



Legal but Unfair: Auditing the Impact of Data Minimization on Fairness and Accuracy Trade-off in Recommender Systems

Salvatore Bufi*
Politecnico di Bari
Bari, Italy
s.bufl@phd.poliba.it

Vito Walter Anelli
Department of Electrical and Information Engineering
Politecnico di Bari
Bari, Italy
vitowalter.anelli@poliba.it

Vincenzo Paparella*†
ISTI-CNR
Pisa, Italy
vincenzo.paparella@isti.cnr.it

Tommaso Di Noia
Department of Electrical and Information Engineering
Politecnico di Bari
Bari, Italy
tommaso.dinoia@poliba.it

Abstract

Data minimization, required by recent data privacy regulations, is crucial for user privacy, but its impact on recommender systems remains largely unclear. The core problem lies in the fact that reducing or altering the training data of these systems can drastically affect their performance. While previous research has explored how data minimization affects recommendation accuracy, a critical gap remains: How does data minimization impact consumers' and providers' fairness? This study addresses this gap by systematically examining how data minimization influences multiple objectives in recommender systems, i.e., the trade-offs between accuracy, user fairness, and provider fairness. Our investigation includes (i) an analysis of how the data minimization strategies affect RS performance across these objectives, (ii) an assessment of data minimization techniques to determine those that can balance better the trade-off among the considered objectives, and (iii) an evaluation of the robustness of different recommendation models under diverse minimization strategies to identify those that best maintain performance. The findings reveal that data minimization can sometimes undermine provider fairness, albeit enhancing group-based consumer fairness to the detriment of accuracy. Additionally, different strategies can offer diverse trade-offs for the assessed objectives. The source code supporting this study is available at <https://github.com/salvatore-bufl/DataMinimizationFairness>.

CCS Concepts

• **Information systems** → **Recommender systems**; • **Social and professional topics** → **Governmental regulations**.

Keywords

Data Minimization, Recommender Systems, Fairness, Multi-Objective Evaluation

*Corresponding authors.

†Work done while being a PhD student at Politecnico di Bari.



This work is licensed under a Creative Commons Attribution 4.0 International License. UMAP '25, New York City, NY, USA

© 2025 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-1313-2/25/06

<https://doi.org/10.1145/3699682.3728356>

ACM Reference Format:

Salvatore Bufl, Vincenzo Paparella, Vito Walter Anelli, and Tommaso Di Noia. 2025. Legal but Unfair: Auditing the Impact of Data Minimization on Fairness and Accuracy Trade-off in Recommender Systems. In *33rd ACM Conference on User Modeling, Adaptation and Personalization (UMAP '25)*, June 16–19, 2025, New York City, NY, USA. ACM, New York, NY, USA, 10 pages. <https://doi.org/10.1145/3699682.3728356>

1 Introduction

Recommender Systems (RSs) have become a cornerstone for modern digital platforms, powering personalized content delivery across e-commerce, media streaming, and social networks [24, 27, 28, 35]. These systems rely on extensive user data to accurately infer preferences and produce highly tailored recommendations. However, growing concerns about user privacy, data security, and system scalability create a significant tension between the need to maintain high-quality recommendations and the increasingly important requirement to minimize data collection [14, 16]. This tension is further intensified by rigid regulatory frameworks like the GDPR [3], CCPA [1], CPRA [2], and PIPL [4], which command principles such as data minimization. The legal principle of data minimization underscores the need to collect only the strictly necessary data for a specific purpose. Given these regulations, many works have started to study the effects and risks of data minimization on machine learning-driven systems [9, 18, 39, 44]. In the realm of recommendation, Biega et al. [7] performed pioneering work to understand the performance-based effects of data minimization on system accuracy. Hence, other works have analyzed the influence of data minimization on privacy [15], algorithmic bias [9], and individual users' accuracy [37] performance of recommendations. Nonetheless, data minimization reshapes available user-item interactions by discarding or selectively retaining interactions, potentially altering the distribution of users and items. These alterations can amplify existing biases and undermine system fairness from consumer and producer perspectives. For this reason, it is pertinent to ask: *To what extent does data minimization risk to impact fair exposure of items? Can data minimization favor privileged user groups with higher-quality recommendations?* For example, a popular item may appear underrepresented in a minimized dataset if many interactions are removed, potentially pushing it into the long-tail category and reducing its visibility. Similarly, specific user groups may receive

less accurate or less diverse recommendations, further widening existing disparities. While these risks are evident, a comprehensive understanding of the broader ramifications of data minimization on accuracy and fairness remains lacking, highlighting a critical gap in current research, given the increasing emphasis on fairness [6, 12, 31, 32, 43] in RSs.

This study bridges the gap by systematically investigating the interplay between data minimization, accuracy, and fairness in recommendation. To this end, we investigate how data minimization singularly impacts the accuracy and fairness of recommendations by identifying observable trade-offs. Then, we adopt a multi-objective evaluation approach to audit to what extent different data minimization strategies can preserve RS performance on multiple dimensions. Finally, we analyze how different recommendation models are robust to data minimization, maintaining a balance between accuracy and fairness. In summary, our key contributions are the following:

- We experimentally evaluate the impact of data minimization on recommendation accuracy, consumer fairness, and provider fairness across various data minimization strategies and real-world datasets. The evaluation demonstrates that the effectiveness of data minimization cannot be evaluated by accuracy alone but must consider how it affects different stakeholders in the system.
- We adopt a multi-objective evaluation approach to study which data minimization strategies can simultaneously preserve recommendation performance when considering multiple aspects. Our findings reveal that strategies that can shape better user profiles are crucial to avoid alterations in recommendation performance.
- We investigate the effects of data reduction on various recommendation algorithms, including neighborhood-based, matrix factorization, autoencoder-based, and graph-based methods. We understand how different model architectures react to data minimization when balancing the competing objectives of accuracy and fairness. We show that graph-based and factorization-based models can preserve recommendation performance under strong minimization conditions.

