
The Data Minimization Principle in Machine Learning

Prakhar Ganesh
McGill University & Mila
prakhar.ganesh@mila.quebec

Cuong Tran
University of Virginia & Dyania Health
cuong@dyaniahealth.com

Reza Shokri
National University of Singapore
reza@comp.nus.edu.sg

Ferdinando Fioretto
University of Virginia
fioretto@virginia.edu

Abstract

The principle of *data minimization* aims to reduce the amount of data collected, processed or retained to minimize the potential for misuse, unauthorized access, or data breaches. Rooted in privacy-by-design principles, data minimization has been endorsed by various global data protection regulations. However, its implementation remains a challenge due to the lack of a rigorous formulation. This paper addresses this gap and introduces an optimization framework to operationalize the legal definitions of data minimization. It adapts several optimization algorithms to perform data minimization and conducts a comprehensive evaluation in terms of their compliance with minimization objectives and their impact on user privacy. Our analysis underscores the mismatch between the privacy expectations of data minimization and the actual privacy benefits, emphasizing the need for approaches that account for multiple facets of real-world privacy risks.

1 Introduction

The proliferation of data-driven systems and machine learning (ML) applications escalates a number of privacy risks, including those related to unauthorized access to sensitive information [18, 29]. In response, international data protection frameworks like the European General Data Protection Regulation (GDPR), the California Privacy Rights Act (CPRA), and the Brazilian General Data Protection Law (LGPD) have adopted *data minimization* as a key principle to mitigate these risks [2].

At its core, the data minimization principle requires organizations to *collect, process, and retain only personal data that is adequate, relevant, and limited to what is necessary for specified objectives* (see Table 1 for further details). It’s grounded in the expectation that not all data is essential for the objective and, instead, contributes to a heightened risk of information leakage [11, 19, 23, 26]. However, despite its legal significance and endorsement by global data protection regulations, the data minimization principle lacks a mathematical formalization suitable for real-world ML applications. In particular, as reviewed in §B, the current discourse on data minimization practices often overlooks two crucial aspects: (1) the individualized nature of minimization (e.g., information that is unimportant for an individual may be critical for another) and (2) its intrinsic link to data privacy.

Contributions. To overcome these limitations, this paper introduces a formal framework to operationalize data minimization in ML as an optimization problem while being faithful to its legal notion (§§2, 3); adapts and evaluates various optimization algorithms to solve this problem (§4); and analyzes their compatibility with real-world privacy (§5). In particular, we seek to answer a critical question: “Do current data minimization requirements in various regulations genuinely meet privacy expectations in legal frameworks?” Our evaluations reveal that the answer is, unfortunately, no. While being an implicit intention, the requirements of data minimization are not necessarily aligned

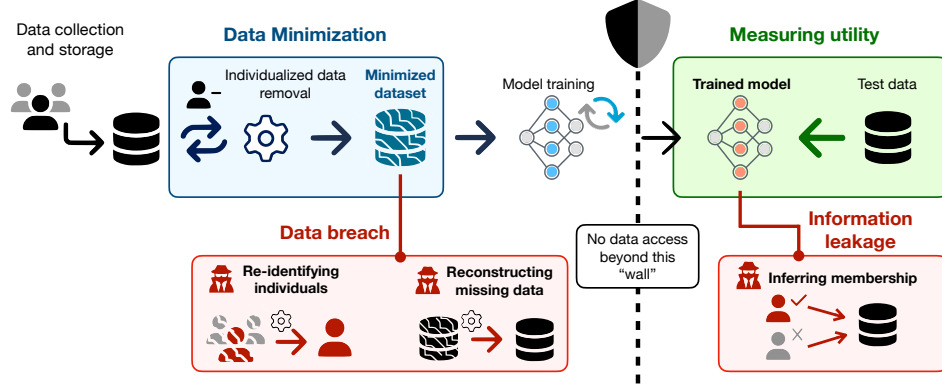


Figure 1: Overview of our framework to study data minimization in a real-world ML pipeline. Behind the data access “wall”, we highlight the formalization of data minimization and quantify the risks of a data breach. Outside the “wall” with no direct data access, we establish the data minimization objectives using utility measurement and study information leakage from the trained model.

with the risk of reconstruction and re-identification and thus may not provide the expected privacy protection. In summary, this paper makes the following contributions:

1. It examines various global data protection regulations and provides the first formalization of data minimization for ML tasks that faithfully incorporates the individualized nature of minimization.
2. It compares several algorithms to solve the optimization problem derived from data minimization and conducts an extensive evaluation focusing on key characteristics of the minimized datasets, including emergent individualization and multiplicity.
3. Through extensive evaluations using privacy attacks on the minimized data, it assesses the shortcomings of current regulations of data minimization in meeting the implicit privacy expectations.
4. It proposes simple yet effective modifications to the data minimization algorithms and demonstrates that attention to privacy during minimization can provide better trade-offs between user privacy and downstream utility.

This work aims to lay down a path for future research for developing privacy-preserving ML systems in compliance with the legal requirements of data minimization.

2 The data minimization principle

To faithfully develop the proposed framework (illustrated in Figure 1), we inspect the data minimization principle from six global data protection regulations. Building on existing work interpreting the legal language [2, 3, 20, 21, 23, 28], we focus instead on operationalizing data minimization, grounded in the common foundations of data privacy and security principles shared by these regulations. Relevant extracts are summarized in Table 1, with more details and an in-depth discussion of our interpretation delegated to the appendix (§A). The table highlights two core pillars of data minimization: (i) *purpose limitation*, and (ii) *data relevance*.

- **Purpose limitation** : Regulations mandate that data be collected for a *legitimate, specific and explicit purpose* (LGPD, Brazil) and prohibit using the collected data for any other incompatible purpose from the one disclosed (CPRA, USA). Thus, data collectors must define a clear, legal objective before data collection and use the data solely for that objective. In machine learning, this *purpose* translates to collecting data solely for training models to achieve optimal performance, where the performance metric needs to be clearly defined [2].
- **Data relevance** : Regulations like the GDPR require that all collected data *be adequate, relevant, and limited to what is necessary for the purposes it was collected for*. In other words, data minimization aims to remove data that does not serve the purpose defined. In ML contexts, this translates to retaining only data that contributes to the performance of the model. We interpret this requirement as an optimization objective, minimizing data size while maintaining model performance, under a binary interpretation that the data is either collected or not [3, 21, 23, 28].

General Data Protection Regulation (GDPR), Europe	gdpr-info.eu/
Article 5(1)(b): Personal data shall be collected for <i>specified, explicit and legitimate purposes and not further processed in a manner that is incompatible with those purposes</i> ;	
Article 5(1)(c): Personal data shall be <i>adequate, relevant and limited to what is necessary in relation to the purposes for which they are processed</i> .	
California Privacy Rights Act (CPRA), USA	cppa.ca.gov/
Section 1798.100 (a)(1) & (a)(2): [...] A business shall <i>not collect additional categories of (sensitive) personal information or use (sensitive) personal information collected for additional purposes that are incompatible with the disclosed purpose for which the (sensitive) personal information was collected without providing the consumer with notice consistent with this section</i> .	
General Personal Data Protection Law (LGPD), Brazil	lgpd-brazil.info/
Article 6: Activities of processing of personal data shall be subject to the following principles, I: <i>processing done for legitimate, specific and explicit purposes of which the data subject is informed, with no possibility of subsequent processing that is incompatible with these purposes</i> ; II: <i>compatibility of the processing with the purposes communicated to the data subject</i> , in accordance with the context of the processing; III: <i>limitation of the processing to the minimum necessary to achieve its purposes</i> , covering data that are <i>relevant, proportional and non-excessive in relation to purposes of the data processing</i> [...];	

Table 1: Excerpts from various data protection regulations on the data minimization principle, highlighting **purpose limitation** and **data relevance**.

2.1 Implications in practice

Building on the principles of purpose limitation and data relevance, this section examines three practical implications of data minimization as framed in the reviewed regulations: **(1) Individualized nature of data minimization:** Notice that different individuals may require different amounts and types of information to fulfill a given purpose. For instance, in a loan approval scenario, personal data such as age, income, and job history are crucial. However, the importance of additional data, like medical history, can vary significantly among applicants. This variability underscores that data relevance—and therefore redundancy—is not universal but contextually and individually defined. **(2) Data minimization is data dependent:** Regulations describe the term “data collection” as acquiring data from entities like government agencies, which typically involves selecting relevant data from an already gathered large dataset. For instance, census data is often used for various analysis tasks. This differs from *field* data collection, which lacks flexibility for individualized minimization adjustments. Our framework focuses on the former interpretation, wherein minimization is contingent on the data itself (as shown by the two distinct stages of “Data collection from individuals” and “Data minimization” in Figure 1). **(3) Privacy expectations through minimization:** There is an implicit expectation of privacy through minimization in data protection regulations [11, 16, 28]. *However, this expectation overlooks a crucial aspect of real-world data—the inherent correlations among various features.* Information about individuals is rarely isolated, thus, merely removing or not collecting data, may allow for confident reconstruction inference [10], as we will show in our work.

3 A formal framework for data minimization in ML

This section presents a data minimization framework for ML that aligns with the regulatory objectives discussed above. Successively, we examine the tension between these minimization objectives and their privacy implication and define various threat models to quantify real-world privacy risks.

