

Operationalizing the Legal Principle of Data Minimization for Personalization

Asia J. Biega
Microsoft Research
Montréal

Peter Potash
Microsoft Research
Montréal

Hal Daumé III
Microsoft Research NYC
University of Maryland

Fernando Diaz
Microsoft Research
Montréal

Michèle Finck
Max Planck Institute for Innovation
and Competition

ABSTRACT

Article 5(1)(c) of the European Union’s General Data Protection Regulation (GDPR) requires that “personal data shall be [...] adequate, relevant, and limited to what is necessary in relation to the purposes for which they are processed (‘data minimisation’).” To date, the legal and computational definitions of ‘purpose limitation’ and ‘data minimization’ remain largely unclear. In particular, the interpretation of these principles is an open issue for information access systems that optimize for user experience through personalization and do not strictly require personal data collection for the delivery of basic service.

In this paper, we identify a lack of a homogeneous interpretation of the data minimization principle and explore two operational definitions applicable in the context of personalization. The focus of our empirical study in the domain of recommender systems is on providing foundational insights about the (i) feasibility of different data minimization definitions, (ii) robustness of different recommendation algorithms to minimization, and (iii) performance of different minimization strategies. We find that the performance decrease incurred by data minimization might not be substantial, but that it might disparately impact different users—a finding which has implications for the viability of different formal minimization definitions. Overall, our analysis uncovers the complexities of the data minimization problem in the context of personalization and maps the remaining computational and regulatory challenges.

KEYWORDS

GDPR, Data Minimization, Purpose Limitation, Personalization

ACM Reference Format:

Asia J. Biega, Peter Potash, Hal Daumé III, Fernando Diaz, and Michèle Finck. 2020. Operationalizing the Legal Principle of Data Minimization for Personalization. In *Proceedings of the 43rd International ACM SIGIR Conference on Research and Development in Information Retrieval (SIGIR '20)*, July 25–30, 2020, Virtual Event, China. ACM, New York, NY, USA, 10 pages. <https://doi.org/10.1145/3397271.3401034>

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

SIGIR '20, July 25–30, 2020, Virtual Event, China

© 2020 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM ISBN 978-1-4503-8016-4/20/07...\$15.00

<https://doi.org/10.1145/3397271.3401034>

1 INTRODUCTION

Personalized services such as recommender systems or search engines collect large amounts of user interaction logs. Such data collection practice is widely accepted to be necessary for platforms to build high-quality models [18, 37]. However, some prior work shows that exact user interaction profiles are not necessary to tailor the results of search or recommendations. For instance, Singla *et al.* show that it is possible to personalize results while storing a reduced user interaction history [34], while Biega *et al.* show that it is possible to shuffle user queries and ratings while preserving the quality of personalized search and recommendations [2].

If results can be personalized without exact user profiles, it is pertinent to ask: *How much information and what information does an individual need to provide to receive quality personalized results?* Note the parallel between this question and the principle of *data minimization* defined in Article 5 of the European Union’s General Data Protection Regulation (GDPR) [33] as well as data protection regimes in other jurisdictions, which requires that a system only retain user data necessary to deliver service. The core idea we explore in this work is whether the principles of purpose limitation and data minimization can be complied with in the context of personalization and what minimizing data in this context entails.

In contrast to other GDPR concepts, such as *the right to be forgotten* or *informed consent*, there is to date only marginal regulatory and judicial guidance on the interpretation of data minimization. Reasoning about data minimization has largely been confined to setups involving immutable or relatively stationary user characteristics. For instance, examples mentioned in the guidelines issued by the UK’s Information Commissioner’s Office [31] discuss scenarios of collecting people’s names by debt collectors, or employee blood types by employers. More recent regulatory guidelines and industrial practice, however, recognize the multitude of challenges related to minimization in data-intensive applications [3, 13].

To the best of our knowledge, this work is the first to operationalize the legal concepts of purpose limitation and data minimization in a scenario where *user data collection is not strictly necessary to deliver a service, but where the collection of such data might improve service quality*. We tie the purpose of data collection to performance metrics, and define *performance-based minimization principles*.

In this study, we investigate two possible technical definitions of performance-based data minimization. The first interpretation, which we refer to as *global data minimization*, minimizes per-user data collection subject to meeting a target mean performance across users. This aligns well with standard empirical risk minimization

approaches in machine learning [41]. Alternatively, *per-user data minimization* minimizes per-user data collection subject to each user meeting a target performance. Equivalently, this aligns with meeting a target performance for the minimum across all users.

We use these interpretations to compare different minimization strategies for personalized recommendations. We demonstrate that quality recommendations can be provided while collecting substantially less user data. However, we also find that the appropriate minimization strategy is sensitive to the base recommendation algorithm used. While our results suggest that systems should be able to achieve global data minimization, we demonstrate that preserving the average performance conceals substantial impact for individual users. To sum up, the salient contributions of this paper are:

- Identifying a lack of a homogeneous interpretation of the GDPR's purpose limitation and data minimization principles in the context of personalization systems and proposing a computational definition of *performance-based data minimization*.
- An analysis of the feasibility of two different data minimization definitions in the domain of recommender systems.
- An analysis of the robustness of different recommendation algorithms to various minimization strategies, both on a population as well as an individual user levels.

2 DATA MINIMIZATION

2.1 A legal perspective

Article 5(1)(c) GDPR requires that personal data be '*adequate, relevant and limited to what is necessary in relation to the purposes for which they are processed*.' Data minimisation is the direct consequence of the legal principle of *purpose limitation*, which requires that personal data only be processed for specified, explicit and legitimate purposes and not further processed in a manner incompatible with these purposes. While these core data protection principles cannot be examined exhaustively here, it is worth noting that general statements such as 'improving user experience' are generally not specific enough to meet the legal threshold of purpose limitation. This raises the question of whether 'personalization' can be a purpose under the GDPR at all.

According to data minimisation, no more personal data than necessary to achieve the purpose can be processed. The first question to ask is thus whether data such as this studied in our paper is personal data. The Article 4 GDPR embraces a very broad definition of personal data as 'any information relating to an identified or identifiable natural person.' In the past, movie ratings such as those in the MovieLens 20M dataset [19] have been shown to allow for identification through linking of private and public datasets [29]. It is thus safe to assume that much of the data used in recommender systems, such as movie ratings, constitutes personal data and is hence subject to the GDPR (where within its geographical scope).

Data minimisation can be broken down into three distinct requirements. First, data must be *adequate* in relation to the purpose that is pursued. Arguably, adequacy is the most interesting of the three criteria as it may actually (and somewhat counterintuitively) require that more data is processed. It is well established that the omission of certain data can limit the usefulness of a dataset and the accuracy of an analysis done on that dataset. As such, to achieve

accurate results, more data may need to be collected. Data minimisation indeed ought to be interpreted in light of the other substantive requirements in Article 5 GDPR such as fairness, transparency and accuracy and there are scenarios, often those involving under-represented groups, where this can only be achieved through the processing of more personal data.

Second, data ought to be *relevant* in light of the purpose, meaning that only data that is pertinent for the purpose can be processed. For example, if an e-commerce provider requested users' date of birth to provide personalised recommendations regarding future purchases, this data is unlikely to be relevant (except where recommendations have an astrological flavor). Relevance thus functions as a safeguard against accumulating data simply for the sake of doing so.