2 Related Work

User privacy [10, 13, 26] and data privacy [22, 45] have become a critical concern in recommender systems research, driven by regulations such as the GDPR, CCPA, CPRA, and PIPL. These legal frameworks mandate that only data necessary for a specific purpose be collected, creating tension between the extensive datasets typically required by recommender systems and the principle of data minimization.

Several works have begun to explore how to reconcile the need for personalized recommendations with the constraints imposed by these regulations. Biega et al. [7] introduced operational definitions of data minimization aimed at maintaining system quality while satisfying legal requirements. However, they also pointed out that minimization can worsen outcomes for specific user groups if not carefully managed. A family of works focuses on designing data minimization techniques to be integrated during model training to retain acceptable accuracy performance. Shanmugam et al. [36] proposed FIDO, a strategy that adaptively restricts data collection based on performance scaling rules, illustrating how data minimization can be integrated directly into training workflows. Sonboli et al. [37]

have explored the trade-off between reducing data and maintaining accuracy at the user level by using active learning approaches that select only highly informative user-item interactions. While these methods can improve efficiency, they can also worsen outcomes for groups with sparse interaction histories. Other works [38, 42] study the effects of data reduction on environmental factors, showing that smaller datasets obtained through data reduction lead to reduced energy consumption but may compromise recommendation accuracy. Regarding the interplay between data minimization and fairness, Clavell et al. [9] highlight the complexities of simultaneously addressing privacy, performance, and fairness, emphasizing the need for more comprehensive frameworks to evaluate these trade-offs.

The existing works concentrate on a single objective—such as privacy, bias, or accuracy—rather than investigating how these factors interact in a unified manner. Our work extends the current state-of-the-art by systematically examining how data minimization strategies simultaneously affect both accuracy and fairness in recommender systems.

3 Experimental Setup

This research carries out a comprehensive investigation into the effects of data minimization on recommender systems, moving beyond a simple accuracy analysis to consider its broader implications for both users and item providers. In this section, we detail the experimental setup employed to address the following research questions:

- RQ1:** How does data minimization impact the accuracy, provider fairness, and consumer fairness of recommendations? What are the observable trade-offs at the strategy level?
- RQ2:** Which data minimization strategies most effectively preserve recommender system performance compared to training on the full dataset?
- RQ3:** How robust are different recommendation models to data minimization, particularly concerning maintaining a balance between accuracy and fairness?

3.1 Dataset

Our experiments were conducted using two datasets to ensure a comprehensive evaluation: *MovieLens1M* (ML1M) [20] and *Ambar* [19]. Setting a minimum interaction threshold is essential when studying data minimization since it ensures each user profile is represented well enough for reliable model training and evaluation [7]. For this reason, we preprocessed both datasets to guarantee sufficient representation for each user and item, aligning with the methodology of Biega et al. [7]. In ML1M, we selected users who have at least 45 ratings in their profile. Furthermore, we subsampled the dataset to 2500 users, covering 3605 items and totaling 522152 ratings. The resulting dataset has a density of 0.0579, with a median of 137 interactions per user and an average of 208.86. In Ambar, we applied an iterative 45-core filtering, ensuring that each user rated at least 45 items and that at least 45 users rated each item. We then subsampled the dataset to 2500 users, 4996 items, and 214678 total ratings, resulting in a density of 0.0172. The median number of interactions per user is 74, with an average of 85.87.

3.2 Evaluation Protocol

Our experimental evaluation aims to emulate a real-world scenario in which users can grant permission for their data storage. Hence, solely data related to users who allowed their data storage explicitly are included in the analysis, resulting in a reduced dataset called *Minimization Data* (D_M).

To create this scenario, the users of the entire dataset are randomly divided into two groups. The first group, representing 70% of all users, simulates those who do not consent to data storage and are excluded from the experiments. The second group consists of the remaining 30% of users who have consented to their data being stored. The interaction data related to these users constitute D_M . Within D_M , each user's ratings are further partitioned into three sets: a candidate set (70%), a validation set (10%), and a test set (20%). The candidate set serves as a pool of user data from which different minimization strategies can select a subset of ratings to train the system. Data is selected from this candidate set using a specific data minimization strategy and a minimization parameter n , which indicates the number of items to choose for each user. The study uses the following values for $n = \{1, 3, 7, 15, 100\}$. This approach draws inspiration from Biega et al. [7] but adapts it to a practical scenario, where individual user choices determine the availability of their data to the system.

3.3 Algorithms and Metrics

To provide a comprehensive view of how data minimization affects different types of recommender systems, we consider five representative algorithms spanning different methodological families:

- **LightGCN** [21] is a graph-based model that simplifies the aggregation process of Graph Convolutional Networks by eliminating the feature transformation step and non-linear activation functions. It aggregates messages from neighboring nodes linearly, assigning equal weights to all edges.
- **User-kNN** [34] is a neighborhood-based model that computes recommendations by identifying the nearest neighbors of a target user and aggregating their preferences. Similarity between users is typically defined based on the overlap of items they have co-rated, and final predictions are derived from a weighted combination of neighbor ratings.
- **BPRMF** [33] is a pairwise learning-to-rank method based on matrix factorization.
- **EASE^R** [41] is a linear model conceptually resembling an autoencoder by computing a closed-form item-item similarity matrix.
- **MultiVAE** [25] harnesses a deep variational autoencoder to learn non-linear latent representations from entire user interactions.

We evaluate the quality of recommendations of the selected algorithms through nDCG [23]. Furthermore, we select two fairness-related metrics to assess the impact of data minimization on recommendation performance on consumer and provider fairness:

- **Mean Absolute Deviation (MAD)** [11]: it measures the disparity in recommendation quality across different user groups. Formally, let Q_u denote the quality of the recommendation for the user u and let $\mathcal{U}_x$ and $\mathcal{U}_y$ be two distinct groups of users. The mean absolute deviation (MAD) between these groups is

defined as:

$$MAD(Q^{(\mathcal{U}_x)}, Q^{(\mathcal{U}_y)}) = \left| \frac{\sum_{x \in \mathcal{U}_x} Q_x}{|\mathcal{U}_x|} - \frac{\sum_{y \in \mathcal{U}_y} Q_y}{|\mathcal{U}_y|} \right|. \quad (1)$$

When the users are partitioned in more than two groups, MAD values are calculated for all possible pairwise group combinations. The final metric is derived as the mean of these computed MAD values. A smaller MAD value implies more equitable treatment across the groups, whereas a larger MAD reveals greater discrepancies and, thus, a higher degree of group-level unfairness. In our experiments, we use nDCG as the quality metric. In the ML1M dataset, users are partitioned by gender, while users are grouped by continent for the Ambar dataset.