3.1 Data minimization as optimization

Consider a dataset $D = \{\mathbf{X}, \mathbf{Y}\} = \{(\mathbf{x}_i, y_i)\}_{i=1}^n$, drawn i.i.d. from an unknown distribution $D \sim (\mathcal{X}, \mathcal{Y})$. Therein, $\mathbf{x}_i \in \mathbf{X}$ is a p -dimensional feature vector and $y_i \in \mathbf{Y}$ is the output label. Given a non-negative *loss function* $\ell : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}_+$, the goal of learning is to learn a predictor $f_\theta : \mathcal{X} \rightarrow \mathcal{Y}$, with real-valued parameters θ that minimizes the empirical risk function (i.e., ERM):

$$\hat{\theta} = \arg \min_{\theta} J(\theta; \mathbf{X}, \mathbf{Y}) \quad \text{where} \quad J(\theta; \mathbf{X}, \mathbf{Y}) = \frac{1}{n} \sum_{i=1}^n \ell(f_\theta(\mathbf{x}_i), y_i). \quad (1)$$

Data minimization can be modelled as minimizing the size of the dataset D by removing entries from the features $\mathbf{x}_i$ (data relevance) while also retaining similar performance as it would have on the complete dataset (purpose). This objective can be formalized in the following bi-level optimization:

$$\underset{\mathbf{B} \in \{\perp, 1\}^{n \times p}}{\text{Minimize}} \quad \|\mathbf{B}\|_1 \quad \text{s.t. :} \quad J(\hat{\theta}; \mathbf{X}, \mathbf{Y}) - J(\theta; \mathbf{X}, \mathbf{Y}) \leq \alpha \quad (2a)$$

$$\hat{\theta} = \arg \min_{\theta} \frac{1}{n} \sum_{i=1}^n \ell \left(f_{\theta}(\mathbf{x}_i \odot \mathbf{B}_i), y_i \right) \quad (2b)$$

where $\mathbf{B}$ is an $n \times p$ binary matrix, which we call the *minimization matrix*, taking values in the set $\{\perp, 1\}$, and the ℓ_1 -norm of $\mathbf{B}$, i.e., $\|\mathbf{B}\|_1$, is simply the count of 1 in the minimization matrix. Here, $\perp$ represents the concealment or removal of a value, i.e., $\forall v \in \mathbb{R} : v \times \perp = \perp$, and $\alpha \geq 0$ is a parameter which bounds the permitted drop in model quality due to data minimization. While we define ERM as our learning objective and the difference in loss on the original dataset as the quality drop (equation 2a), these choices can be easily generalized.

The minimized input features $\mathbf{X}'$ are defined as the element-wise product of the original features $\mathbf{X}$ and the minimization matrix $\mathbf{B}$, i.e., $\mathbf{X}' = \mathbf{X} \odot \mathbf{B}$. Note that, data minimization produces features in the space $\mathcal{X} \cup \{\perp\}$. While some learning algorithms can handle missing values ($\perp$), data imputation is needed in other cases. A discussion of imputation methods is delegated to §4.1.

The optimization problem in equation 2 defines an operational method to remove entries from the feature set $\mathbf{X}$ in an *individualized* manner, while adhering to pre-specified quality drop requirements on the original dataset (equation 2a), for the final model trained on the minimized dataset (equation 2b). While this optimization captures the legal formulation of data minimization, it is however intractable to solve in practice. A discussion of tractable approximate algorithms is delegated to §4 and we discuss next the privacy implications of this data minimization process.

3.2 Privacy leakage

The implicit objective of data minimization is to enhance privacy. Thus, the paper focuses on assessing the minimized data using several threat models (Figure 1), to measure inference attacks under various adversaries.

Note that the expectation of data privacy from minimization in various regulations is focused on the events of a data breach, i.e., situations in which the adversary has access to the minimized data. For this reason, we primarily focus on reconstruction and re-identification risks. This is different from conventional private ML settings [9, 22], which focuses on information leakage due to ML inference (i.e., outside the "wall" in Figure 1). Nonetheless, we also examine how data minimization privacy promise may hold under membership inference attacks and delegate this analysis to §E.

Reconstruction Risk (RCR). Real-world datasets often exhibit underlying associations between various features, making it possible to reconstruct minimized data [10]. Reconstruction attacks thus aim to recover missing information from a target dataset. Given the minimized dataset $\mathbf{X}'$, the attacker's goal is to generate their reconstruction $\mathbf{X}^R$ of the original set of features $\mathbf{X}$. The ReConstruction Risk (RCR) can be evaluated by measuring the similarity between the original features $\mathbf{x}_i \in \mathbf{X}$ and the reconstructed features $\mathbf{x}_i^R \in \mathbf{X}^R$, computed using a gaussian kernel with $\sigma = 1$ [27] as:

$$\text{RCR} = \frac{1}{n} \sum_{i=1}^n e^{-\|\mathbf{x}_i - \mathbf{x}_i^R\|_2}. \quad (3)$$

If required, the reconstruction risk metric can also be adjusted further to prioritize the reconstruction of certain features over others by introducing appropriate weights to the similarity measurement.

Re-identification Risk (RIR). Data breaches often lead to re-identification of individuals using partial or anonymized data matched with an auxiliary dataset $\mathbf{X}^A$. The success of these re-identification attacks can be measured by the mean reciprocal rank (MRR) scores. For each data point $\mathbf{x}_i^A$ in $\mathbf{X}^A$, the adversary ranks data points in the minimized dataset $\mathbf{X}'$ based on the likelihood of matching identities. The Re-Identification Risk (RIR) is calculated as the average MRR score, assigning a

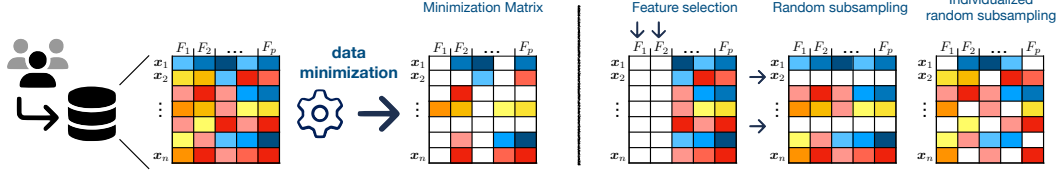


Figure 2: **(Left)** Data minimization as a set of individualized binary decisions, visualizing the minimization matrix B . **(Right)** Three baseline algorithms for data minimization.

score of 1 for a correct match and 0 otherwise, defined by the formula:

$$\text{RIR} = \frac{1}{n} \sum_{i=1}^n \frac{1}{\text{index}(i, \text{rank}(\mathbf{x}_i^A, \mathbf{X}'))}. \quad (4)$$

In this, $\text{rank}(\mathbf{x}_i^A, \mathbf{X}')$ is the adversary’s predicted ranking, and $\text{index}(e, A)$ gives the index of element e in array A . It assumes a one-to-one match between auxiliary and minimized datasets.

4 Data minimization and utility

Having interpreted data minimization for ML as a bilevel optimization problem, this section outlines how to operationalize data minimization, with three strong baselines and three algorithms from the bi-level optimization literature to address the challenge formulated in equation 2. Note that although the data minimization principle aims to optimize dataset size while maintaining model quality (see Constraint 2a), these algorithms adopt a dual approach, optimizing model quality when trained on minimized data under a given sparsity constraint $\|\mathbf{B}\|_1 \leq k$. This dual approach allows these algorithms to find a sparsity parameter k that meets the desired α -drop in performance.

Baseline techniques. The experiments implement three baseline methods, depicted pictorially in Figure 2: (1) feature selection, (2) random subsampling, and (3) individualized random subsampling. *Feature selection* [5] employs a breadth-based strategy that identifies and minimizes less important features within the dataset. *Random subsampling* is a depth-based strategy that randomly selects a subset of data points, thereby reducing the dataset size by excluding specific rows. *Individualized random subsampling* further tailors this approach by randomly selecting specific entries (feature, sample) for each individual, aiming to reduce dataset size in a more personalized manner. While these baselines help to assess model quality degradation, they do not fully comply with legal standards of data minimization, which require adherence to principles of purpose and data relevance, as discussed earlier and further expanded in the appendix (§C).

Data minimization algorithms. We briefly introduce three classes of algorithms adopted to solve the optimization posed in Problem 2:

1. *Approximating the target utility* [14], which focuses on finding an approximate closed-form solution to the lower-level problem, thus simplifying the original optimization.
2. *Modelling the target utility* [32], which instead attempts to model the lower-level mapping and estimate it online again without solving the optimization.
3. *Evolutionary algorithms* [25], which trade the advantage of carrying no assumption with slow convergence and high computational demand.

Extensive details about these algorithms, their theoretical underpinning, and comparative analysis can be found in the appendix (§D). Throughout our experimental evaluation, baseline methods are depicted with blue colours while the optimization algorithms with red colours.

4.1 Experiment setup

Datasets. Our evaluations focus on classification for both tabular data (marked with $\diamond$) and image data ($\heartsuit$), to cover a diverse set of distributions and modalities. In this section, we use **(i) the bank marketing dataset** [17], taken from a telemarketing campaign with financial information of 11,162 customers, and **(ii) the handwritten digits dataset** [33], containing a total of 1,797 handwritten digit images across 10 different classes (i.e., digits). Additional results on other datasets, including text modality, can be found in the appendix (§F).