Third, the GDPR requires that data be *limited* to what is necessary, meaning that controllers ought to identify the minimum amount of personal data required to fulfil the stated purpose. Thus, where similarly robust results can be achieved through the processing of less personal data, the processing of personal data can likely not be accepted as being necessary. Where possible, only anonymised data should be used. However, given the practical limitations of achieving anonymisation, the latter cannot be assumed as a viable alternative to minimisation in many contexts [10].

2.2 Performance-Based Data Minimization

Our focus in this paper is on operationalizing the third requirement of data minimization, namely that of limitation. According to the legal considerations detailed in the previous subsection, generic statements such as 'improving user experience' are not specific enough to be used as a purpose of data collection. Thus, we propose to reason about data minimization by tying the *purpose* to *performance metrics*. While there are many ways in which this proposition might be operationalized, in this paper, we begin investigating this space with an empirical study of two definitions.

Let $\mathcal{U}$ be a set of users for whom the system needs to minimize the data and let $\mathcal{I}$ be the set of items that a system can recommend. Each user has *rated* some subset $\mathcal{I}_u \subseteq \mathcal{I}$ of items. Let r_u be the $|\mathcal{I}| \times 1$ vector of ratings for these items. Of the rated items in $\mathcal{I}_u$, in a minimization setting, the system only sees a subset $\tilde{\mathcal{I}}_u \subseteq \mathcal{I}_u$, referred to as the observational pool for user u . Let $\hat{r}_u$ be the ratings for these observations. Given $\hat{r}_u$, a system generates $\tilde{r}_u$, its predicted ratings for u . The quality metric for u is defined as $\sigma(\tilde{r}_u)$.

DEFINITION 1 (GLOBAL DATA MINIMIZATION). A system satisfies *global data minimization* if it minimizes the amount of per-user data while achieving the quality of a system with access to the full data on average,

$$\min_k \text{ s.t. } \forall u, |\tilde{\mathcal{I}}_u| = k \text{ and } \mathbb{E}_{\mathcal{U}}[\sigma(\tilde{r}'_u)] - \mathbb{E}_{\mathcal{U}}[\sigma(\tilde{r}_u)] \leq \lambda$$

where $\tilde{r}'$ is the prediction using the ratings in $\mathcal{I}_u$ and λ is a threshold difference in the expected per-user performance.

DEFINITION 2 (PER-USER DATA MINIMIZATION). A system satisfies *per-user data minimization* if it minimizes the amount of per-user data while achieving the quality of a system with access to the full data for each user,

$$\min_k \text{ s.t. } \forall u, |\tilde{\mathcal{I}}_u| = k \text{ and } \forall u, \sigma(\tilde{r}'_u) - \sigma(\tilde{r}_u) \leq \lambda$$

where $\tilde{r}'$ is the prediction using the ratings in $\mathcal{I}_u$ and λ is a threshold difference in the per-user performance.

3 EXPERIMENTAL SETUP

3.1 Datasets

We run our analyses using (1) the MovieLens 20M dataset [19] and (2) the Google Location dataset [20]. Because of the space constraints, we report the results using dataset (1), and use dataset (2) for validation, reporting differences in observations where applicable. To properly reason about data minimization, we only select users who have at least 45 ratings in their profile. For efficiency reasons, we further subsample the users, creating (1) a MovieLens dataset containing around 2.5k users, 170k ratings, and 20k unique movies; the mean and median number of ratings in a user profile are 69.5 and 59, respectively, and (2) a Google Location dataset containing around 2.2k users, 185k ratings, and 150k unique items; the mean and median number of ratings in a user profile are 85.2 and 64, respectively.

3.2 Recommendation algorithms

We analyze data minimization properties for two fundamental classes of recommendation algorithms - neighborhood-based (k-nearest-neighbors) and matrix-factorization-based (SVD) [12], both as implemented in the Surprise library [21].

3.2.1 Notation. For a user u and item i , we use r_{ui} to denote the true rating given by the user for the item and $\tilde{r}_{ui}$ as the predicted rating by the user for the item from a predictive model,

3.2.2 Neighborhood-based. For the neighborhood-based recommendations, we use the user-user k-nearest-neighbors algorithm setting $k = 30$, as per prior studies investigating the recommendation performance in the MovieLens dataset [8]. Rating prediction $\tilde{r}_{ui}$ for user u and item i is computed as a weighted sum of the ratings of i made by u 's top- k nearest neighbors among users who rated item i :

$$\tilde{r}_{ui} = \frac{\sum_{v \in N_i^k(u)} \text{sim}(u, v) \cdot r_{vi}}{\sum_{v \in N_i^k(u)} \text{sim}(u, v)} \quad (1)$$

where $N_i^k(u)$ is the set of users v who have rated item i and who are most similar to user u by the value of similarity measure $\text{sim}(u, v)$.

User similarity is computed as the inverse of the mean squared difference of ratings (with add-1 smoothing) over set $I_u \cap I_v$.

3.2.3 Matrix-factorization-based. For the matrix-factorization-based recommendations, we use an implementation of the FunkSVD algorithm [12] with 30 latent factors. Rating prediction for user u and item i is computed as:

$$\tilde{r}_{ui} = \mu + b_u + b_i + q_i^\top p_u \quad (2)$$

where q_i is a 30-dimensional latent vector representing item i , p_u is a 30-dimensional latent vector representing user u , μ is a global mean, and b_i and b_u are item and user biases, respectively.

3.3 Error measures

We measure the quality of recommendations using: RMSE (comparing the differences between the predicted and true ratings for all items in the test set and thus assuming a user consumes the whole recommendation set) and NDCG (measuring the quality of the top

results with a logarithmic discounting factor for errors in lower ranking positions [22]). In our experiments, we set $k = 10$.

3.4 Protocol

We explore data minimization in the context of a system that begins with extensive data collection for a starting set of users. This may be gathered in-house or from a representative market not subject to data minimization constraints. While there will be situations where seed data is unavailable, we leave that for future work.

To simulate this situation, we randomly split the full dataset into two parts: the system data D_S (70% of all users), and the minimization data D_M (30% of all users). Users are randomly assigned to one of these groups. For minimizing users in D_M , we further randomly split their ratings into candidate (70% of all ratings) and test data (30% of all ratings). Different minimization strategies will select different subsets of each user's candidate data for use by the system. Recommendations generated based on the selected data from the candidate user data are evaluated using the remaining test data.

Data is selected from the candidate user data using a chosen minimization strategy and a minimization parameter n (the number of items to select). We run experiments for $n = \{1, 3, 7, 15, 100\}$.

3.5 Data minimization strategies

When minimizing data, we select a subset of user candidate items to present to the recommendation algorithm. While approaches with similar problem structure have used greedy algorithms modeling the information-theoretic utility of data [26], greedy algorithms are less practical in a data minimization scenario. Since utility of data is tied to a specific recommendation performance metric rather than modeled as information gain, the submodularity and monotonicity properties upon which guarantees on greedy algorithms are based do not necessarily hold. Moreover greedy selection is costly in terms of runtime, since the recommendation algorithm needs to be run for every possible selection. This section presents the selection strategies we study in this paper.