- **Ranking-based Statistical Parity (RSP)** [46]: it measures how evenly different item categories are exposed to users in the top- k positions. Let us denote c_i as the item categories, and $P(R@k | c_i)$ the probability of a user encountering items in category c_i among the top- k recommendations ($R@k$). The RSP@ k is defined as:

$$RSP@k = \frac{std(P(R@k | c_1), \dots, P(R@k | c_n))}{mean(P(R@k | c_1), \dots, P(R@k | c_n))} \quad (2)$$

where $std(\cdot)$ and $mean(\cdot)$ denote the standard deviation and mean, respectively. A lower RSP indicates that the ranking probabilities are more evenly distributed, signifying less bias in how items are exposed to users. For ML1M, we categorize items into four release-year ranges. In the Ambar dataset, we group items by their continent of origin.

We also employ the **Hypervolume (HV)** [17] as a multi-objective metric to assess the model performance on accuracy, consumer, and provider fairness simultaneously. We choose the reference point $r = (1, 1.5, 1)$ for the three objectives: nDCG, RSP, and MAD.

3.4 Reproducibility Details

We adopt the all-unrated-items protocol [40], where the recommendable items are those not previously interacted with for each user. In our experiments, the embedding size is fixed at 64, and the batch size and the maximum number of epochs are 1024 and 300, respectively. Furthermore, we tune the model's hyper-parameters through grid search. For BPRMF and LightGCN, the learning rate and the regularization term are tuned in $\{0.01, 0.005, 0.001\}$. The number of layers in LightGCN is set to 2. EASE^R is tuned by varying its regularization factor in $\{500, 1000, 5000\}$. Moreover, for UserKNN, the neighborhood size is searched in $\{20, 30, 50\}$, while we use Cosine and Jaccard as similarity measures. Finally, for MultiVAE, the learning rate is tuned in $\{0.01, 0.001\}$ while the dimension of both the latent representation and the hidden layer is searched in $\{100, 200\}$. Combining all hyperparameter explorations for every model and strategy resulted in 2345 total trained models. We select the best models by adopting nDCG@10 as the validation metric, employing a patience value of 5 epochs for early stopping. Then, model performance results are computed on the test set. All the metrics are measured with a cutoff at 10, i.e., considering the top 10 suggested items to each user. To foster reproducibility, the models are trained using the reproducibility framework Elliot [5].

3.5 Data Minimization Strategies

In our study, we evaluate several approaches to minimize user data. These strategies are designed to select subsets of user-item interactions that serve as input to the system. In this paper, we study the data minimization strategies explored by Biega et al. [7], aiming at broadening their assessment to fairness issues. We briefly describe such strategies in the following.

Full data. This baseline utilizes the complete set of user interactions without minimization. This strategy serves as a reference point against which the effectiveness of all other strategies is measured. Formally, the system observes all interactions in the user history.

Random selection. This approach randomly selects n interactions from each user's history. Although randomness may introduce variability, the average profile characteristics of users are often preserved due to the sampling process.

Most recent. This strategy selects the n most recent user interactions by assuming that recent behaviors are more relevant to the recommendation task. The effectiveness of this method depends on the temporal data distribution. Long observation periods may lead to performance variation when compared to random selection. We do not apply this strategy to the Ambar dataset due to lack of timestamp information.

Most and least favorite. These techniques select the n items with the highest or lowest scores for users, respectively. While the former emphasizes user preferences, the latter explores how negative interactions influence recommendations.

Most rated. This method selects n items rated most frequently in the dataset for each user. This global popularity measure identifies items that are statistically significant within the system.

Most characteristic. It selects items based on their proximity to the average item profile. Each item is represented as a binary vector indicating which users have interacted with it. The average item profile is computed by averaging all these binary vectors in the system. We calculate the Euclidean distance for each item between its binary representation and this average vector. Items with lower distances are deemed more representative of the overall trends and are therefore considered “most characteristic.” For each user, we select the n items they have interacted with closest to the system-wide profile, ensuring that the retained items best capture standard user behavior while avoiding niche or rarely rated items.

Highest variance. This strategy prioritizes items with the highest variability in their ratings across all users. Items with high variance are considered more informative as they capture diverse or polarizing opinions. The n items with the highest rating variance are selected for each user.

4 Results and Discussion

This section presents the empirical findings that address our three research questions.

4.1 Impact of Data Minimization (RQ1)

We answer RQ1 by discussing the impact of data minimization strategies on the accuracy, provider fairness, and consumer fairness of recommendation algorithms. Tables 1 and 2 report the model's

performance measured with nDCG, RSP, MAD, and HV (see Section 3.3) for various minimization strategies and parameter n values for the ML1M and Ambar datasets, respectively. For each metric, we report the percentage differences in the performance between the considered strategy and the “full” strategy in parenthesis. Following Biega et al. [7], we say that minimization has a low impact on accuracy if the difference in nDCG is statistically insignificant or minimal. From the fairness perspective, minimization positively (negatively) impacts fairness if it strongly reduces (increases) the values of RSP and MAD, for which the lower, the better.