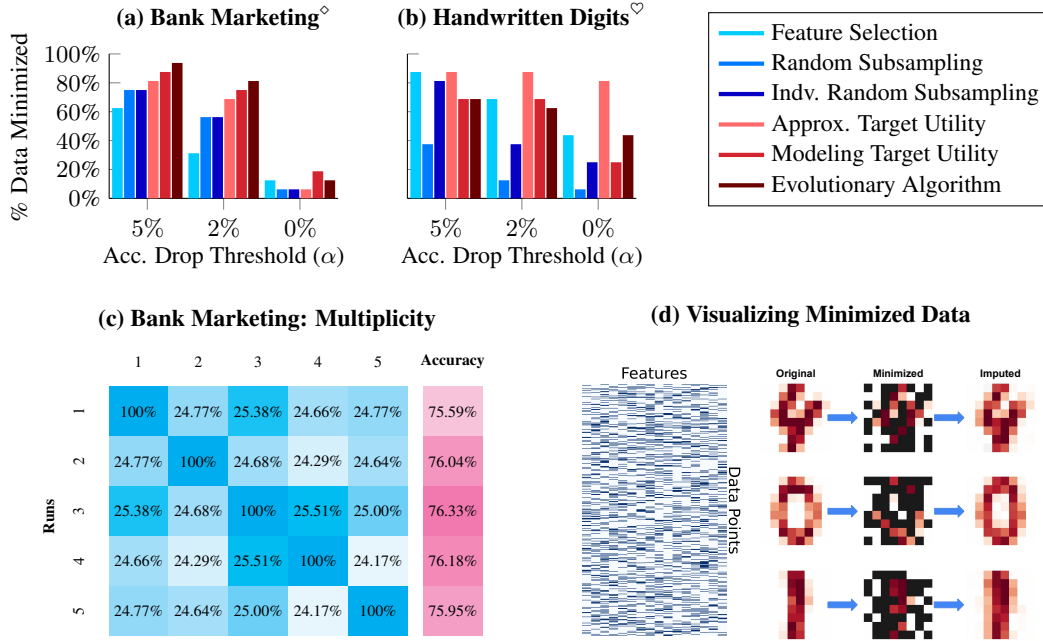


Figure 3: **(a, b)** Percentage of dataset minimized under different accuracy drop thresholds (α). **(c)** Pairwise overlap between minimized datasets across five runs with changing randomness. **(d)** Visualizing minimized datasets across modalities showing: (Left) emergent individualized minimization in a compressed view of minimized bank dataset; and (Right) misalignment of minimization with privacy as imputed images have reconstructed missing information despite minimization.

Data splits. Datasets are split equally into private and public data. The public data is used as the test data, to calculate statistics for data imputation, and train reference models for membership inference. The private data is further split equally into members and non-members. The members are our training data $\mathbf{X}$, i.e., also the data that will be minimized in our pipeline. The non-members are only used for evaluating membership inference attacks (details in §E).

Learning setup and data imputation. The experiments use logistic regression for the tabular dataset, and a fully connected neural network for the handwritten digits dataset. Both datasets are normalized using MinMaxScaler. Data imputation is performed under the assumption of a multivariate Gaussian distribution, utilizing the mean and covariance matrix of the public data to fill missing values with the mean of the conditional distribution [13].

4.2 Efficacy of data minimization

First, this section assesses the ability of data minimization algorithms to reduce the dataset size across several modalities. Figure 3 summarizes the results, comparing the amount of data that could be minimized (y-axis) given a maximum drop in accuracy (x-axis) when compared to the accuracy returned by the classifier trained over the original, not-minimized, dataset. The results indicate a strong trend: *A substantial amount of data can often be removed without sacrificing utility, suggesting that much of the collected data is superfluous in the datasets analyzed.*

Notably, the baseline methods for minimization are found effective in reducing substantial amounts of data, possibly because they leverage existing structural redundancies in the dataset. For instance, feature selection excels with the handwritten digit dataset, where outer boundary pixels are never covered by the digits, and thus can be easily removed with minimal impact on model performance. Nonetheless, the optimization algorithms consistently outperform the baselines across all datasets, with different algorithms outshining others in specific scenarios. Evolutionary algorithms, for instance, surpass others on the bank dataset, driving minimization to extreme sparsity. It highlights a strength of minimization at its core: reducing the dataset size to a mere 18.75% (and 6.25%) subset of the original bank dataset still allows the model to maintain utility with only a 2% (and 5%) drop in accuracy. Similarly, approximating the target utility excels on the handwritten digits dataset, presumably due to

its compatibility with the values of redundant features in the dataset (empty pixels are represented with 0, compatible with the zero-imputation assumption of the algorithm). An in-depth discussion of the strengths and weaknesses of these algorithms is provided in §D.

4.3 Multiplicity, emergent individualization, and privacy in data minimization

Building on the previous section’s demonstration of data minimization’s effectiveness in various settings, this section explores *the concept of multiplicity of data minimization*. We applied the evolutionary algorithm five times to the bank dataset under different randomness settings, each targeting 75% sparsity. Notably, the results, illustrated in Figure 3(c), show minimal overlap among the datasets retained from each run, with the highest overlap being only 25.51%. Clearly, there are many distinct ways to achieve data minimization while maintaining utility that meets data relevance constraints (i.e., no further data can be removed without compromising utility). Such variability underscores the flexibility of achieving data minimization through diverse algorithms. However, this also suggests that different minimized datasets may pose varying privacy risks related to the undisclosed features, as discussed in §2.1, pointing to a misalignment in privacy outcomes.

Next, we visualize the minimized dataset and highlight some intriguing characteristics. Starting with the bank dataset, minimized to 75% sparsity using the evolutionary algorithm, Figure 3(d)(Left) shows that individualization trends naturally emerge from the optimization process, where no single feature is consistently favoured over others. Similarly, the minimized images from the handwritten digits dataset, processed by the evolutionary algorithm at 50% sparsity, are shown in Figure 3(d)(Right). The trends here are more interpretable; for example, the central vertical line is preserved in the image of the digit ‘1’, while the outer curves are retained for ‘0’. Notice however that despite reducing the dataset to half its original size, a significant portion of the images can still be reconstructed using statistics learned from public data. This provides a strong indication of privacy risks and suggests that, as we will show next, a minimized dataset does not equate to enhanced privacy.

5 Data minimization and privacy

In the previous section, we showed that various algorithms can effectively minimize data while focusing on utility. However, as discussed in §2.1, there is an expectation of privacy associated with minimization, currently unexplored in literature [3, 21, 23]. Next, we conduct an important empirical assessment of minimization algorithms on real-world privacy attacks and their alignment with the minimization objective.

Defining inference attacks. During a reconstruction attack, the attacker aims to recover missing information from the minimized dataset. To achieve this, we use the data imputation method described in §4.1. In contrast, a re-identification attack involves the attacker attempting to identify an individual using only partial information. Specifically, the attacker aims to find the best match in the minimized dataset $\mathbf{X}'$ for a data point from the auxiliary dataset (in our setup, $\mathbf{X}^A \equiv \mathbf{X}$). For this purpose, we use a modified Euclidean distance that disregards missing values and adjusts the distance scaling accordingly, to accurately rank the best matches.

5.1 Evaluating privacy leakage

Data minimization and re-identification risk. We first focus on the re-identification risks for the handwritten digits dataset, shown in Figure 4(b). A key observation is that re-identification risk remains remarkably high, even after extensive minimization. For instance, the re-identification risk is close to 1 for most algorithms even when the dataset is reduced to approximately 20% of its original size. This behaviour is due to the dataset’s large feature space. Thus, while minimization algorithms were able to reduce dataset size significantly while maintaining utility (as shown in Figure 3), they do not achieve comparable reductions in re-identification risk. *This underscores a fundamental misalignment between the goals of data minimization and the actual privacy goals.*

While this misalignment is less pronounced in the bank dataset (Figure 4(a)), there are still regions in which reducing dataset size does not correspondingly decrease re-identification risk. Notably, the algorithm that approximates the target utility demonstrates a closer alignment with privacy goals compared to other algorithms. Considering the lower amount of minimization by approximating the target utility as observed in Figure 3, we can infer that the additional data minimized by the other two

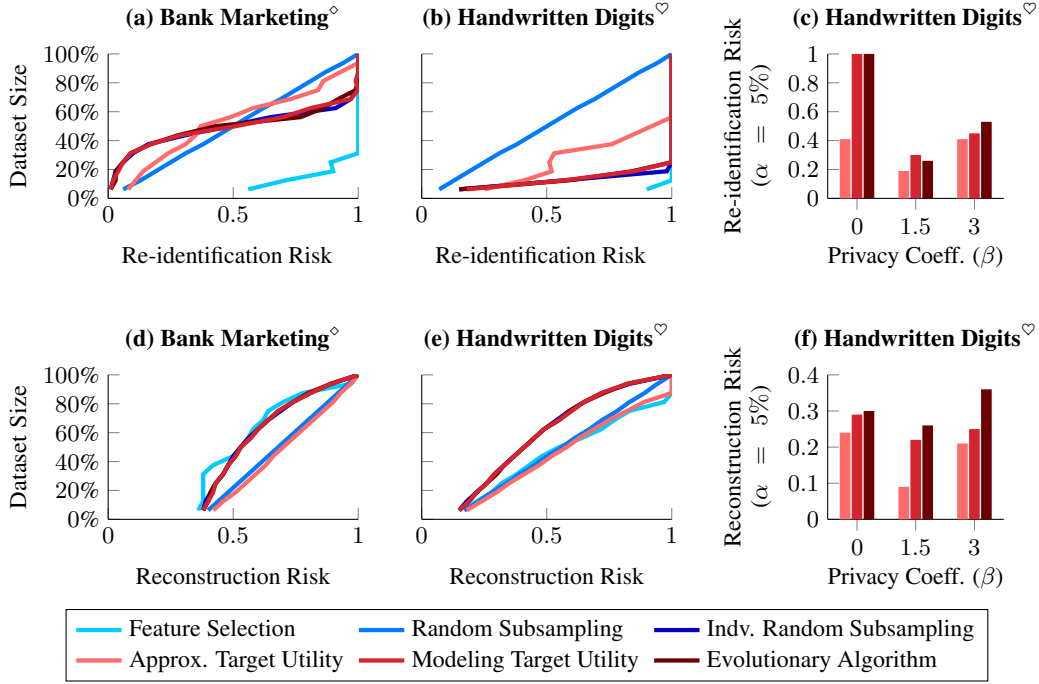


Figure 4: **(a, b, d, e)** Re-identification and reconstruction risks under changing sparsity. While the overall trends show improvement in privacy with data minimization, they are not aligned with the trends of dataset size. **(c, f)** Privacy risks under feature-level privacy scores. The lowest risk values can be seen at $\beta = 1.5$, highlighting the importance of considering privacy during minimization.

algorithms was not aligned with re-identification risk. Thus, while these algorithms would have been preferred if we only considered utility and dataset size (as we did in Figure 3), they do not provide proportionate privacy improvements, highlighting a critical concern of misalignment.