3.5.1 Full. We compare other minimization strategies against a baseline generating predictions based on full observational pools of users from D_M . Formally, $\tilde{I}_u = I_u$.

3.5.2 Empirical bounds. We compare the minimization results against brute-force baselines that select 1 item from a user's profile that lead to (i) the highest prediction RMSE (One item worst), (ii) the lowest prediction RMSE (One item best). We also compute (iii) the average RMSE error over all possible 1-item selections (One item avg); this value can be thought of as an empirical expected value of RMSE over 1-item random selections.

3.5.3 Random. This strategy selects n ratings uniformly at random from the observational pools of the minimizing users. The key observation to make here is that this method *will not create random user profiles* as a result, but minimized *average profiles* of each user. That is, if ratings of certain types (e.g., of a certain genre) are common in the full observational profile, they are likely to be preserved through the random sampling.

3.5.4 Most recent. This strategy selects n most recent ratings from the observational pools of the minimizing users. Note that one can

expect this method to behave similarly to the random method in case the data is collected over a period of time short enough for the user tastes to stay intact. In case the observational data of each user spans a very long time, we could expect the predictions to be better than random in case the test data is also sampled from the most recent ratings, and worse than random otherwise.

3.5.5 Most/least favorite. This strategies select the n ratings that have the highest/lowest value for a given user, respectively.

3.5.6 Most Rated. This method uses the system data to determine the selection method. For a given user, we select the n items that have been rated the most often (by the number of times an item has been rated by all users in the system data).

3.5.7 Most characteristic. This method uses the system data to determine the selection method for a given user. We create binary vector representations of items b_i by allocating each system data user to a dimension of b_i and setting the value to 1 if the user has rated item i , and 0 otherwise. We then take the average of all the item representations b_{avg} . Finally, for a given user we select the n items with the closest Euclidean distance to the average item representation. Whereas the most rated strategy treats all users the same when creating its counts, this strategy rewards items for being rated by users who have rated many items and penalizes items that have been rated by user who have rated few items. Formally, $\tilde{I}_u = \arg \min_{\{i\}} \sum d(b_i, b_{avg})$ s.t. $|\{i\}| = n$, where $d()$ is the Euclidean distance between two vectors, b_i is the binary representation of item i , and b_{avg} is the average item vector; all vectors are computed using the system data.

3.5.8 Highest variance. This method is based on one of the standard approaches for feature selection in machine learning [17]. It uses the system data to determine the selection method for each user by looking at which items have the highest variance in their ratings. Formally, $\tilde{I}_u = \arg \max_{\{i\}} \sum \sigma(\{r_{*i}\})^2$ s.t. $|\{i\}| = n$, where σ is standard deviation, and $\{r_{*i}\}$ is the set of all ratings for item i in the system data.

4 GLOBAL DATA MINIMIZATION

To guide the interpretation of the results, we want to make the following remarks. Reasoning about feasibility of data minimization, it is important to understand what quality loss we would incur if we based personalized recommendations on minimized user profiles. The main purpose of our experimental study is thus to measure and compare the quality of recommendations under different minimization conditions.

To reason about the efficacy of a minimization condition (maximum size of user profile n and a minimization strategy) for a given recommendation algorithm, we measure the difference in the quality of recommendations obtained under the minimization condition, and the quality of recommendations obtained if the recommendation algorithm sees all available user data (the Full strategy). We conclude that minimization is feasible if this difference is not statistically significant, or if the difference is minimal (low RMSE increase, and low NDCG decrease).

Table 1: Minimization performance for k-NN recommendations macro-averaged over all users. * denotes cases when the difference between a given strategy and the 'full' strategy is statistically significant under a two-tailed t-test with $p < 0.01$ and the Bonferroni correction. Average RMSE and NDCG@10 for non-minimized data is 0.915 and 0.777, respectively. Note that the lack of statistical significance suggests a minimization technique is performing well.

	n=1	n=3	n=7	n=15	n=100
RMSE					
random	1.062*	1.051*	1.013*	0.963*	0.915
most recent	1.044*	1.060*	1.028*	0.974*	0.915
most favorite	1.053*	1.046*	1.000*	0.957*	0.915
least favorite	1.049*	1.077*	1.039*	0.983*	0.915
most watched	1.064*	1.007*	0.966*	0.935*	0.914
most characteristic	1.008*	1.044*	1.073*	1.024*	0.915
highest variance	1.055*	1.071*	1.020*	0.955*	0.915
NDCG@10					
random	0.681*	0.721*	0.743*	0.762*	0.777
most recent	0.678*	0.708*	0.734*	0.760*	0.777
most favorite	0.697*	0.730*	0.751*	0.767	0.777
least favorite	0.662*	0.700*	0.733*	0.752*	0.777
most watched	0.721*	0.746*	0.764*	0.772	0.777
most characteristic	0.637*	0.656*	0.690*	0.737*	0.777
highest variance	0.664*	0.708*	0.744*	0.766*	0.777

Table 2: Minimization performance for SVD recommendations macro-averaged over all users. * denotes cases when the difference between a given strategy and the 'full' strategy is statistically significant under a two-tailed t-test with $p < 0.01$ and the Bonferroni correction. Average RMSE and NDCG@10 for non-minimized data is 0.818 and 0.793, respectively. Note that the lack of statistical significance suggests a minimization technique is performing well.

	n=1	n=3	n=7	n=15	n=100
RMSE					
random	0.876*	0.861*	0.843*	0.828*	0.818
most recent	0.875*	0.864*	0.851*	0.837*	0.820
most favorite	0.886*	0.913*	0.974*	0.999*	0.820
least favorite	0.888*	0.934*	1.015*	1.036*	0.824*
most watched	0.874*	0.864*	0.849*	0.835*	0.818
most characteristic	0.873*	0.862*	0.847*	0.837*	0.818
highest variance	0.874*	0.860*	0.842*	0.830*	0.819
NDCG@10					
random	0.793	0.793	0.794	0.793	0.795
most recent	0.792	0.795	0.792	0.792	0.791
most favorite	0.793	0.794	0.793	0.794	0.791
least favorite	0.794	0.793	0.792	0.793	0.793
most watched	0.794	0.792	0.792	0.790	0.792
most characteristic	0.794	0.792	0.793	0.793	0.794
highest variance	0.793	0.793	0.794	0.792	0.791

4.1 Feasibility of global data minimization

Tables 1 and 2, as well as Figure 1 present the performance of the k-NN and SVD recommendation algorithms for various minimization strategies and intensity (parameter n denotes the number of items from observational pools that were shown to the recommendation

