where customers of similar characteristics (demographics, transaction history, income, etc.) are clustered together such that each segment gets offered similar loans. However, due to reasons such as errors in data gathering or abnormal spending habits of a few customers within a segment, the loans offered to that particular segment could be very different from the loans that would have been offered if such outliers were not present. That is, due to the presence of these outliers, the loans offered are not fair to each individual. We are interested in the problem of segmenting the customers in such a way that the loans offered to these customers are individually fair, which requires us to exclude the outliers before segmentation.

This can be modelled as a problem of individually fair clustering after excluding outliers. Formally, given a set of n points, which is known to contain a set of m outliers, we want to cluster the $n-m$ points into k clusters in an individually fair manner. The notion of individual fairness in clustering states that each of the $n-m$ points must have a center within its n/k neighbours.

While previous works have looked into the problem of individually fair clustering, this is the first work that has considered the problem in the presence of outliers. The previous algorithms are not suited to handle outliers, as naively applying those algorithms results in clusters that could potentially have a large fairness violation. That is, in the presence of outliers, the output of the previous algorithms becomes suboptimal. In order to tackle this problem, we propose a local search based algorithm. We propose a novel center initialization algorithm that is able to discard outliers. This is followed by the local search algorithm. Local search based algorithms are known to be computationally expensive as it involves swapping each candidate center and a datapoint multiple times. Instead, we adopt a randomized constrained local search algorithm that makes our algorithm scalable to large datasets.

Using our method, we can show that we discard only a bounded number of points as outliers. Also, the solution computed by our algorithm gives an $O(1)$ approximation to the cost of the optimal individually fair clustering, excluding outliers. Empirically, we demonstrate our method on three datasets from the UCI Repository: Adult, Bank, and Skin. We considered the k -means clustering objective for our problem. Since we are excluding outliers, we are able to beat the baseline state-of-the-art algorithms in terms of clustering cost as well as the maximum fairness violation for a varying number of clusters. Many of the baseline algorithms are not scalable, hence we were only able to experiment on approximately 4000 datapoints (for datasets larger than this, some algorithms did not even converge after 30+ hours). However, given that our algorithm is scalable, our algorithm, when applied to the Skin dataset (262k points), it converged within 1 hour.