Finally, observe that our baseline algorithms display notable trends across both datasets. Random subsampling emerges as an effective strategy to reduce dataset size while being aligned with re-identification risks. Indeed, completely removing a data point also eliminates any risk of its re-identification. In contrast, feature selection, despite preserving reasonable utility post-minimization, fails to effectively mitigate re-identification risks in any dataset. This inadequacy can be attributed to the persistence of certain features in the minimized dataset that, even when isolated, can uniquely identify individuals. This phenomenon suggests a general misalignment between breadth-based minimization techniques [11, 21, 28] and the expected reduction in re-identification risks. Once again, these observations highlight the disparity between mere reductions in dataset size and actual improvements in privacy outcomes.

Data minimization and reconstruction risk. Next, we shift focus to the results of reconstruction risks, shown in Figure 4(d, e). Notably, even at the highest levels of data minimization (smallest dataset size), reconstruction risks remain significant for both domains. This aspect is linked to the ability to reconstruct features using overall dataset statistics, which, although decreasing in accuracy as minimization increases, still retain some reconstructive value. Interestingly, we find that the handwritten digits dataset is comparatively easier to protect against reconstruction, likely due to its higher feature variance relative to the bank dataset. While the same high variance in features led to worse re-identification risk, it instead helped the handwritten digits dataset achieve better alignment with reconstruction risk. Thus, not only is the dataset size not aligned with privacy, but even different privacy risks are not aligned with each other, motivating the explicit involvement of appropriate privacy constraints during minimization.

The trends observed in re-identification risk also manifest in reconstruction risks, albeit with some unique differences. Firstly, both random subsampling and individualized random subsampling have linear relationships between dataset size and privacy, highlighting a strong alignment with privacy. Secondly, the best-performing algorithm for dataset size might not always be the best choice overall.

Specifically, while the algorithm that approximates the target utility shows strong performance on the bank marketing dataset and holds a significant advantage on the handwritten digits dataset, it consistently presents higher reconstruction risks for both datasets. These trends emphasize the importance of incorporating privacy considerations when selecting the most appropriate minimization algorithm for a specific scenario.

5.2 Adapting minimization for better privacy

We now present a simple yet effective adjustment to the minimization problem to align better with privacy and demonstrate the feasibility of a more comprehensive minimization objective, managing both utility and privacy. Given the output of a minimization algorithm $B^a = [B]_{\perp \rightarrow 0}$, we define privacy scores V , such that the score matrix $C_{ij} = B_{ij}^a + \beta V_{ij}$ determines the indices that should be minimized. For simplicity, we will only define feature-level privacy scores, which are not personalized, i.e., $V_{ij} = V_{qj}$, $\forall i, q \in [n], j \in [p]$. Here, β serves as a hyperparameter to tune the emphasis on privacy. Ultimately, we minimize values with the lowest C_{ij} scores to achieve the target sparsity. The privacy scores V are normalized to $[0, 1]$ before being combined with the minimization matrix.

- **Privacy scores to reduce re-identification.** In re-identification attacks, the risk arises from disclosing *unique* features (e.g., phone number, SSN) rather than non-unique features (e.g., gender, race), as they make it easier to identify individuals. Following this rationale, we propose using the negative of the number of unique values for a feature as its privacy score. Minimizing distinctive features can increase the difficulty of re-identification.
- **Privacy scores to reduce reconstruction.** In reconstruction attacks, the risk stems from features that exhibit a high correlation with other features, exploited by adversaries to deduce missing values. Thus, to mitigate reconstruction risks, we propose using the negative of the average correlation of each feature with every other feature as a measure of its independence and thus its privacy score.

Evaluating privacy score modifications. For a given level of sparsity, it can be expected that incorporating privacy scores will decrease data breach risks, but potentially reduce accuracy. The critical question is whether we can attain a more favourable trade-off between privacy and utility, regardless of the sparsity. We present the results in Figure 4(c, f) for the handwritten digits dataset, varying the hyperparameter β . Results for other datasets are present in §F.

At $\beta = 1.5$, integrating privacy scores enhances the privacy-utility trade-off, cutting the re-identification risk by more than half while maintaining the same accuracy drop. The impact on reconstruction risk is less dramatic, but improvements are evident. Increasing the focus on privacy to $\beta = 3$ results in a less optimal trade-off, yet still better than at $\beta = 0$. This notable effect of a basic feature-level privacy score underscores the necessity of directly addressing privacy in the minimization process, instead of relying on its incidental emergence.

6 Discussion and conclusion

In proposing a formal framework for data minimization in ML, this paper reveals a disconnect between legal mandates and their practical implementation. While data protection regulations aim to limit data collection with an expectation of privacy, current objectives of minimization fall short of providing robust privacy safeguards. Notice, however, that this is not to say that minimization is incompatible with privacy; instead, we emphasize the need for approaches that incorporate privacy into their objectives (as done in §5.2), rather than treating them as an afterthought.

Addressing this concern also brings forth a variety of optimization challenges. The optimization problem is complicated further under alternative interpretations of data collection, such as gathering a range of data rather than exact values [11, 28]. These challenges are particularly significant in large-scale applications where both time and accuracy are critical factors. Therefore, future work should focus on the development of efficient minimization algorithms able to make good trade-offs between utility and privacy. The ethical and fairness considerations of data minimization also add to its complexity. By design, data minimization is likely to remove data that resembles the majority [3], leaving the minority more vulnerable to privacy risks. Thus, developing fairness-aware mechanisms for minimization is an important avenue for future research.

In summary, our study marks a crucial step in aligning the legal requirements with practical, technical solutions for interpreting and operationalizing data minimization in ML. It sets the stage for future

work aimed at developing comprehensive, efficient, and ethically sound methodologies for minimization. It hopes to guide the creation of ML systems that are both respectful of individual privacy and compliant with legal norms, thereby emphasizing the need for multidisciplinary approaches.

Acknowledgments and Disclosure of Funding

This research is partially supported by NSF grants 2133169, 2232054 and NSF CAREER Award 2143706. Fioretto is also supported by an Amazon Research Award and a Google Research Scholar Award. Its views and conclusions are those of the authors only.

References

- [1] M. Abadi, A. Chu, I. Goodfellow, H. B. McMahan, I. Mironov, K. Talwar, and L. Zhang. Deep learning with differential privacy. In *Proceedings of the 2016 ACM SIGSAC conference on computer and communications security*, pages 308–318, 2016.
- [2] A. J. Biega and M. Finck. Reviving Purpose Limitation and Data Minimisation in Data-Driven Systems. *Technology and Regulation*, pages 44–61 Pages, Dec. 2021.
- [3] A. J. Biega, P. Potash, H. Daumé, F. Diaz, and M. Finck. 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*, pages 399–408, 2020.
- [4] C. M. Bishop and N. M. Nasrabadi. *Pattern recognition and machine learning*, volume 4. Springer, 2006.
- [5] A. L. Blum and P. Langley. Selection of relevant features and examples in machine learning. *Artificial intelligence*, 97(1-2):245–271, 1997.
- [6] Z. Chen, F. Sun, Y. Tang, H. Chen, J. Gao, and B. Ding. Studying the impact of data disclosure mechanism in recommender systems via simulation. *ACM Transactions on Information Systems*, 41(3):1–26, 2023.
- [7] P. Cortez, A. Cerdeira, F. Almeida, T. Matos, and J. Reis. Modeling wine preferences by data mining from physicochemical properties. *Decision support systems*, 47(4):547–553, 2009.
- [8] F. Ding, M. Hardt, J. Miller, and L. Schmidt. Retiring adult: New datasets for fair machine learning. *Advances in neural information processing systems*, 34:6478–6490, 2021.
- [9] C. Dwork. Differential privacy. In *International colloquium on automata, languages, and programming*, pages 1–12. Springer, 2006.
- [10] S. Garfinkel, J. M. Abowd, and C. Martindale. Understanding database reconstruction attacks on public data. *Communications of the ACM*, 62(3):46–53, 2019.
- [11] A. Goldsteen, G. Ezov, R. Shmelkin, M. Moffie, and A. Farkash. Data minimization for gdpr compliance in machine learning models. *AI and Ethics*, pages 1–15, 2021.
- [12] J. Henriksen-Bulmer and S. Jeary. Re-identification attacks—a systematic literature review. *International Journal of Information Management*, 36(6):1184–1192, 2016.
- [13] P. D. Hoff. Extending the rank likelihood for semiparametric copula estimation. *Annals of Applied Statistics*, 1(1):265–283, 2007.
- [14] G. Hongli, L. Juntao, and G. Hong. A survey of bilevel programming model and algorithm. In *2011 Fourth International Symposium on Computational Intelligence and Design*, volume 2, pages 199–203, 2011. doi: 10.1109/ISCID.2011.151.
- [15] H. James, C. Nagpal, K. A. Heller, and B. Ustun. Participatory personalization in classification. In *ICML 3rd Workshop on Interpretable Machine Learning in Healthcare (IMLH)*, 2023.
- [16] T. Leemann, M. Pawelczyk, C. T. Eberle, and G. Kasneci. I prefer not to say: Operationalizing fair and user-guided data minimization. *arXiv preprint arXiv:2210.13954*, 2022.
- [17] S. Moro, P. Cortez, and P. Rita. A data-driven approach to predict the success of bank telemarketing. *Decision Support Systems*, 62:22–31, 2014.
- [18] M. Nasr, R. Shokri, and A. Houmansadr. Comprehensive privacy analysis of deep learning: Passive and active white-box inference attacks against centralized and federated learning. In *2019 IEEE symposium on security and privacy (SP)*, pages 739–753. IEEE, 2019.