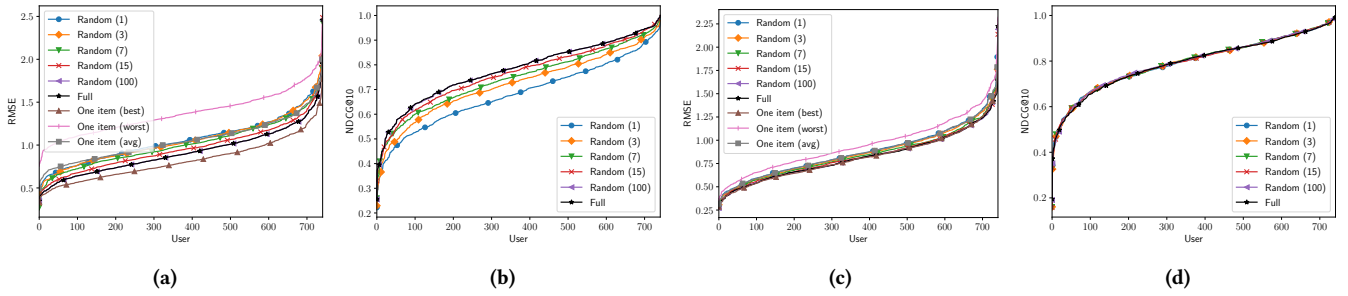


Figure 1: Sorted RMSE (a, c) and NDCG (b, d) values for all users when selecting random subsets of items of varying sizes as input to the kNN (a, b) and SVD (c, d) recommendation algorithms. Higher values on the y-axis in plots (a, c) are worse, while higher values on the y-axis in plots (b, d) are better. SVD is more robust to minimization than kNN, with aggressive minimization incurring low quality loss. While error increases as we minimize, the distribution of remains the same.

algorithm). The numbers show the RMSE and NDCG values of the minimized recommendations, averaged over all minimizing users.

For both recommendation algorithms, we observe that the recommendation quality decreases as we present the algorithm with less data to base personalized recommendations on. We attribute the few exceptions (e.g., the increase of RMSE between $n=3$ and $n=7$) to the inherent noisiness of data and effects of sampling strategies.

We would like to highlight two observations. First, the overall loss incurred by minimization is relatively low when compared to the variation of error across users — see Figure 1 for a visualization of sorted error values for all users in the minimization dataset for random minimization strategies. It is important not to overinterpret these results based on measures like RMSE, though. Ratings in recommendation datasets are often relatively homogenous in terms of absolute values: In the MovieLens dataset, for instance, they vary between 0.5 and 5 in 0.5 increments. Moreover, most users abstain from using extreme values in their ratings: In our system data, out of 10 different values in the rating scale, the three most used rating values of 3, 3.5, and 4, make 61% of all ratings.

Second, the distribution of error in the population remains the same even when the recommendations are based on minimized data. We observe that the shapes of the error value curves are similar for different minimization methods beyond random (effects similar to those in Figure 1). We exclude additional plots for lack of space.

4.1.1 Withheld data. While our experiments explicitly controlled the size of user interaction logs available to a recommendation algorithm, the data *withheld* from the algorithm can be substantial. On average, minimization with $n = \{1, 3, 7, 15, 100\}$ leads to 99%, 96%, 90%, 79%, 6% of data withheld from the recommendation algorithm, respectively. Note that this is not a comment about the total amount of data available to the system: In the setup we consider in this paper, the recommendation algorithm is trained on full data of 70% of users, which means that the effective percentage of the withheld data is lower.

4.2 Algorithm robustness to data minimization

We find that SVD is more robust to data minimization according to both quality measures. In case of RMSE, metric differences between the Full strategy and any other strategy and minimization

parameter n are lower for SVD than for kNN. This observation also holds for NDCG; moreover, the differences in NDCG between the performance of SVD on full data and minimized data are not significant (under a two-tailed t-test and $p < 0.01$ with the Bonferroni correction). Note that the SVD robustness result is partly explained by our experimental protocol—the minimized observed data of each test user is ‘folded in’ into the matrix one user at a time. While this approach is more computationally expensive than folding in all test users at once, the resulting decomposition is computed for a matrix where only one row is different from the full data condition. On top of that, the NDCG measure is not sensitive to differences in predicted rating values as long as the predicted ranking of items remains the same (which is likely to happen when the decomposed matrix is similar to the full data matrix). The lower minimization robustness of kNN can furthermore be explained by the fact that user similarities are computed over rating sets joint with other system users ($I_u \cap I_v$, see 3.2), and minimization thus leads to computing predictions over noisier neighbour sets.

4.2.1 Comparison to prior work. Note that these findings are consistent with prior work. First, Chow *et al.* [5] demonstrate that, for similarity-based recommendations, performance often does not differ after removing random data points. Further, different data removal strategies can improve or degrade predictive performance relative to random removal; in some cases, strategies can improve over the non-minimized predictions [5, Fig. 1].

Second, Wen *et al.* [43] analyzed performance decreases in a recommendation privacy scenario where users provide an algorithm with their recommendation data from the most recent N days. This filtering strategy is similar to the Most Recent minimization strategy we introduce in Sec. 3.5. Wen *et al.* showed that predictions of matrix-factorization-based methods are robust, with performance not degrading even when data is limited to ratings from the previous one to seven days and especially when the percentage of minimizing users is low [43, Fig. 2].¹

4.2.2 Factors influencing rating changes when minimizing for k-NN. Recall Eq. 1. What will influence the difference between an item prediction $\hat{r}_{ui}$ under the minimization condition and the prediction

¹The experimental protocol used in our paper maps to a setting in Wen *et al.* [43] where the percentage P of minimizing users is much lower than 0.25.

based on the full observational pool? Since the system data remains intact under our experimental setup, the values of r_{vi} will remain intact as well. The value of $\hat{r}_{ui}$ will be changed, though, when u 's relative similarity to other users changes. This might happen when:

- The set of nearest neighbors $N_i^k(u)$ changes and user u is placed in a different neighborhood for item i . The nearest neighbor summation of r_{vi} ratings happens over a different set of users v (even if the relative similarities to those users stay the same).
- The set of nearest neighbors $N_i^k(u)$ changes and user u is placed in a neighborhood where the relative similarities to other users $\text{sim}(u, v)$ are different (even if the neighbor rating values r_{vi} are the same).
- The set of nearest neighbors $N_i^k(u)$ stays the same but the similarity of u to other users within the neighborhood changes. Note that this is very likely to happen since the similarities will be computed over u 's minimized data.

While it is possible to enumerate these error contributing factors, analysis of how exactly they impact overall minimization error is challenging because the different dimensions (user similarity, neighborhoods, item popularity, etc.) all influence each other.

4.2.3 Factors influencing rating changes when minimizing for SVD. When will an item prediction $\hat{r}_{ui}$ under the minimization condition and the prediction based on the full observational pool? Note that latent item representations q_i and biases b_i will largely stay intact – during training, most updates to q_i 's and b_i 's will come from the data of the system users. The rating change will primarily be influenced by a change in the latent user representation p_u and bias b_u – during training, updates to these components will come from the latent factors of minimized observational items. Thus, we can expect biggest rating differences if the items in the minimized user profile don't reflect the full user profile. To examine the relative importance of p_u and b_u , we run minimization for recommendations generated using an unbiased SVD (removing μ , b_u , and b_i from Eq. 2). We find that errors incurred by minimization for this setup increase, suggesting that recommendation performance might be preserved by the bias terms when data is minimized.