Impact on accuracy. Data minimization seriously alters the accuracy of recommendation algorithms. Reducing the parameter n affects recommendation accuracy, measured by nDCG, across datasets, models, and minimization strategies. This analysis corroborates the insights from Biega et al. [7], revealing that smaller n leads to a consistent drop in accuracy. Specifically, extreme minimization ($n = 1$) results in accuracy reductions that generally exceed 90% for many configurations, rendering recommendations impractical, as users cannot receive relevant suggestions. The accuracy performance loss is smaller but still noteworthy for moderate minimization ($n = 3$ and $n = 7$). At $n = 15$, dataset-specific differences emerge. In ML1M, the drop in accuracy remains pronounced, while the nDCG values become more acceptable for specific strategies in Ambar. Interestingly, the “Most Favorite” strategy even improves accuracy compared to the “full” strategy in some cases. At $n = 100$, accuracy differences become minimal or statistically insignificant in Ambar, demonstrating the feasibility of data minimization in this dataset under moderate conditions. In ML1M, performance differences remain measurable but less severe than smaller n values.

Impact on provider fairness. Data minimization introduces variability in provider fairness, as measured by RSP. Lower n values lead to inconsistent changes in RSP, with both improvements (negative changes compared to the “full” strategy) and deteriorations (positive changes) observed. This variability indicates a loss of control over provider fairness, particularly at smaller n , where fairness outcomes become unpredictable. Regarding the Ambar dataset, the relationship between accuracy and provider fairness is evident: strategies achieving higher nDCG often induce worse provider fairness. For example, the “Most Rated” strategy, which excels in accuracy, consistently increases RSP compared to the “full” strategy, indicating fairness degradation. In contrast, the “Most Characteristic” strategy sacrifices accuracy but consistently improves provider fairness, highlighting a trade-off between these two objectives.

Impact on consumer fairness. Data minimization generally enhances consumer fairness, as lower MAD values indicate. This improvement mainly stems from a reduction in overall nDCG, which levels the quality of recommendations across user groups to the detriment of overall recommendation relevance.

Impact of strategies. The interplay between data minimization and the objectives of accuracy, provider fairness, and consumer fairness is intricate, with varying impacts on recommendation performance. To unravel these complexities, we analyze the minimization strategies to identify patterns in their performance when $n \leq 15$. The “Most Rated” strategy, which emphasizes the presence of popular items in the training dataset, generally achieves the best or

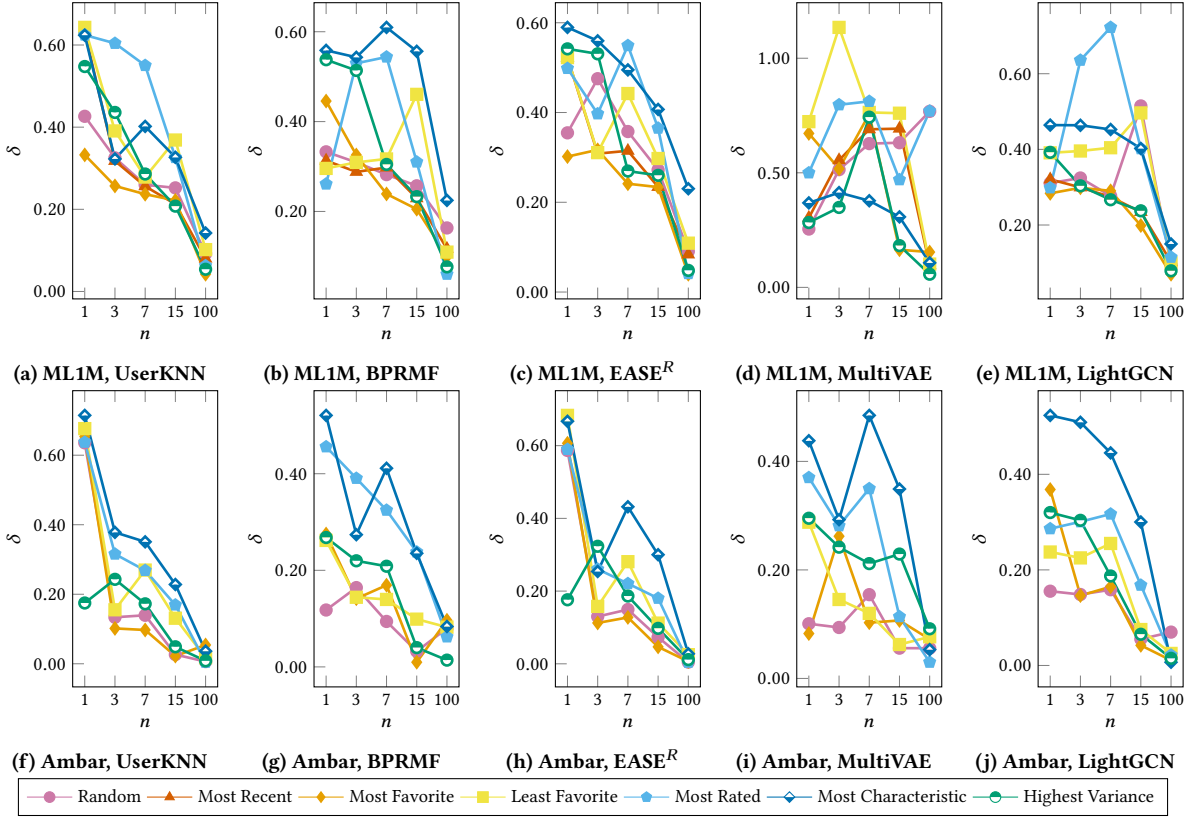


Figure 1: Comparison of the Euclidean distance (δ) between the point representing the model performance under a minimization strategy with value n and the point representing the model performance under the “full” data strategy. The colored lines represent the different minimization strategy. Smaller δ values indicate better preservation of the model’s overall performance.

second-best nDCG values. However, this focus on popularity comes at a cost: provider fairness, as measured by RSP, deteriorates significantly, especially at smaller n . As n increases, the items are exposed more fairly due to the inclusion of a broader range of items in the training set. Nevertheless, the low values of HV indicate that this strategy predominantly optimizes for accuracy, with limited alignment to beyond-accuracy objectives. The “Most Favorite” strategy prioritizes the items most liked by individual users, resulting in competitive nDCG values (often the second-best) while mitigating the fairness degradation seen in “Most Rated.” This strategy strikes a better balance between accuracy and fairness since the items most liked by users are not necessarily the most popular. Indeed, the acceptable HV values suggest a more harmonious fulfillment of accuracy and fairness objectives. In contrast, the “Most Characteristic” strategy consistently achieves the worst performance in terms of accuracy but demonstrates improvements in provider fairness. This strategy favors the representation of items most interacted with by the most active users, effectively tailoring user profiles for this subgroup. However, on most platforms, most users are less active, resulting in a decline in overall recommendation quality when adopting this strategy. The best HV values associated with this strategy indicate that it excels in achieving beyond-accuracy objectives, albeit with poor accuracy performance. The remaining

strategies introduce higher variability in user profiles and tend to achieve a better balance between accuracy and fairness. This balance is reflected in their HV values, which suggest moderate fulfillment of all objectives.