- [19] M. Paul, S. Ganguli, and G. K. Dziugaite. Deep learning on a data diet: Finding important examples early in training. *Advances in Neural Information Processing Systems*, 34:20596–20607, 2021.
- [20] P. Quinn and G. Malgieri. The difficulty of defining sensitive data—the concept of sensitive data in the eu data protection framework. *German Law Journal*, 22(8):1583–1612, 2021.
- [21] B. Rastegarpanah, K. Gummadi, and M. Crovella. Auditing black-box prediction models for data minimization compliance. *Advances in Neural Information Processing Systems*, 34: 20621–20632, 2021.
- [22] M. Rigaki and S. Garcia. A survey of privacy attacks in machine learning. *ACM Computing Surveys*, 56(4):1–34, 2023.
- [23] D. Shanmugam, F. Diaz, S. Shabanian, M. Finck, and A. Biega. Learning to limit data collection via scaling laws: A computational interpretation for the legal principle of data minimization. In *2022 ACM Conference on Fairness, Accountability, and Transparency*, pages 839–849, 2022.
- [24] R. Shokri, M. Stronati, C. Song, and V. Shmatikov. Membership inference attacks against machine learning models. In *2017 IEEE symposium on security and privacy (SP)*, pages 3–18. IEEE, 2017.
- [25] A. Sinha, P. Malo, A. Frantsev, and K. Deb. Finding optimal strategies in a multi-period multi-leader–follower stackelberg game using an evolutionary algorithm. *Computers & Operations Research*, 41:374–385, 2014.
- [26] B. Sorscher, R. Geirhos, S. Shekhar, S. Ganguli, and A. Morcos. Beyond neural scaling laws: beating power law scaling via data pruning. *Advances in Neural Information Processing Systems*, 35:19523–19536, 2022.
- [27] N. Srebro. How good is a kernel when used as a similarity measure? In *Learning Theory: 20th Annual Conference on Learning Theory, COLT 2007, San Diego, CA, USA; June 13-15, 2007. Proceedings 20*, pages 323–335. Springer, 2007.
- [28] R. Staab, N. Jovanović, M. Balunović, and M. Vechev. From principle to practice: Vertical data minimization for machine learning. In *2024 IEEE Symposium on Security and Privacy (SP)*, pages 89–89. IEEE Computer Society, 2024.
- [29] L. Thomas, I. Gondal, T. Oseni, and S. S. Firmin. A framework for data privacy and security accountability in data breach communications. *Computers & Security*, 116:102657, 2022.
- [30] R. Tibshirani. Modeling basics: Assessment, selection, and complexity. 2015. URL www.stat.cmu.edu/~ryantibs/statml.
- [31] C. Tran and F. Fioretto. Personalized privacy auditing and optimization at test time. *arXiv preprint arXiv:2302.00077*, 2023.
- [32] G. G. Wang and S. Shan. Review of metamodeling techniques in support of engineering design optimization. In *International Design Engineering Technical Conferences and Computers and Information in Engineering Conference*, volume 4255, pages 415–426, 2006.
- [33] L. Xu, A. Krzyzak, and C. Y. Suen. Methods of combining multiple classifiers and their applications to handwriting recognition. *IEEE transactions on systems, man, and cybernetics*, 22(3):418–435, 1992.
- [34] S. Yeom, I. Giacomelli, M. Fredrikson, and S. Jha. Privacy risk in machine learning: Analyzing the connection to overfitting. In *2018 IEEE 31st computer security foundations symposium (CSF)*, pages 268–282. IEEE, 2018.
- [35] S. Zarifzadeh, P. Liu, and R. Shokri. Low-cost high-power membership inference attacks. *arXiv preprint arXiv:2312.03262*, 2023.

A Legal language

Table 2 summarizes relevant excerpts from six data protection regulations from around the world. Along with [purpose limitation](#) and [data relevance](#), we also highlight excerpts that show the expectation of [data privacy](#) through minimization.

Why interpret data minimization as dataset size? As discussed in §2, one of the core principles of data minimization is data relevance (the other one is purpose limitation). Data relevance suggests that data collected should be limited to what is necessary, should be non-excessive, etc. This is interpreted as limiting the quantity of the data, because if the dataset size is minimized under a certain utility constraint, then no more data can be removed without hurting utility, hence the data is indeed limited to what is necessary. Our interpretation of data relevance as limiting data collection to dataset size, and purpose limitation as an explicitly defined utility function, are built on existing work in the literature [2, 3, 20, 21, 23, 28].

Biega and Finck [2] interpret data relevance as reducing either the data quantity (amount of data) or data quality (coarseness of the data), while purpose limitation is a legal obligation made of three distinct requirements, that the purpose be specified, that it be explicit, and that it be legitimate. Quinn and Malgieri [20] discuss in detail the definition of sensitive or personal data, providing two broad categories, contextual, where sensitive data is information that can be perceived as sensitive in context, or purposeful, where data is sensitive if it is used to perform sensitive actions on individuals. Notably, both of these papers focus on removing data without consideration for whether the data can be recovered or not, and thus, neither of these works directly signal any discussion of the privacy risks for an individual under data breach. This interpretation has also driven most attempts at operationalizing the data minimization principle [3, 21, 23, 28], focusing heavily on reducing data size.

Why expect privacy from data minimization? Data minimization has been hailed by many in the public discourse¹²³⁴⁵ as a principle to protect privacy in case of a potential data breach. For example, the EU AI Act introduces data minimization in Recital 69⁶ as *‘The right to privacy and to protection of personal data must be guaranteed throughout the entire lifecycle of the AI system. In this regard, the principles of data minimisation and data protection by design and by default, as set out in Union data protection law, are applicable when personal data are processed’*. Sentiments such as, *‘Data minimization is a direct way to limit privacy leakage. Intuitively, the less data there is to collect, store, and share by smart devices, the easier it is for users to protect their personal information’*⁷, have also become commonplace.

Moreover, as we highlight in Table 2, data minimization in various data protection regulations is defined on personal data, emphasizing that it is the personal data that needs to be minimized. Personal data is defined as any information that directly *identifies a particular individual*, or can be indirectly *combined with other information to identify a particular individual* (PIPA, South Korea). In other words, the term personal data is often used in various data protection regulations to mark information that can hurt an individual’s privacy in some way. Hence, the reduction of personal data can be tied to an expectation of privacy from data minimization. Thus, both the general discourse around minimization, as well as the language in various data protection regulations, create an expectation of privacy from minimization.

But this is in contrast to the interpretation of data minimization as reducing the dataset set and the work in the literature on data minimization that does not inspect the true privacy risks of these models. This creates a gap, where individuals or organizations using these operationalization attempts of data minimization, might expect improved privacy, despite using a framework that is limited to only dataset size reduction.

¹https://www.edps.europa.eu/data-protection/data-protection/glossary/d_en

²<https://www.kiteworks.com/risk-compliance-glossary/data-minimization/>

³<https://therecord.media/lawmakers-set-sights-on-data-minimization-with-new-bills/>

⁴<https://www.skadden.com/insights/publications/2024/04/cppas-first-enforcement-advisory>

⁵<https://www.k2view.com/blog/data-minimization/>

⁶<https://www.euaiact.com/recital/69>

⁷<https://www.sciencedirect.com/topics/computer-science/data-minimization>

General Data Protection Regulation (GDPR), Europe	gdpr-info.eu/
Article 4(1): “personal data” means any information relating to an identified or identifiable natural person (“data subject”) [...]; Article 5(1)(b): Personal data shall be collected for specified, explicit and legitimate purposes and not further processed in a manner that is incompatible with those purposes; Article 5(1)(c): Personal data shall be adequate, relevant and limited to what is necessary in relation to the purposes for which they are processed.	
California Privacy Rights Act (CPRA), USA	cpa.ca.gov/
Section 1798.100 (a)(1) & (a)(2): [...] A business shall not collect additional categories of (sensitive) personal information or use (sensitive) personal information collected for additional purposes that are incompatible with the disclosed purpose for which the (sensitive) personal information was collected without providing the consumer with notice consistent with this section. Section 1798.140 (v)(1): “Personal information” means information that identifies, relates to, describes, is reasonably capable of being associated with, or could reasonably be linked, directly or indirectly, with a particular consumer or household.	
General Personal Data Protection Law (LGPD), Brazil	lgpd-brazil.info/
Article 5(I): personal data: information regarding an identified or identifiable natural person; Article 6: Activities of processing of personal data shall be subject to the following principles, I: processing done for legitimate, specific and explicit purposes of which the data subject is informed, with no possibility of subsequent processing that is incompatible with these purposes ; II: compatibility of the processing with the purposes communicated to the data subject , in accordance with the context of the processing; III: limitation of the processing to the minimum necessary to achieve its purposes , covering data that are relevant, proportional and non-excessive in relation to purposes of the data processing [...];	
Protection of Personal Information Act (POPIA), South Africa	popia.co.za/
Section 1: “personal information” means information relating to an identifiable, living, natural person , [...] Section 10: Personal information may only be processed if, given the purpose for which it is processed, it is adequate, relevant and not excessive . Section 13(1): Personal information must be collected for a specific, explicitly defined and lawful purpose [...]	
Consumer Data Rights (CDR), Australia	www.legislation.gov.au/Details/F2023C00735
Rule 1.8(a): An accredited person complies with the data minimisation principle if when making a consumer data request on behalf of a CDR consumer, it does not seek to collect: (i) more CDR data than is reasonably needed ; [...] Rule 4.11(1)(a): When asking a CDR consumer to give a consent, an accredited person must allow the CDR consumer to choose the types of CDR data to which the consent will apply by enabling the CDR consumer to actively select or otherwise clearly indicate: (i).; and (ii) in the case of a use consent—the specific uses of collected data to which they are consenting ; [...]	
Personal Information Protection Act (PIPA), South Korea	www.pipc.go.kr/eng/index.do
Article 2(1): The term “personal information” means any of the following information relating to a living individual: (a) Information that identifies a particular individual [...]; (b) Information which, even if it by itself does not identify a particular individual, may be easily combined with other information to identify a particular individual [...]; Article 3(1): The personal information controller shall specify explicitly the purposes for which personal information is processed ; and shall collect personal information lawfully and fairly to the minimum extent necessary for such purposes .	