4.3 Best and worst minimization strategies

4.3.1 Random minimization strategy. Figure 1 presents sorted RMSE (a, c) and NDCG (b, d) error values per user in the MovieLens dataset population, respectively, when minimizing data using random selection strategies. Unsurprisingly, on average, recommendation error increases as we observe fewer items. The error increase is, however, not substantial. There a number of factors that contribute to this effect. First, note that the random minimization strategy does not create random user profiles, but average user profiles, and the rating distributions over salient categories are likely to remain the same. Second, user profiles are of varying sizes and for some methods minimizing methods already access full observational pools. We tried to alleviate this effect by inclusion of users whose observational pools have at least 45 ratings. To understand these limitations better, we also plot the empirical lower bound on the error for predictions based on empty observational pools (non-personalized predictions based on the system data only). While the random minimization strategy performs reasonably well, there exist better and worse minimization strategies for both recommendation algorithms.

4.3.2 Strategies performing better than random minimization. For kNN recommendations, Most Favorite and Most Watched strategies perform better than Random. Movies users like most likely lead to highest contributions to user-user similarity, and thus the Most Favorite strategy tends to quickly place users in the right neighborhoods. Most Watched, by asking about the most rated movies, will quickly place users belonging to large clusters of popular movie watchers in the right neighborhood. Since there are many users with a taste for most popular movies, this strategy overall leads to a good global minimization performance.

4.3.3 Strategies performing worse than random minimization. For kNN recommendations, the Highest Variance selection strategy performs worse than the random selection strategy for the lowest n values ($n = 1, 3, 7$). One hypothesis is that the items selected by this strategy often have very high or very low ratings for a given user, causing this strategy to effectively interpolate between the performance of the Most Favorite and Least Favorite strategies. Whereas Most Favorite usually performs slightly better than random, Least Favorite often performs far worse, and when observed together explains why the Highest Variance strategy often performs worse than the Random selection strategy. We believe this is because the most characteristic score is inversely correlated with the most watched count.

For SVD recommendations, Most Favorite and Least Favorite strategies perform significantly worse than Random. We hypothesize that asking a user for ratings from just one side of their taste spectrum fails to populate all latent dimensions with relevant information. Moreover, since the most and least favorite items of a given user are likely correlated, asking for more items corroborates this effect by constructing an increasingly skewed user taste representations. This skew potentially leads to a reversal effect we have observed—Most Favorite and Least Favorite strategies initially decrease in performance as we increase n .

4.3.4 Other strategies. For kNN recommendations, Most Recent strategy performs on average worse than Random, likely due to the fact that the MovieLens-20M data was collected over a long period of time, yet our testing sample was random. Relatively bad performance of the Least Favorite strategy is related to insensitivity of standard recommendation algorithms to negative rating feedback; systems generally need to be tuned to be able to learn from negative ratings [11].

4.4 Differences between datasets

As described in Sec. 3.1, we run the the same experiments with two different datasets, using the Google Location dataset for validation. We observe the same trends in terms of the performance of different minimization strategies. One major difference in the results is that we observe similar robustness to minimization for KNN and SVD recommendations. We attribute this fact to the key difference between the two datasets—in Google Location dataset item ratings are sparser (20k vs. 150k unique items for a similar total number of users) thus minimization is less likely lead to overall change in similarities to other users.

5 PER-USER DATA MINIMIZATION

5.1 Feasibility of per-user data minimization

Figure 2 shows the error variation when the data is sorted only by the error value of the Full method - other error values correspond to users at the ranking positions determined by the sorting for Full. Note that the data plotted here is exactly the same as the data in Figure 1 — only the sorting differs. These results suggest that, while the distribution of error in the population across users remains largely similar irrespective of recommendation algorithm or minimization strategy (see Figure 1), *errors incurred to individuals can be substantial*. We observe this behavior for all tested minimization methods and recommendation algorithms, although the per-user variations are lower when minimizing for SVD recommendations.

This finding suggests that, for a fixed quality threshold, data can be less effectively minimized if the loss requirement applies to every individual as opposed to the population on average.

Since the error is not uniformly distributed, we dive deeper to try to understand which users are most impacted. The following sections analyze a number of user characteristics and their correlations with error deltas.

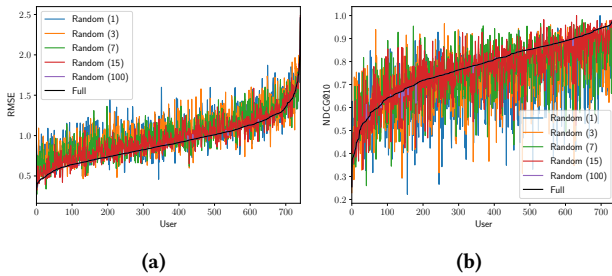


Figure 2: RMSE (a) and NDCG (b) variation over the population of users when selecting random subset of items of varying sizes as an input to the kNN algorithm. The underlying data presented here is the same as in Figure 1, but the data points are sorted by the y-axis value of the Full strategy only. Data points of other selection methods are unsorted and match the users at the ranking positions defined by the sorting of the Full strategy. This result shows that, while the overall quality loss is low and the error distribution remains the same, the quality loss for individuals can be substantial.

5.2 User characteristics vs. minimization error

We investigate whether the minimization errors (the difference in the quality metric when comparing the recommendations over minimized profiles and the recommendations over full profiles) are correlated with different user characteristics. For each user, we consider the following characteristics (measured over the user’s full profile, before minimization): (1) Number of ratings, (2) Average value of the ratings in a user’s profile, (3) Average popularity of items in a user’s profile (measured as the number of users in the system data who have rated a given item), (4) Profile diversity measured by the number of genres the movies in a user’s profile belong to, (5) Average similarity to all users in the system data, (6) Average similarity to the 30 most similar users in the system data.

5.2.1 Regression analysis. For each pair of recommendation algorithm and minimization strategy, we run an Ordinary Least Squares regression with the error delta as the dependent variable, and the above user characteristics as independent variables. Error delta is computed in two versions as: (i) $\Delta_3 = \text{RMSE}(3) - \text{RMSE}(\text{Full})$, and (ii) $\Delta_{15} = \text{RMSE}(15) - \text{RMSE}(\text{Full})$. We compute the coefficient of determination (R^2) to measure what proportion of variance in the dependent variable can be explained by the independent variables.

We find that the variance in neither Δ_3 nor Δ_{15} is well explained by the selected user variables, across recommendation and minimization strategies. For kNN and Δ_3 , we get the highest R^2 at 0.102 for the Most Recent strategy, followed by 0.0935 for the Most Characteristic System strategy, and 0.061 for the Least Favorite strategy. For kNN and Δ_{15} , R^2 values are even lower. For SVD and Δ_3 , we get the highest R^2 values for the Most and Least Favorite strategies, at 0.396 and 0.364, respectively. For SVD and Δ_3 , R^2 values follow similar trends.

5.2.2 A closer look. For a closer look into the complex dependencies between user characteristics, minimization strategies, recommendation algorithms, and minimization errors, we plot the most interesting cases in Figure 3.

Figure 3a shows the dependency between the number of ratings in a user’s full profile and the error delta (kNN+Random). The plot suggests that the smaller a user’s observational pool, the higher variation in the incurred minimization error. We conjecture that the reason for this effect is that sparse profiles with little data are likely to misrepresent true user tastes.

Figure 3b shows the dependency between a user’s average similarity to all users in the system data and the error delta (kNN+Random). We observe a similar trend – lower global similarity means higher variance in minimization error. However, the reason for this effect is likely different. Users who are similar to many system users are likely to end up in a neighborhood with accurate recommendations irrespective of which items are minimized out of their profiles.