To answer RQ1, we conclude that data minimization can lead to a decrease in the accuracy performance of recommendation models. While this outcome directly offers fairer recommendations across different user groups, data minimization can be responsible for unfair exposure of items. Hence, these objectives potentially show a critical trade-off, underscoring the need for nuanced strategies to balance these competing objectives in recommendation systems.

4.2 Data Minimization Strategies Performance (RQ2)

To answer RQ2, we evaluate how effectively data minimization strategies preserve the performance of recommender systems compared to the “full” strategy. We adopt a multi-objective evaluation approach [29, 30] that simultaneously considers multiple dimensions of analysis. Fixing a value of n and a recommendation model, we measure the model’s performance regarding nDCG, RSP, and MAD for each minimization strategy applied. These metrics define a three-dimensional objective function space, where each model trained under a specific minimization strategy is represented as

Table 3: Mean (μ) and standard deviation (σ) of the distances among the points represented by each model performance under different strategies and the point corresponding to the model performance under the “full” data strategy, across different n . For each scenario, bold and underline stand for best and second-to-best values, respectively.

Model	$n = 1$		$n = 3$		$n = 7$		$n = 15$		$n = 100$		Global	
	$\mu \downarrow$	$\sigma \downarrow$	$\mu \downarrow$	$\sigma \downarrow$	$\mu \downarrow$	$\sigma \downarrow$	$\mu \downarrow$	$\sigma \downarrow$	$\mu \downarrow$	$\sigma \downarrow$	$\mu \downarrow$	$\sigma \downarrow$
Movielens1M												
UserKNN	0.547	0.1216	0.379	<u>0.1146</u>	0.3241	0.1134	0.2736	0.0656	0.0821	<u>0.0339</u>	0.3212	0.1777
BPRMF	<u>0.3921</u>	0.1211	0.4027	<u>0.1188</u>	<u>0.3705</u>	0.1443	0.3217	0.1346	0.1216	0.056	0.3217	0.1537
EASER	0.4758	<u>0.1058</u>	0.4141	0.1082	0.3814	<u>0.1169</u>	<u>0.2957</u>	<u>0.0666</u>	<u>0.0918</u>	0.0667	0.3317	0.1622
MultiVAE	0.4435	0.1916	0.6115	0.2701	0.6814	0.1461	0.4583	0.2455	0.2935	0.3259	0.4977	0.2676
LightGCN	0.3515	0.0657	<u>0.3887</u>	0.1253	0.384	0.1662	0.3551	0.1309	0.1014	0.0261	0.3161	<u>0.1538</u>
Ambar												
UserKNN	0.5821	0.2017	0.2212	0.1098	0.2163	0.0956	0.1043	0.085	<u>0.0212</u>	<u>0.02</u>	0.229	0.2236
BPRMF	0.3166	0.1467	0.2223	0.0969	<u>0.2245</u>	0.1203	<u>0.109</u>	0.1031	0.0698	0.0294	<u>0.1884</u>	<u>0.1339</u>
EASER	0.5516	0.1884	0.2063	<u>0.0847</u>	0.2328	0.1116	0.135	<u>0.0926</u>	0.0133	0.0105	0.2278	0.21
MultiVAE	0.2625	<u>0.1433</u>	<u>0.2196</u>	0.0809	0.2372	0.1506	0.1529	0.1145	0.0632	0.0216	0.1871	0.1275
LightGCN	<u>0.3151</u>	0.1254	0.2722	0.1351	0.2546	<u>0.1113</u>	0.1177	0.1	0.0249	0.0231	0.1969	0.1483

a point identified by the corresponding metric values. The model trained using the “full” strategy serves as a reference point in this space. To quantify the similarity of performance between the “full” strategy and each minimization strategy, we compute the Euclidean distance δ of every point (i.e., the model trained under a minimization strategy) from the reference point. A smaller Euclidean distance indicates that the minimization strategy better preserves the model’s overall performance compared to the “full” strategy. Figure 1 presents the Euclidean distance results for each model and dataset as n varies.

Consistent with the trends observed in Tables 1 and 2, setting $n = 100$ enables most strategies to effectively preserve model performance, as evidenced by the lowest δ values. No single strategy emerges as a definitive winner at this level, reflecting similar performance across strategies. Conversely, as n decreases, a more evident hierarchy among the strategies emerges.

The “Most Rated” and “Most Characteristic” strategies generally exhibit higher δ values, indicating poorer performance preservation—except for MultiVAE trained on ML1M (Figure 1d). This behavior aligns with the observations in Section 4.1: the “Most Rated” strategy degrades model performance primarily regarding fairness, while the “Most Characteristic” strategy compromises accuracy.

In contrast, the “Most Favorite” strategy consistently preserves model performance, frequently achieving the lowest distance values. As discussed in Section 4.1, this strategy effectively delineates user profiles without adversely affecting fairness. Similar performance trends can be observed for the “Random,” “Most Recent,” and “Highest Variance” strategies. These strategies introduce variability into user profiles, yet this variability enables the retention of average user profile characteristics. Consequently, these strategies maintain acceptable accuracy levels while avoiding a disproportionate emphasis on specific items in the training dataset. The resulting datasets closely resemble those generated by the “full” strategy, which explains their superior performance compared to the “Most

Rated” and “Most Characteristic” strategies. Among these, the “Random” strategy is the most effective. As noted in Section 3.5, it best maintains the average profile of each user. Finally, the “Least Favorite” strategy generally underperforms compared to the “Random” strategy. This behavior is likely because models find learning effectively from negative feedback more challenging, resulting in more significant performance degradation.