Table 2: Excerpts from various data protection regulations from across the globe on the principle of data minimization, highlighting language on **purpose limitation** , **data relevance** , and references to the expectations of **data privacy** .

B Related work

Data Minimization. Existing research on minimization in the context of data protection regulations can be broadly divided into *breadth-based* and *depth-based* techniques. Breadth-based strategies aim to minimize data by limiting the number of features [21] or introducing feature generalization [11, 28]. On the other hand, depth-based approaches focus on reducing the number of unique data points by using methods like data pruning [19, 23, 26]. While there are some discussions on individualized minimization for recommender systems [3, 6], they are limited in their ability to generalize to other settings in ML.

On a separate note, most studies in data minimization aim to simply reduce the raw size of the datasets [3, 6, 21, 23] and don't give any attention to privacy concerns [16]. Although some works do go beyond dataset size and discuss other aspects of information leakage [11], they still lack connections with real-world privacy risks. The work closest to ours is a concurrent work by Staab et al. [28], which also introduces real-world privacy attacks to quantify privacy leakage after minimization. However, unlike our approach, Staab et al. [28] concentrates on breadth-based methods, thus missing the individualized nature of minimization.

Some studies have also formalized data minimization during inference, emphasizing the personalized nature of the process and delving into its privacy implications [15, 31]. However, data minimization during inference is distinctly different from data minimization during training, which is the primary focus of our paper.

Privacy Auditing with Inference Attacks. Privacy auditing often involves the application of inference attacks to assess real-world privacy leakage. These attacks serve as powerful tools to expose potential vulnerabilities and by simulating realistic scenarios [12, 24, 34], auditors can evaluate the effectiveness of privacy protection mechanisms and identify areas where sensitive information may be revealed.

C Theoretical analysis of baseline techniques

We provide the theoretical properties of the model learned on minimized data for feature selection and individualized random subsampling in this section.

C.1 Feature selection

As introduced in the main text, the feature selection framework sorts the importance of features based on their importance in learning task and then remove the least important features. We denote $S \in [p]$ to be the subset of the weakest features that will be not used. The following Theorem 1 provides the Bayes Mean Squared Error (MSE) when using all features $[p]$ and using a subset of features $[p] \setminus S$.

Theorem 1. *Suppose all input features and labels are jointly Gaussian, i.e., $[\mathbf{x}, y] \sim \mathcal{N}(\mu, \Sigma)$, where $\Sigma = \begin{bmatrix} \Sigma_{\mathbf{x}, \mathbf{x}} & \Sigma_{\mathbf{x}, y} \\ \Sigma_{y, \mathbf{x}} & \Sigma_{y, y} \end{bmatrix}$. Furthermore, we assume that all input features are mutually independent, i.e., the covariance matrix $\Sigma_{\mathbf{x}, \mathbf{x}} = \text{diag}([\sigma_i^2]_{i=1}^p)$ is a diagonal matrix and $\sigma_i^2 = \text{Var}[\mathbf{x}^i]$ the variance of i^{th} feature. Then the Bayes MSE when using all input features $[p]$ and using a subset of input features in $[p] \setminus S$ in turn are: $\text{Var}[y] - \sum_{i=1}^p \frac{(\text{Cov}(y, \mathbf{x}^i))^2}{\sigma_i^2}$ and $\text{Var}[y] - \sum_{i \in [p] \setminus S} \frac{(\text{Cov}(y, \mathbf{x}^i))^2}{\sigma_i^2}$*

Theorem 1 suggests data minimization procedure based on feature selection introduces an additional MSE of $\sum_{i=1}^p \frac{(\text{Cov}(y, \mathbf{x}^i))^2}{\sigma_i^2} - \sum_{i \in [p] \setminus S} \frac{(\text{Cov}(y, \mathbf{x}^i))^2}{\sigma_i^2} = \sum_{i \in S} \frac{(\text{Cov}(y, \mathbf{x}^i))^2}{\sigma_i^2}$. As long as all features in the removed feature set S have small correlation with the label, i.e. $\text{Cov}(y, \mathbf{x}^i) \approx 0$, such additional MSE is negligible.

Proof. Our proof relies on the properties of multivariate Gaussian variables. In particular, when feature and label are jointly Gaussian, $[\mathbf{x}, y] \sim \mathcal{N}(\mu, \Sigma)$ then for any subset of features $A \in [p]$ we can derive the following conditional density of label y given partial input features $\mathbf{x}_A$ [4]:

$$P(y|\mathbf{x}_A) = \mathcal{N}(\mu_y + \Sigma_{\mathbf{x}_A, y} \Sigma_{y, y}^{-1} \Sigma_{y, \mathbf{x}} - \Sigma_{\mathbf{x}_A, y} \Sigma_{\mathbf{x}, \mathbf{x}}^{-1} \Sigma_{\mathbf{x}, y}, \Sigma_{y, y} - \Sigma_{\mathbf{x}_A, y} \Sigma_{\mathbf{x}, \mathbf{x}}^{-1} \Sigma_{\mathbf{x}, y}).$$

In the above equation, $\Sigma_{y,y}$ is the variance of the label y , i.e., $\Sigma_{y,y} = \text{Var}(y)$, while $\Sigma_{\mathbf{x}_A,y}$ or $\Sigma_{y,\mathbf{x}_A}$ is the covariance between y and a subset of features $\mathbf{x}_A$. The Bayes MSE for Gaussian distribution is just the conditional variance [30]. Hence the Bayes MSE when using all features, i.e $A = [p]$ is: $\Sigma_{y,y} - \Sigma_{y,\mathbf{x}} \Sigma_{\mathbf{x},\mathbf{x}}^{-1} \Sigma_{\mathbf{x},y}$. By the assumption that the input features are mutually independent $\Sigma_{\mathbf{x},\mathbf{x}} = \text{diag}([\sigma_i^2]_{i=1}^p)$, the Bayes MSE using all features can be further reduced as: $\text{Var}[y] - \Sigma_{y,\mathbf{x}} \text{diag}([\frac{1}{\sigma_i^2}]_{i=1}^p) \Sigma_{\mathbf{x},y} = \text{Var}[y] - \sum_{i \in [p]} \frac{(\text{Cov}(y, \mathbf{x}_i))^2}{\sigma_i^2}$.

Similarly, we can derive the Bayes MSE when using a subset of features $d \setminus S$ as: $\text{Var}[y] - \sum_{i \in [d] \setminus S} \frac{(\text{Cov}(y, \mathbf{x}_i))^2}{\sigma_i^2}$. $\square$

C.2 Personalized random subsampling

As introduced in the main text, the personalized random subsampling method works by randomly setting the entries of the minimization $\mathbf{B}_{ij} = \perp$ with a probability k_p , which controls the sparsity of the minimized dataset (Here, $k_p = \frac{k}{np}$). This can also be rewritten as randomly selecting a subset $S \in \{(i, j) | i \in [n], j \in [p]\}$ of a given size $|S| = np - k$, and remove the entries in S . To understand the theoretical behaviour of this method, we first consider the optimal model parameter $\theta^*(\mathbf{X}) = \arg \min_{\theta} J(\theta; \mathbf{X}, \mathbf{Y})$ as a function of the data $\mathbf{X}$. Using this notation, the model learnt on minimized data can be represented as $\hat{\theta} = \theta^*(\mathbf{X}')$ while the model parameter learned on original data $\hat{\theta} = \theta^*(\mathbf{X})$. We then have the following Lemma 1 that derives the gradient of the optimal model parameter $\theta^*(\mathbf{X})$ w.r.t data $\mathbf{X}$.

Lemma 1 (Sensitivity of model parameter w.r.t input $\mathbf{X}$). *Assume the loss function $J(\theta; \mathbf{X}, \mathbf{Y})$ is differentiable w.r.t θ and $\mathbf{X}$, suppose $\theta^*(\mathbf{X}) = \arg \min_{\theta} J(\theta; \mathbf{X}, \mathbf{Y})$ then the following holds:*

$$\frac{\partial \theta^*(\mathbf{X})}{\partial \mathbf{X}} = -H^{-1}G, \quad (5)$$

where $H = \frac{\partial^2 J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y})}{\partial \theta^2}$ is the Hessian matrix of the loss w.r.t model parameter, while $G = \frac{\partial^2 J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y})}{\partial \theta \partial \mathbf{X}}$ is the second derivative of the loss w.r.t model parameter and input data.

Proof. Since the model parameter $\theta^*(\mathbf{X}) = \arg \min_{\theta} J(\theta; \mathbf{X}, \mathbf{Y})$ hence the gradient of the loss $J(\cdot)$ w.r.t model parameter vanishes at $\theta^*(\mathbf{X})$. In other words

$$\frac{\partial J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y})}{\partial \theta} = \mathbf{0}^T.$$

Take the derivative w.r.t $\mathbf{X}$ on both sides of the above equation, it follows:

$$\frac{\partial}{\partial \mathbf{X}} \frac{\partial J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y})}{\partial \theta} = \mathbf{0}^T. \quad (6)$$

The L.H.S can be rewritten as

$$\begin{aligned} \frac{\partial}{\partial \theta} \frac{\partial J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y})}{\partial \mathbf{X}} &= \frac{\partial}{\partial \theta} \left[\frac{\partial J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y})}{\partial \theta} \frac{\partial \theta^*(\mathbf{X})}{\partial \mathbf{X}} \right. \\ &\quad \left. + \frac{\partial J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y})}{\partial \mathbf{X}} \frac{\partial \mathbf{X}}{\partial \mathbf{X}} \right] \\ &= \frac{\partial^2 J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y})}{\partial \theta^2} \frac{\partial \theta^*(\mathbf{X})}{\partial \mathbf{X}} \\ &\quad + \frac{\partial^2 J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y})}{\partial \theta \partial \mathbf{X}}, \end{aligned}$$

where the first equation is due to the chain rule. If we put $H = \frac{\partial^2 J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y})}{\partial \theta^2}$ (the Hessian matrix) and $G = \frac{\partial^2 J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y})}{\partial \theta \partial \mathbf{X}}$, then the L.H.S of Equation 6 can be rewritten as:

$$H \frac{\partial \theta^*(\mathbf{X})}{\partial \mathbf{X}} + G = \mathbf{0}^T,$$

which implies $\frac{\partial \theta^*(\mathbf{X})}{\partial \mathbf{X}} = -H^{-1}G$.