Figure 3c shows the dependency between a user’s RMSE error for recommendations over the full observational pool the error delta (kNN+Random). We observe that lower RMSE values over the full data tend to imply higher error deltas, suggesting that users who are underserved by a system will be harmed the most when minimizing data.

Figures 3d and 3e reveal a curious observation about the dependency between the average value of ratings in a user profile and the error delta incurred by the Most and Least Favorite strategies for SVD. Users who tend to give lower movie ratings on average will receive worse results when minimizing using the Most Favorite strategy – likely because the movies they like the most will look like neutral movies when compared to the absolute values of ratings of other users. For a similar reason, though inverted, users who tend to give higher ratings on average will receive worse results when minimizing using the Least Favorite strategy. Figure 3f shows that for the Random strategy the effect is symmetric and less pronounced.

6 DATA MINIMIZATION VS. PRIVACY

The operational definitions of data minimization proposed in this paper, as shown in the experiments, will often lead to a decrease of

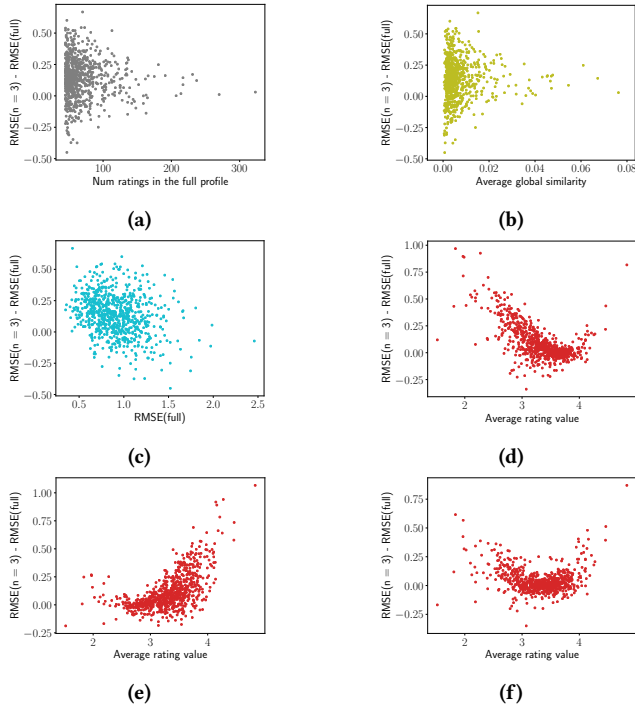


Figure 3: Selected cases of the dependencies between user characteristics and the RMSE delta (higher value means higher quality loss for minimization) incurred by minimization: (a) Number of ratings in the full profile, kNN, Random, (b) Average similarity to all users in the data, kNN, Random, (c) RMSE of recommendations over full profile, kNN, Random (d) Average rating value in the full profile, SVD, Most Favorite, (e) Average rating value in the full profile, SVD, Least Favorite, (f) Average rating value in the full profile, SVD, Random. In each of the plots, a dot corresponds to one minimizing user.

data collection. However, it is a feasible scenario that that each data point positively contributes to a system’s performance and data collection will not be decreased as a result. Rather than thinking of minimization as another computational definition of privacy, we look at data protection more broadly. For instance, the UK Information Commissioner’s Office defines data protection as ‘*the fair and proper use of information about people*’ [30]. Nevertheless, because of the potential decrease in the collected data, the proposed definitions of data minimization are related to different computational concepts of privacy. We briefly discuss some of these relationships. **Identifiability.** Presence of a unique combination of items in an anonymous user profile poses a deanonymization risk: if an attacker has the background knowledge that a user has rated these items, they can uniquely identify their profile and thus gain access to the rest of the user’s data. Analogous scenarios motivated the work on *k-anonymity* and related concepts [38]. One way of quantifying identifiability without access to external datasets is through a lower bound on the number of items an attacker would need to know to identify a user in a given dataset. More specifically, we compute, for

each user u , the minimum size of a subset of her ratings that does not exist in a profile of any other user: $\min_{I \in \mathcal{P}(I_u)} |I|$ s.t. $\forall v \neq u, I \not\subseteq I_v$. The higher the value of the above measure, the bigger the number of items an attacker would need to know to uniquely identify a user profile, and thus the lower the identifiability risk.

Table 3 presents the identifiability statistics for user profiles minimized using different strategies, averaged over all users. The results suggest that minimization strategies selecting items based on the characteristics of system data (Most Watched, Highest Variance, Most Characteristic) lead to lower profile identifiability than minimization methods based on an individual’s preferences (Most and Least Favorite). The most Recent strategy leads to the lowest identifiability across different values of the minimization parameter n . We conjecture this is because at a given time, many users rate the same new releases.

Profiling. Another computational privacy concept is that of profiling—collecting detailed topical profiles of users [2, 4, 44]. Should data minimization lead to decrease of collected data, it is likely that profiling risks also decrease. For instance, in our experiments, decreasing the number of movie ratings in all user’s profile to a maximum of 100 already reduces the average number of different genres in a user profile from 28.2 down to 25.1 according to the best strategy. **Other.** While decreasing the size of data might also help with other privacy dimensions, such as protection from inference [6] or differential privacy [7] in case aggregate data is released, analysis of these dimensions is more complex and might lead to removal of different data points.

Table 3: Identifiability (the minimum number of items necessary to uniquely identify a user) for user profiles minimized using different strategies, averaged over all users.

	n=3	n=7	n=15	n=100
random	2.02	1.89	1.76	1.55
most recent	1.91	1.79	1.71	1.55
most favorite	2.01	1.88	1.79	1.55
least favorite	1.92	1.81	1.71	1.55
most watched	2.28	2.33	2.00	1.57
most characteristic	1.99	2.00	2.00	1.57
highest variance	2.04	2.00	2.00	1.57

7 RELATED WORK

Interpreting GDPR principles in practice. The core contribution of this paper is in pointing out the gap between the current understanding of GDPR’s data minimization principle and the reality of personalization systems and proposing possible adequate re-interpretations. In this context, our work is related to other efforts to translate GDPR’s principles into data science practice. Prior work in this space has explored practical challenges behind revoking consent to data processing[32, 40], and explored what *the right to be forgotten* [25] means in practice. Recent work proposes practical solutions for removing data points from trained machine learning models in case an individual included in the training data requests deletion [15]. The right to explanation[23], requiring service providers to be able to explain algorithmic decisions and results to their users, motivated the active are of explainability and

transparency. Another line of work analyzes changes to the online ecosystem incurred by GDPR, including the presence of consent notices [39], or tracking scripts [35, 36].

Privacy. As discussed in Sec. 6, data minimization is related to some of the computational concepts of privacy. In the context of personalized search, many works proposed mechanisms for perturbing user search logs while preserving the search quality, including mixing and merging queries into synthetic profiles [2, 9], grouping user profiles [28], or splitting them [4, 44, 47]. Privacy has also been interpreted as a probabilistic guarantee on data retention [34]. To preserve the privacy of recommender system users, it has been proposed to prevent the collection of ratings locally if they are predicted to lead to privacy loss [16], or to store the ratings of different users intermingled [2]. Research in privacy-preserving information retrieval [45] moreover investigates problems related to search log anonymization [46], or the relation between user behavior and privacy attitudes [48].