To answer RQ2, we say that data minimization strategies that can shape better user profiles and introduce variability strike a better balance among the considered objectives, being able to retain the model’s performance compared to the “full” strategy.

4.3 Robustness of Recommenders to Data Minimization (RQ3)

To address RQ3, we evaluate the robustness of recommendation models under data minimization while considering both accuracy and fairness. Here, robustness is defined as a model’s ability to preserve its performance under data minimization compared to the full-data setting. We adopt the multi-objective approach outlined in Section 4.2 to quantify this. We compute the performance achieved under each data minimization strategy for each model in the three-dimensional objective function space—defined by the nDCG, RSP, and MAD metrics. Given a model and fixed n , we measure the Euclidean distances of these performance points from the reference point, i.e., the model’s performance under the “full” strategy. The mean μ and standard deviation σ of these distances capture the model’s robustness: a lower μ indicates better overall robustness, while a lower σ suggests more consistent performance across strategies. Table 3 reports μ and σ for each model across different values of n and their global averages across all settings (i.e., without fixing n).

In Ambar and ML1M, UserKNN demonstrates strong robustness for $n \geq 3$, consistently achieving the lowest (or second-lowest) values of μ and σ . This finding aligns with previous works [7, 8], which

showed that neighborhood-based models can maintain performance even after removing random interactions. Similarly, EASE^R is robust only when data minimization is moderate. Its reliance on item co-occurrence statistics makes it highly sensitive to severe sparsity. If the goal is to ensure robustness even with minimal data ($n = 1$), LightGCN emerges as the most reliable model across both datasets. This finding aligns with its graph-based structure, which propagates information from neighboring users and items, mitigating the impact of data sparsity. Moreover, LightGCN maintains low μ and σ across all scenarios, making it the most robust model overall. BPRMF follows closely behind LightGCN, demonstrating strong generalization under data minimization. Its factorization-based approach enables it to retain performance even with limited interactions, though it lacks the graph-based smoothing effect of LightGCN, leading to slightly higher variability. Finally, MultiVAE exhibits the most dataset-dependent behavior, showing weaker robustness in ML1M but achieving the best performance in Ambar at the global level. This observation suggests that its variational latent space representations are highly sensitive to dataset characteristics, which may impact its stability under data minimization.

To answer RQ3, we conclude that graph-based modeling is crucial for robustness under extreme data sparsity when considering accuracy and fairness objectives, enhancing the classic factorization-based representation of users and items. Additionally, traditional neighborhood-based and non-linear methods show limitations in such scenarios, suffering from high data reduction.

5 Conclusion and Future Work

In this study, we systematically investigated the impact of data minimization on recommender systems, with particular attention to the trade-offs among accuracy, consumer fairness, and provider fairness. Our contributions extend the current understanding of data minimization in three key ways. First, we empirically evaluate how different data minimization strategies across real-world datasets impact the accuracy, provider, and consumer fairness of recommendations. Second, we employ a multi-objective evaluation to investigate which data minimization strategies can simultaneously preserve recommendation performance across multiple objectives. Third, we explore how different model classes respond to varying degrees of data reduction, shedding light on the relationship between architectural design and data scarcity.

Our findings show that although data minimization can diminish accuracy performance, it levels the recommendation quality of different user groups, enhancing consumer fairness. However, reducing data can inadvertently lead to the unfair exposure of specific items, highlighting a critical trade-off between data minimization, accuracy, and fairness. Further analysis highlights specific behaviors of data minimization strategies. Particularly, approaches that introduce variability when shaping the user profiles can strike a better balance among data minimization, accuracy, and fairness of recommendations. Finally, our insights indicate that graph-based recommender systems remain more resilient under extreme data minimization when fairness and accuracy are jointly considered, in contrast to traditional methods that exhibit lower robustness.

Future works may seek to devise more refined data minimization strategies that promote equitable recommendations while minimizing adverse effects on accuracy performance. In addition, our insight can guide the development of recommender systems specifically devised to be robust under severe data minimization conditions while accomplishing fairness objectives.

Acknowledgments

The authors acknowledge partial support of the following projects: OVS: Fashion Retail Reloaded, ePansa, Lutech Digitale 4.0, Secure Safe Apulia, Patti Territoriali WP1, BIO-D, Meditech. We acknowledge the CINECA award under the ISCR initiative for the availability of high-performance computing resources and support.