□

Lemma 1 tells us how much the optimal model parameter changes when the data input changes. Based on this Lemma 1 we can prove the following Lemma 2 regarding the target utility.

Lemma 2 (Sensitivity of target utility w.r.t input $\mathbf{X}$). *Given the same settings and conditions as in Lemma 1, then the following holds:*

$$\frac{\partial J(\theta^*(\mathbf{X}), \mathbf{X}, \mathbf{Y})}{\partial \mathbf{X}} = -LH^{-1}G, \quad (7)$$

where H and G were provided in Lemma 1, and $L = \frac{\partial J(\theta^*(\mathbf{X}), \mathbf{X}, \mathbf{Y})}{\partial \theta}$ is the gradient of loss w.r.t model parameter.

Proof. The proof is based directly on the chain rule:

$$\frac{\partial J(\theta^*(\mathbf{X}), \mathbf{X}, \mathbf{Y})}{\partial \mathbf{X}} = \frac{\partial J(\theta^*(\mathbf{X}), \mathbf{X}, \mathbf{Y})}{\partial \theta^*(b\mathbf{X})} \frac{\partial \theta^*(b\mathbf{X})}{\partial \mathbf{X}} = -LH^{-1}G,$$

where the last equation is by Lemma 1. □

Based on Lemma 2 we have the following Theorem 2 that derives the target utility $J(\hat{\theta}; \mathbf{X}, \mathbf{Y})$ bound based on the original utility $J(\theta; \mathbf{X}, \mathbf{Y})$ as follows:

Theorem 2. *Consider the personalized random subsampling framework, in which all features in a random set $S \in \{(i, j) | i \in [n], j \in [p]\}$ is removed to form the minimized dataset $\mathbf{X}'$, then the following holds:*

$$J(\hat{\theta}; \mathbf{X}, \mathbf{Y}) \leq J(\theta; \mathbf{X}, \mathbf{Y}) + \sqrt{2|S|} \|\mathbf{X}\|_\infty \|LH^{-1}G\|_2 + \mathcal{O}(|S| \|\mathbf{X}\|_\infty^2). \quad (8)$$

Proof. The proof relies on the first-order Taylor approximation. First, we consider both optimal model parameters $\hat{\theta}, \dot{\theta}$ as a function of the input data, i.e., $\hat{\theta} = \theta^*(\mathbf{X}')$ and $\dot{\theta} = \theta^*(\mathbf{X})$. Based on the first-order Taylor approximation around $\mathbf{X}$ it follows that:

$$\begin{aligned} J(\theta^*(\mathbf{X}'); \mathbf{X}, \mathbf{Y}) &\approx J(\theta^*(\mathbf{X}); \mathbf{X}, \mathbf{Y}) \\ &+ (\mathbf{X}' - \mathbf{X})^T \frac{\partial J(\theta^*(\mathbf{X}), \mathbf{X}, \mathbf{Y})}{\partial \mathbf{X}} \\ &+ \mathcal{O}(\|\mathbf{X}' - \mathbf{X}\|_2^2) \end{aligned} \quad (9)$$

By Lemma 2, $\frac{\partial J(\theta^*(\mathbf{X}), \mathbf{X}, \mathbf{Y})}{\partial \mathbf{X}} = -LH^{-1}G$. Furthermore

$$\begin{aligned} \|\mathbf{X}' - \mathbf{X}\|_2^2 &= \sum_{i=1}^n \sum_{j=1}^p (\mathbf{X}'_{i,j} - \mathbf{X}_{i,j})^2 = \sum_{(i,j) \in S} (\mathbf{X}'_{i,j} - \mathbf{X}_{i,j})^2 \\ &\leq \sum_{(i,j) \in S} 2\|\mathbf{X}\|_\infty^2 = 2|S| \|\mathbf{X}\|_\infty^2, \end{aligned} \quad (10)$$

where the second equation is due to the fact we only remove features in S while the other entries are kept the same. The inequality is due to the fact that $|\mathbf{X}'_{i,j} - \mathbf{X}_{i,j}| \leq 2 \max_{i,j} |\mathbf{X}_{i,j}| = 2\|\mathbf{X}\|_\infty$, since the imputed data is in the range $\mathbf{X}'_{i,j} \in [\min_{i,j} \mathbf{X}_{i,j}, \max_{i,j} \mathbf{X}_{i,j}]$.

By Equation 9 and Cauchy-Schwarz inequality for vectors it follows that:

$$\begin{aligned} (\mathbf{X}' - \mathbf{X})^T \frac{\partial J(\theta^*(\mathbf{X}), \mathbf{X}, \mathbf{Y})}{\partial \mathbf{X}} &\leq \|\mathbf{X}' - \mathbf{X}\|_2 \|LH^{-1}G\|_2 \\ &\leq \sqrt{2|S|} \|\mathbf{X}\|_\infty \|LH^{-1}G\|_2. \end{aligned} \quad (11)$$

Applying the results from Equation 10 and Equation 11 to Equation 9 we verify the correctness of the statement. □

D Data minimization algorithms

We provide additional details on various algorithms used in the paper.

D.1 Baseline techniques

Feature Selection [5]. Feature selection is a breadth-based minimization strategy that retains only the most important features in the data. The algorithm works by first sorting the features of the dataset in order of their importance, using a pre-established criterion, to identify a subset S of the least important features. The minimized dataset $\mathbf{X}'$ is formed from $\mathbf{X}$ by removing all features in S . In our paper, the importance criterion is the absolute correlation between each feature and the output label, and the algorithm sets:

$$B_{ij} = \perp \quad \forall i \in [n], j \in S \quad (12)$$

where $|S|$ is a parameter controlling the minimization sparsity.

Random Subsampling. Random subsampling is a depth-based minimization strategy that randomly chooses a subset of data points from the original dataset. In the context of data minimization, it sets the minimization matrix B as:

$$B_{ij} = \perp \quad \forall j \in [p], \text{ with probability } k_p, \quad (13)$$

where $0 \leq k_p \leq 1$ is a probability value chosen so that nk_p rows of the dataset D are minimized, in expectation.

Individualized Random Subsampling. Individualized random subsampling is an extension of random subsampling, but it removes individual entries (feature, sample) rather than complete rows, i.e., it performs individualized minimization. The minimization matrix B is defined as:

$$B_{ij} = \perp \quad \text{with probability } k_p, \quad (14)$$

where $0 \leq k_p \leq 1$ is chosen so that npk_p elements of the dataset D are minimized, in expectation.

D.2 Data minimization algorithms

Approximating the Lower Level Program [14]. Solving the bi-level optimization discussed is challenging due to the nested structure of the problem which requires solving a non-convex lower-level optimization within a non-convex upper-level optimization. When the lower-level problem has a unique solution that can be explicitly expressed in a closed form, then the overall bilevel program can be rewritten as a single-level program which is much simpler to solve. The underlying idea behind the proposed framework lies in approximating the target utility via the original utility by the first Taylor approximation (see again Equation 9). Assuming that the second order component associated with $\|\mathbf{X}' - \mathbf{X}\|_2^2$ is negligible, the difference in model's utility is

$$J(\hat{\theta}; \mathbf{X}, \mathbf{Y}) - J(\hat{\theta}^*; \mathbf{X}, \mathbf{Y}) \approx -(\mathbf{X}' - \mathbf{X})^T L H^{-1} G.$$

where L, H, G are the gradients w.r.t. the model's parameters, the Hessian w.r.t. the model's parameters, and the second-order derivative w.r.t. the model's parameters and the dataset, respectively, of the original utility on the complete dataset $J(\hat{\theta}; \mathbf{X}, \mathbf{Y})$. This simplifies the original optimization as:

$$\underset{B \in \{\perp, 1\}^{n \times p}}{\text{Minimize}} \quad \|B\|_1 \quad \text{s.t. : } (\mathbf{X} - \mathbf{X}')^T L H^{-1} G \leq \alpha. \quad (15)$$

Suppose, we perform a simple zero imputation, i.e., setting $\perp = 0$, then $\mathbf{X}' = \mathbf{X} \odot B$. The above problem now is a binary integer linear programming and the dual problem can be rewritten as below:

$$\underset{B \in \{\perp, 1\}^{n \times p}}{\text{Minimize}} \quad (\mathbf{X} \odot B - \mathbf{X})^T L H^{-1} G \quad (16a)$$

$$\text{s.t. : } \|B\|_1 \leq k. \quad (16b)$$

Recall again that it is often convenient to view the optimization expressed in equation 2 as a program that minimizes the loss $J(\hat{\theta}; \mathbf{X}, \mathbf{Y})$ under sparsity constraint over the minimization matrix B . It turns

out that for this binary linear programming, we can obtain a closed-form solution by setting $B_{ij} = 1$ for k largest entries of $(\mathbf{X}_{ij} - \mathbf{X}'_{ij}) \odot (LH^{-1}G)_{ij}$, and $B_{ij} = \perp$ otherwise. Note that this proposed method is applicable when the imputed data is given (e.g., zero imputation) before running the method. Complicated data imputation like mean/mode/median imputation by the non-missing entries of the same column in the minimized data will result in a non-linear objective in the dual problem. This poses a difficulty since there is no efficient method to solve the binary integer non-linear programming in general.