Performance of recommender systems under varying conditions. Analyses we perform in this paper are related to a line of work analyzing the success and failure of recommender systems under changing conditions. Ekstrand et al. [8] analyze data factors that cause different recommendation algorithms to fail. Chow et al. [5] propose techniques to estimate the contributions of different data points to the overall recommendation quality. Vincent et al. [42] propose 'data strikes' as a form of collective action where users protest by withholding their data from recommendation provider. Note that, while the goal of data strikes is to limit availability of data to *reduce* recommendation performance, the goal of performance-based data minimization is to limit availability of data while *preserving* recommendation performance. Wen et al. [43] analyzed performance decrease in a recommendation privacy scenario where users provide an algorithm with their recommendation data from the most recent N days.

Relation to other disciplines We believe that further work on data minimization would lead to synergies not only with the legal community, but also with other computer science subdisciplines. While the focus of data minimization is on minimizing features rather than data points, the problem is related to works studying the relationships between training examples and algorithmic performance. This abstract description includes, for instance, the problems of data influence [24], data valuation [14], active learning [1], or budgeted learning [27].

8 DISCUSSION AND CONCLUSIONS

8.1 Summary of the findings

In this paper, we have identified a lack of a homogeneous interpretation of the GDPR's purpose limitation and data minimization principles in the domain of personalization systems. We argue that these systems do not necessarily need to collect user data, but that they do so in order to improve the quality of the results. Thus, we propose two performance-based interpretations of the data minimization principle that tie the limitations of data collection to quality metrics. The first interpretation focuses on the global average algorithm performance, while the second focuses on the local per-user minimum performance.

We found SVD (FunkSVD with user and item biases) to be more robust to minimization than kNN user-user collaborative filtering across different minimization strategies. Among the minimization strategies, we found the random strategy to perform well, likely due to the fact that it preserves average user characteristics. However, for each recommendation algorithm, it is possible to find strategies that perform better or worse than random.

While the results suggest global data minimization can be quite successful (in some cases we can withhold as much as 90% of the user data incurring RMSE loss as low as 0.025), we show that quality difference can be substantial for individual users. Furthermore, our analysis with Ordinal Least Squares regression shows that the variation in individual-level error is not well explained by standard user features. The complex interaction between the individual-level error and recommendation algorithms, minimization strategies, system data, and individual data, require further study, also from a legal perspective. Indeed, further research should evaluate the desirability of both approaches, considering that, on the one hand, the GDPR requires that each data processing operation be examined on its own merits, yet on the other purpose limitation or data protection by design and by default ought to be evaluated from a global perspective.

8.2 Potential negative impacts

Based on our observations about varying user-level errors, it is plausible that data minimization hurts marginalized groups, in particular if those groups form a minority of the data—the members of majority population will be well served with just a few features (because there is sufficient statistical support), while minority populations will *need* to provide more features to get service of comparable quality. A scenario like this would further harm marginalized populations through decreased data protection.

Furthermore, our analysis assumes service providers have a collection of background data to base personalization on (purchased or collected from markets that are not legally obliged to data minimization). Companies might also need personal data to develop new services. In this work, we did not consider such provider costs.

8.3 Challenges for data minimization

While this paper enhances our understanding of what performance-based data minimization means in practice, a number of challenges emerge. Practical minimization mechanisms would not be able to measure quality loss directly, nor easily adapt selection mechanisms to each user if necessary without access to candidate user data. To support minimization, we need to design new protocols for user-system interaction, and new learning mechanisms that select data while respecting specific minimization requirements. Last but not least, further interdisciplinary work with the legal community is necessary to develop data minimization interpretations that are verifiable and viable, both legally and computationally.

ACKNOWLEDGMENTS

We wish to thank Solon Barocas for the discussion that helped shape the interdisciplinary direction of this work.