References

- [1] [n. d.]. California Consumer Privacy Act (CCPA). <https://cloud.google.com/security/compliance/ccpa>. accessed 2024-02-13.
- [2] [n. d.]. The California Privacy Rights Act of 2020. <https://thecpra.org/>. accessed 2024-02-13.
- [3] [n. d.]. Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016 on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation). <https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=CELEX:32016R0679>. accessed 2024-03-15.
- [4] [n. d.]. Translation: Personal Information Protection Law of the People's Republic of China – Effective Nov. 1, 2021. <https://digichina.stanford.edu/work/translation-personal-information-protection-law-of-the-peoples-republic-of-china-effective-nov-1-2021/>. accessed 2024-02-13.
- [5] Vito Walter Anelli, Alejandro Bellogin, Antonio Ferrara, Daniele Malatesta, Felice Antonio Merra, Claudio Pomo, Francesco Maria Donini, and Tommaso Di Noia. 2021. Elliot: A Comprehensive and Rigorous Framework for Reproducible Recommender Systems Evaluation. In *SIGIR*. ACM, 2405–2414.
- [6] Vito Walter Anelli, Yashar Deldjoo, Tommaso Di Noia, Daniele Malatesta, Vincenzo Paparella, and Claudio Pomo. 2023. Auditing Consumer- and Producer-Fairness in Graph Collaborative Filtering. In *Advances in Information Retrieval - 45th European Conference on Information Retrieval, ECIR 2023, Dublin, Ireland, April 2–6, 2023, Proceedings, Part I (Lecture Notes in Computer Science, Vol. 13980)*, Jaap Kamps, Lorraine Goeuriot, Fabio Crestani, Maria Maistro, Hideo Joho, Brian Davis, Cathal Gurrin, Udo Kruschwitz, and Annalina Caputo (Eds.). Springer, 33–48. doi:10.1007/978-3-031-28244-7_3
- [7] Asia J. Biega, Peter Potash, Hal Daumé III, Fernando Diaz, and Michèle Finck. 2020. Operationalizing the Legal Principle of Data Minimization for Personalization. In *SIGIR*. ACM, 399–408.
- [8] Richard Chow, Hongxia Jin, Bart P. Knijnenburg, and Gökyay Saldamli. 2013. Differential data analysis for recommender systems. In *Seventh ACM Conference on Recommender Systems, RecSys '13, Hong Kong, China, October 12–16, 2013*, Qiang Yang, Irwin King, Qing Li, Pearl Pu, and George Karypis (Eds.). ACM, 323–326. doi:10.1145/2507157.2507190
- [9] Gemma Galdon Clavell, Mariano Martín Zamorano, Carlos Castillo, Oliver Smith, and Aleksandar Matic. 2020. Auditing Algorithms: On Lessons Learned and the Risks of Data Minimization. In *AIES*. ACM, 265–271.
- [10] Enyan Dai, Tianxiang Zhao, Huaisheng Zhu, Junjie Xu, Zhimeng Guo, Hui Liu, Jiliang Tang, and Suhang Wang. 2024. A Comprehensive Survey on Trustworthy Graph Neural Networks: Privacy, Robustness, Fairness, and Explainability. *Mach. Intell. Res.* 21, 6 (2024), 1011–1061.
- [11] Yashar Deldjoo, Vito Walter Anelli, Hamed Zamani, Alejandro Bellogin, and Tommaso Di Noia. 2021. A flexible framework for evaluating user and item fairness in recommender systems. *User Model. User Adapt. Interact.* 31, 3 (2021), 457–511. doi:10.1007/S11257-020-09285-1
- [12] Yashar Deldjoo, Dietmar Jannach, Alejandro Bellogin, Alessandro Difonzo, and Dario Zanzonelli. 2024. Fairness in recommender systems: research landscape and future directions. *User Model. User Adapt. Interact.* 34, 1 (2024), 59–108.
- [13] Yashar Deldjoo, Tommaso Di Noia, and Felice Antonio Merra. 2022. A Survey on Adversarial Recommender Systems: From Attack/Defense Strategies to Generative Adversarial Networks. *ACM Comput. Surv.* 54, 2 (2022), 35:1–35:38.
- [14] Tobias Eichinger and Axel Küpper. 2023. Distributed Data Minimization for Decentralized Collaborative Filtering Systems. In *ICDCN*. ACM, 140–149.
- [15] Tobias Eichinger and Axel Küpper. 2024. On data minimization and anonymity in pervasive mobile-to-mobile recommender systems. *Pervasive Mob. Comput.* 103 (2024), 101951.
- [16] Michèle Finck and Asia Biega. 2021. Reviving Purpose Limitation and Data Minimisation in Personalisation, Profiling and Decision-Making Systems. *CoRR*

- abs/2101.06203 (2021).
- [17] M. Fleischer. 2003. The Measure of Pareto Optima. In *EMO (Lecture Notes in Computer Science, Vol. 2632)*. Springer, 519–533.
 - [18] Prakhhar Ganesh, Cuong Tran, Reza Shokri, and Ferdinando Fioretto. 2024. The Data Minimization Principle in Machine Learning. *CoRR* abs/2405.19471 (2024).
 - [19] Elizabeth Gómez, David Contreras, Ludovico Boratto, and Maria Salamó. 2024. AMBAR: A dataset for Assessing Multiple Beyond-Accuracy Recommenders. In *RecSys*. ACM, 137–147.
 - [20] F. Maxwell Harper and Joseph A. Konstan. 2016. The MovieLens Datasets: History and Context. *ACM Trans. Interact. Intell. Syst.* 5, 4 (2016), 19:1–19:19.
 - [21] Xiangnan He, Kuan Deng, Xiang Wang, Yan Li, Yong-Dong Zhang, and Meng Wang. 2020. LightGCN: Simplifying and Powering Graph Convolution Network for Recommendation. In *SIGIR*. ACM, 639–648.
 - [22] Yassine Himeur, Shahab Saquib Sohail, Faycal Bensaali, Abbas Amira, and Mamoun Alazab. 2022. Latest trends of security and privacy in recommender systems: A comprehensive review and future perspectives. *Comput. Secur.* 118 (2022), 102746.
 - [23] Walid Krichene and Steffen Rendle. 2020. On Sampled Metrics for Item Recommendation. In *KDD '20: The 26th ACM SIGKDD Conference on Knowledge Discovery and Data Mining, Virtual Event, CA, USA, August 23–27, 2020*, Rajesh Gupta, Yan Liu, Jiliang Tang, and B. Aditya Prakash (Eds.). ACM, 1748–1757. doi:10.1145/3394486.3403226
 - [24] Nian Li, Chen Gao, Depeng Jin, and Qingmin Liao. 2023. Disentangled Modeling of Social Homophily and Influence for Social Recommendation. *IEEE Trans. Knowl. Data Eng.* 35, 6 (2023), 5738–5751.
 - [25] Dawen Liang, Rahul G. Krishnan, Matthew D. Hoffman, and Tony Jebara. 2018. Variational Autoencoders for Collaborative Filtering. In *WWW*. ACM, 689–698.
 - [26] Jianghao Lin, Xinyi Dai, Yunjia Xi, Weiwen Liu, Bo Chen, Hao Zhang, Yong Liu, Chuhan Wu, Xiangyang Li, Chenxu Zhu, et al. 2023. How can recommender systems benefit from large language models: A survey. *ACM Transactions on Information Systems* (2023).
 - [27] Qin Liu, Xuan Feng, Tianlong Gu, and Xiaoli Liu. 2024. FairCRS: Towards User-oriented Fairness in Conversational Recommendation Systems. In *RecSys*. ACM, 126–136.
 - [28] Alberto Carlo Maria Mancino, Antonio Ferrara, Salvatore Bui, Daniele Malatesta, Tommaso Di Noia, and Eugenio Di Sciascio. 2023. KGTOR: Tailored Recommendations through Knowledge-aware GNN Models. In *RecSys*. ACM, 576–587.
 - [29] Vincenzo Paparella, Vito Walter Anelli, Franco Maria Nardini, Raffaele Perego, and Tommaso Di Noia. 2023. Post-hoc Selection of Pareto-Optimal Solutions in Search and Recommendation. In *CIKM*. ACM, 2013–2023.
 - [30] Vincenzo Paparella, Dario Di Palma, Vito Walter Anelli, and Tommaso Di Noia. 2023. Broadening the Scope: Evaluating the Potential of Recommender Systems beyond prioritizing Accuracy. In *Proceedings of the 17th ACM Conference on Recommender Systems, RecSys 2023, Singapore, Singapore, September 18–22, 2023*, Jie Zhang, Li Chen, Shlomo Berkovsky, Min Zhang, Tommaso Di Noia, Justin Basilico, Luiz Pizzato, and Yang Song (Eds.). ACM, 1139–1145. doi:10.1145/3604915.3610649
 - [31] Hossein A. Rahmani, Mohammadmehdi Naghiaei, and Yashar Deldjoo. 2024. A Personalized Framework for Consumer and Producer Group Fairness Optimization in Recommender Systems. *Trans. Recomm. Syst.* 2, 3 (2024), 19:1–19:24.
 - [32] Theresia Veronika Rampisela, Maria Maistro, Tuukka Ruotsalo, and Christina Lioma. 2023. Evaluation Measures of Individual Item Fairness for Recommender Systems: A Critical Study. *CoRR* abs/2311.01013 (2023).
 - [33] Steffen Rendle, Christoph Freudenthaler, Zeno Gantner, and Lars Schmidt-Thieme. 2012. BPR: Bayesian Personalized Ranking from Implicit Feedback. *CoRR* abs/1205.2618 (2012).
 - [34] Paul Resnick, Neophytos Iacovou, Mitesh Suchak, Peter Bergstrom, and John Riedl. 1994. GroupLens: An Open Architecture for Collaborative Filtering of Netnews. In *CSCW*. ACM, 175–186.
 - [35] Abdulaziz Samra, Evgeny Frolov, Alexey Vasilev, Alexander Grigorevskiy, and Anton Vakhruhev. 2024. Cross-Domain Latent Factors Sharing via Implicit Matrix Factorization. In *RecSys*. ACM, 309–317.
 - [36] Divya Shanmugam, Fernando Diaz, Samira Shabanian, Michèle Finck, and Asia Biega. 2022. Learning to Limit Data Collection via Scaling Laws: A Computational Interpretation for the Legal Principle of Data Minimization. In *FAccT*. ACM, 839–849.
 - [37] Nasim Sonboli, Sipei Li, Mehdi Elahi, and Asia Biega. 2024. The trade-off between data minimization and fairness in collaborative filtering. *CoRR* abs/2410.07182 (2024).
 - [38] Giuseppe Spillo, Allegra De Filippo, Cataldo Musto, Michela Milano, and Giovanni Semeraro. 2024. Towards Green Recommender Systems: Investigating the Impact of Data Reduction on Carbon Footprint and Algorithm Performances. In *RecSys*. ACM, 866–871.
 - [39] Robin Staab, Nikola Jovanovic, Mislav Balunovic, and Martin T. Vechev. 2024. From Principle to Practice: Vertical Data Minimization for Machine Learning. In *SP*. IEEE, 4733–4752.
 - [40] Harald Steck. 2013. Evaluation of recommendations: rating-prediction and ranking. In *RecSys*. ACM, 213–220.
 - [41] Harald Steck. 2019. Embarrassingly Shallow Autoencoders for Sparse Data. In *WWW*. ACM, 3251–3257.
 - [42] Tobias Vente, Lukas Wegmeth, Alan Said, and Joeran Beel. 2024. From Clicks to Carbon: The Environmental Toll of Recommender Systems. In *Proceedings of the 18th ACM Conference on Recommender Systems, RecSys 2024, Bari, Italy, October 14–18, 2024*, Tommaso Di Noia, Pasquale Lops, Thorsten Joachims, Katrien Verbert, Pablo Castells, Zhenhua Dong, and Ben London (Eds.). ACM, 580–590. doi:10.1145/3640457.3688074
 - [43] Yifan Wang, Weizhi Ma, Min Zhang, Yiqun Liu, and Shaoping Ma. 2023. A Survey on the Fairness of Recommender Systems. *ACM Trans. Inf. Syst.* 41, 3 (2023), 52:1–52:43.
 - [44] Hongyi Wen, Longqi Yang, Michael Sobolev, and Deborah Estrin. 2018. Exploring recommendations under user-controlled data filtering. In *RecSys*. ACM, 72–76.
 - [45] Keping Yu, Zhiwei Guo, Yu Shen, Wei Wang, Jerry Chun-Wei Lin, and Takuro Sato. 2022. Secure Artificial Intelligence of Things for Implicit Group Recommendations. *IEEE Internet Things J.* 9, 4 (2022), 2698–2707.
 - [46] Ziwei Zhu, Jianling Wang, and James Caverlee. 2020. Measuring and Mitigating Item Under-Recommendation Bias in Personalized Ranking Systems. In *Proceedings of the 43rd International ACM SIGIR conference on research and development in Information Retrieval, SIGIR 2020, Virtual Event, China, July 25–30, 2020*, Jimmy X. Huang, Yi Chang, Xueqi Cheng, Jaap Kamps, Vanessa Murdock, Ji-Rong Wen, and Yiqun Liu (Eds.). ACM, 449–458. doi:10.1145/3397271.3401177