REFERENCES

- [1] Philip Bachman, Alessandro Sordani, and Adam Trischler. 2017. Learning algorithms for active learning. In *Proceedings of the 34th International Conference on Machine Learning-Volume 70*. 301–310.
- [2] Asia J Biega, Rishiraj Saha Roy, and Gerhard Weikum. 2017. Privacy through solidarity: A user-utility-preserving framework to counter profiling. In *Proceedings of the 40th International ACM SIGIR Conference on Research and Development in Information Retrieval*. ACM, 675–684.
- [3] Reuben Binns and Valeria Gallo. 2019. *Data minimisation and privacy-preserving techniques in AI systems*. Retrieved May 26, 2020 from <https://ico.org.uk/about-the-ico/news-and-events/ai-blog-data-minimisation-and-privacy-preserving-techniques-in-ai-systems/>
- [4] Gang Chen, He Bai, Lidan Shou, Ke Chen, and Yunjun Gao. 2011. UPS: efficient privacy protection in personalized web search. In *Proceedings of the 34th international ACM SIGIR conference on Research and development in Information Retrieval*. ACM, 615–624.
- [5] Richard Chow, Hongxia Jin, Bart Knijnenburg, and Gokay Saldamli. 2013. Differential data analysis for recommender systems. In *Proceedings of the 7th ACM conference on Recommender systems*. 323–326.
- [6] Flávio du Pin Calmon and Nadia Fawaz. 2012. Privacy against statistical inference. In *2012 50th annual Allerton conference on communication, control, and computing (Allerton)*. IEEE, 1401–1408.
- [7] Cynthia Dwork. 2008. Differential privacy: A survey of results. In *International conference on theory and applications of models of computation*. Springer, 1–19.
- [8] Michael Ekstrand and John Riedl. 2012. When recommenders fail: predicting recommender failure for algorithm selection and combination. In *Proceedings of the sixth ACM conference on Recommender systems*. ACM, 233–236.
- [9] Sedigheh Eslami, Asia J Biega, Rishiraj Saha Roy, and Gerhard Weikum. 2017. Privacy of hidden profiles: Utility-preserving profile removal in online forums. In *Proceedings of the 2017 ACM on Conference on Information and Knowledge Management*. ACM, 2063–2066.
- [10] Michèle Finck and Frank Pallas. 2019. They Who Must Not Be Identified-Distinguishing Personal from Non-Personal Data Under the GDPR. *Max Planck Institute for Innovation & Competition Research Paper* 19-14 (2019).
- [11] Evgeny Frolov and Ivan Oseledets. 2016. Fifty Shades of Ratings: How to Benefit from a Negative Feedback in Top-N Recommendations Tasks. In *Proceedings of the 10th ACM Conference on Recommender Systems*. ACM, 91–98.
- [12] Simon Funk. 2006. Netflix update: Try this at home (December 2006). URL <http://sifter.org/~simon/journal/20061211.html> (2006).
- [13] 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 *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*. 265–271.
- [14] Amirata Ghorbani and James Zou. 2019. Data Shapley: Equitable Valuation of Data for Machine Learning. In *International Conference on Machine Learning*. 2242–2251.
- [15] Antonio Ginart, Melody Guan, Gregory Valiant, and James Zou. 2019. Making AI Forget You: Data Deletion in Machine Learning. *NIPS* (2019).
- [16] Rachid Guerraoui, Anne-Marie Kermarrec, and Mahsa Taziki. 2017. The Utility and Privacy Effects of a Click. In *Proceedings of the 40th International ACM SIGIR Conference on Research and Development in Information Retrieval*. ACM, 665–674.
- [17] Isabelle Guyon and André Elisseeff. 2003. An introduction to variable and feature selection. *Journal of machine learning research* 3, Mar (2003), 1157–1182.
- [18] A. Halevy, P. Norvig, and F. Pereira. 2009. The Unreasonable Effectiveness of Data. *IEEE Intelligent Systems* 24, 2 (March 2009), 8–12.
- [19] F Maxwell Harper and Joseph A Konstan. 2016. The movielens datasets: History and context. *Acm transactions on interactive intelligent systems (tiis)* 5, 4 (2016), 19.
- [20] Ruining He, Wang-Cheng Kang, and Julian McAuley. 2017. Translation-based recommendation. In *Proceedings of the Eleventh ACM Conference on Recommender Systems*. 161–169.
- [21] Nicolas Hug. 2017. Surprise, a Python library for recommender systems. <http://surpriselib.com>.
- [22] Kalervo Järvelin and Jaana Kekäläinen. 2002. Cumulated gain-based evaluation of IR techniques. *TOIS* 20, 4 (2002), 422–446.
- [23] Margot E Kaminski. 2019. The right to explanation, explained. *Berkeley Tech. LJ* 34 (2019), 189.
- [24] Pang Wei Koh and Percy Liang. 2017. Understanding black-box predictions via influence functions. In *Proceedings of the 34th International Conference on Machine Learning-Volume 70*. 1885–1894.
- [25] Bert-Jaap Koops. 2011. Forgetting footprints, shunning shadows: A critical analysis of the right to be forgotten in big data practice. *SCRIPTed* 8 (2011), 229.
- [26] Andreas Krause and Eric Horvitz. 2010. A utility-theoretic approach to privacy in online services. *Journal of Artificial Intelligence Research* 39 (2010), 633–662.
- [27] Daniel J Lizotte, Omid Madani, and Russell Greiner. 2002. Budgeted learning of naive-bayes classifiers. In *Proceedings of the Nineteenth conference on Uncertainty in Artificial Intelligence*. 378–385.
- [28] Pol Mac Aonghusa and Douglas Leith. 2018. 3PS-Online Privacy through Group Identities. *arXiv preprint arXiv:1811.11039* (2018).
- [29] Arvind Narayanan and Vitaly Shmatikov. 2008. Robust de-anonymization of large sparse datasets. In *2008 IEEE Symposium on Security and Privacy*, SP. 111–125.
- [30] UK Information Commissioner's Office. 2018. *Guide to Data Protection. Some basic concepts*. Retrieved Jan 22, 2020 from <https://ico.org.uk/for-organisations/guide-to-data-protection/introduction-to-data-protection/some-basic-concepts/>
- [31] UK Information Commissioner's Office. 2018. *Guide to the General Data Protection Regulation (GDPR). Principle (c): Data minimisation*. Retrieved Jan 22, 2020 from <https://ico.org.uk/for-organisations/guide-to-data-protection/guide-to-the-general-data-protection-regulation-gdpr/principles/data-minimisation/>
- [32] Eugenia Politou, Efthimios Alepis, and Constantinos Patsakis. 2018. Forgetting personal data and revoking consent under the GDPR: Challenges and proposed solutions. *Journal of Cybersecurity* 4, 1 (2018), tty001.
- [33] Protection Regulation. 2016. REGULATION (EU) 2016/679 OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL. *Official Journal of the European Union* (2016).
- [34] Adish Singla, Eric Horvitz, Ece Kamar, and Ryan White. 2014. Stochastic privacy. In *Twenty-Eighth AAAI Conference on Artificial Intelligence*.
- [35] Konstantinos Solomos, Panagiotis Ilia, Sotiris Ioannidis, and Nicolas Kourtellis. 2019. Clash of the Trackers: Measuring the Evolution of the Online Tracking Ecosystem. *arXiv preprint arXiv:1907.12860* (2019).
- [36] Jannick Sørensen and Sokol Kosta. 2019. Before and After GDPR: The Changes in Third Party Presence at Public and Private European Websites. In *The World Wide Web Conference*. ACM, 1590–1600.
- [37] C. Sun, A. Shrivastava, S. Singh, and A. Gupta. 2017. Revisiting Unreasonable Effectiveness of Data in Deep Learning Era. In *2017 IEEE International Conference on Computer Vision (ICCV)*. 843–852.
- [38] Latanya Sweeney. 2002. k-anonymity: A model for protecting privacy. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems* 10, 05 (2002), 557–570.
- [39] Tobias Urban, Dennis Tatang, Martin Degeling, Thorsten Holz, and Norbert Pohlmann. 2019. A Study on Subject Data Access in Online Advertising After the GDPR. In *Data Privacy Management, Cryptocurrencies and Blockchain Technology*. Springer, 61–79.
- [40] Christine Utz, Martin Degeling, Sascha Fahl, Florian Schaub, and Thorsten Holz. 2019. (Un) informed Consent: Studying GDPR Consent Notices in the Field. In *ACM SIGSAC Conference on Computer and Communications Security (CCS '19)*.
- [41] Vladimir Vapnik. 1992. Principles of risk minimization for learning theory. In *Advances in neural information processing systems*. 831–838.
- [42] Nicholas Vincent, Brent Hecht, and Shilad Sen. 2019. "Data Strikes": Evaluating the Effectiveness of a New Form of Collective Action Against Technology Companies. In *The World Wide Web Conference*. ACM, 1931–1943.
- [43] Hongyi Wen, Longqi Yang, Michael Sobolev, and Deborah Estrin. 2018. Exploring recommendations under user-controlled data filtering. In *Proceedings of the 12th ACM Conference on Recommender Systems*. ACM, 72–76.
- [44] Yabo Xu, Ke Wang, Benyu Zhang, and Zheng Chen. 2007. Privacy-enhancing personalized web search. In *Proceedings of the 16th international conference on World Wide Web*. ACM, 591–600.
- [45] Hui Yang, Ian Soboroff, Li Xiong, Charles LA Clarke, and Simson L Garfinkel. 2016. Privacy-preserving ir 2016: Differential privacy, search, and social media. In *Proceedings of the 39th International ACM SIGIR conference on Research and Development in Information Retrieval*. 1247–1248.
- [46] Sicong Zhang, Hui Yang, and Lisa Singh. 2016. Anonymizing query logs by differential privacy. In *Proceedings of the 39th International ACM SIGIR conference on Research and Development in Information Retrieval*. 753–756.
- [47] Yun Zhu, Li Xiong, and Christopher Verdery. 2010. Anonymizing user profiles for personalized web search. In *Proceedings of the 19th international conference on World wide web*. ACM, 1225–1226.
- [48] Steven Zimmerman, Alistair Thorpe, Chris Fox, and Udo Kruschwitz. 2019. Investigating the Interplay Between Searchers' Privacy Concerns and Their Search Behavior. In *Proceedings of the 42nd International ACM SIGIR Conference on Research and Development in Information Retrieval*. 953–956.