Modeling the Target Utility [32]. Instead of relying on the assumptions of low sparsity to approximate the target utility as above, we can take a more general approach by directly modelling and learning the mapping between the target utility and the minimized dataset. In other words, we want to learn a parametrized function $m_\omega(\mathbf{B}) \approx J(\hat{\theta}; \mathbf{X}, \mathbf{Y})$, where ω is a vector of parameters, to estimate the target utility without solving the lower-level optimization for $\hat{\theta}$.

To learn a tractable mapping $m_\omega(\mathbf{B})$, we restrict ourselves to linear functions of $\mathbf{B}$ and assume that each index of the dataset has an independent influence on the target utility. Therefore, if we quantify the influence of each index on the target utility as $\mathcal{I}_{ij}$, the mapping $m_\omega(\mathbf{B})$ becomes:

$$J(\hat{\theta}; \mathbf{X}, \mathbf{Y}) \approx m_\omega(\mathbf{B}) = \frac{1}{np} \sum_{i=1}^n \sum_{j=1}^p \mathcal{I}_{ij}. \quad (17)$$

Here $\mathcal{I}_{ij}$ can only exist in one of two binary states, i.e., $\mathcal{I}_{ij} = \mathbf{1}_{[B_{ij}=1]} \cdot \mathcal{I}_{ij}^1 + \mathbf{1}_{[B_{ij}=\perp]} \cdot \mathcal{I}_{ij}^\perp$. To learn the values of $\mathcal{I}_{ij}^1$ and $\mathcal{I}_{ij}^\perp$, we generate a large number of minimization matrix-target utility pairs by solving the lower-level optimization for each pair. We can then learn these parameters for each index ij by averaging the target utility when $B_{ij} = 1$ and $\perp$, respectively. The final optimization thus becomes:

$$\underset{\mathbf{B} \in \{\perp, 1\}^{n \times p}}{\text{Minimize}} \quad \|\mathbf{B}\|_1 \quad \text{s.t. :} \quad \frac{1}{np} \sum_{i=1}^n \sum_{j=1}^p \mathcal{I}_{ij} - J(\hat{\theta}; \mathbf{X}, \mathbf{Y}) \leq \alpha. \quad (18)$$

When considering a sparsity constraint $\|\mathbf{B}\|_1 \leq k$ on the minimization matrix, the solution to this formulation is retaining the k entries with the highest value of the term $\mathcal{I}_{ij}^\perp - \mathcal{I}_{ij}^1$. Although the assumption of independent influence may not hold in all cases, our results (illustrated in the next section) show that this approach can substantially improve the accuracy of existing baselines.

We need to generate minimization matrix-target utility pairs. We start by generating a random minimization matrix $\mathbf{B}$, with the target sparsity k . This is the same as performing personalized random subsampling. We then solve the lower-level program for this $\mathbf{B}$ and obtain the final target utility of the trained model. We repeat these steps multiple times to create a dataset to learn the parameters of the mapping $m_\omega(\mathbf{B})$.

To learn the mapping $m_\omega(\mathbf{B})$, we simply need to learn the parameters $\mathcal{I}_{ij}^1, \mathcal{I}_{ij}^\perp$. Given the assumption of the independent influence of each index on the minimization matrix, we can learn these parameters by calculating the average loss when $B_{ij} = 1$ and $B_{ij} = \perp$:

$$\mathcal{I}_{ij}^1 = \frac{1}{|P^1|} \sum_{q=1}^{|P^1|} J(\hat{\theta}_{P_q^1}; \mathcal{X}, \mathcal{Y}) \quad \text{where} \quad P^1 := \{\mathbf{B} \mid B_{ij} = 1\} \quad (19)$$

$$\mathcal{I}_{ij}^\perp = \frac{1}{|P^\perp|} \sum_{q=1}^{|P^\perp|} J(\hat{\theta}_{P_q^\perp}; \mathcal{X}, \mathcal{Y}) \quad \text{where} \quad P^\perp := \{\mathbf{B} \mid B_{ij} = \perp\} \quad (20)$$

$$\text{where} \quad \hat{\theta}_P = \arg \min_{\theta} \frac{1}{n} \sum_{i=1}^n \ell(f_\theta(\mathbf{x}_i \odot P_i), y_i). \quad (21)$$

Evolutionary Algorithms [25]. So far, we have discussed methods that approximate the original optimization problem under various assumptions, enabling us to solve it more easily. An orthogonal class of algorithms, often applied to solve bi-level programs, are evolutionary algorithms. They trade the advantage of carrying no assumption with slow convergence, i.e., high computational demand, and higher risks of overfitting.

In our implementation of the evolutionary algorithm, we begin with a population of randomly generated minimization matrices B and then evolve them across iterations to reach our objective. We mutate and breed the current population at each stage of the process to create a larger pool of choices and carry out the lower-level optimization for every member of this pool, only retaining the best performers for the next generation. The entire process is repeated until convergence, or for a fixed number of iterations.

Mutation: In our paper, we mutate a minimization matrix B by flipping the value of exactly 10 randomly chosen indices from $1 \rightarrow \perp$, and exactly 10 randomly chosen indices from $\perp \rightarrow 1$.

Breeding: When breeding between two parent minimization matrices B^1 and B^2 , we keep the same value in the child B^c at indices where both parents agree to be the same, while we randomly choose values for indices where they don't agree to maintain target sparsity. In simpler terms,

$$B_{ij}^c \leftarrow B_{ij}^1 \text{ if } B_{ij}^1 = B_{ij}^2, \quad (22)$$

$$B_{ij}^c \leftarrow \perp \text{ with probability } k', \text{ otherwise} \quad (23)$$

where $0 \leq k' \leq 1$ is a value chosen so that the sparsity k is maintained, in expectation.

D.3 Summary

We provide a summary of various strengths and weaknesses of all algorithms used in our paper in Table 3. To support our discussion of runtime complexity, we also quantify the runtime for minimizing the Bank dataset for all 6 minimization techniques in our paper in Table 4.

E Membership inference attacks and data minimization

In this section, we focus on threat models outside the “wall”, specifically addressing inference attacks on the trained model without direct access to the minimized dataset.

E.1 Membership inference risk and inference attack

Membership Inference Risk (MIR). Membership inference attacks [24] aim to discern if an individual’s data was in the dataset before minimization, focusing on an attacker accessing the model $\hat{\theta}$ trained on this minimized dataset. In these attacks, the adversary calculates a likelihood score $L(x_q)$ for each query x_q , representing its probability of being in the original dataset $\mathbf{X}$ using the model $\hat{\theta}$. The score is given by $L(x_q) = \Pr[x_q \in \mathbf{X} | \hat{\theta}]$. Using these scores for both original dataset $\mathbf{X}$ and non-members $\mathbf{X}_{nm}$, binary membership predictions can be made at any threshold t , denoted as $\mathbf{1}_{L(x_q) \geq t}$. The overall Membership Inference Risk (MIR) is assessed as the area under the curve (AUC) of true positive rates ($\overrightarrow{\text{tpr}}$) and false positive rates ($\overrightarrow{\text{fpr}}$) across various thresholds:

$$\text{MIR} = \text{AUC}(\overrightarrow{\text{tpr}}, \overrightarrow{\text{fpr}}) \quad (24)$$

Attack Details. The attacker aims to determine the presence or absence of an individual in the original training dataset. We use the SOTA membership inference attack RMIA [35], by training 8 reference models on random 50% subsets of the public data split. We operate under the practical assumption that the adversary is unaware of the data minimization applied before training the target model and evaluate the membership inference on the original dataset.

E.2 Privacy leakage through membership inference

Finally, we assess the membership inference risk under various minimization algorithms in Figure 5(a). As previously mentioned, information leakage through a trained model is not an expected benefit of data minimization, which mainly aims to address data breach scenarios. Nevertheless, we still observe that certain minimization algorithms are effective at reducing membership inference risks with decreasing dataset size, i.e., models trained on minimized datasets leak less information. Yet, these improvements are not perfectly aligned, carrying forward the same trends we saw in the two data breach scenarios above.

	Approximating Lower-Level Program	Modeling Target Utility	Evolutionary Algorithms
Assumptions	All errors between the original and minimized data beyond the first order are considered insignificant and ignored. This might not hold when sparsity is high.	The influence of the presence or absence of every value in the data is modelled independently. This can break for datasets with highly correlated features.	No assumptions are made about the structure of the problem setting.
Hyperparameter Sensitivity	No hyperparameters.	Large number of iterations required to create better data when modelling the lower optimization.	Large number of iterations as well as a large active population size will facilitate better solutions.
Consistency	No randomness in the process, but it inherits the randomness of the lower-level learning model.	Highly inconsistent across changing randomness. But consistent across sparsity, i.e. data minimized at lower sparsity will also be minimized at higher sparsity.	Highly inconsistent across both randomness and sparsity. Data minimized may not remain minimized under changing randomness or increasing sparsity.
Convergence Behavior	Provides an exact solution under the given assumptions.	No convergence guarantees.	Convergence guarantees under a sufficiently large number of iterations.
Runtime Considerations	Closed form solution, but requires second-order derivatives. Fast for simpler settings, but does not scale well with either dataset size or model complexity.	Requires training the learning model multiple times. Relatively slower for simpler settings, but scales better with increasing complexity.	Requires training the learning model a significantly larger number of times. Furthermore, needs to be repeated for every unique output sparsity required.
Selection Criteria	Choose in simple settings with a small amount of minimization required for fast and accurate results.	Choose in more complex settings and when algorithm runtime is as important as the accuracy of the method.	Choose in a setting where the accuracy of the method is of the utmost importance, even at the cost of compute.

Table 3: A summary of strengths and weaknesses of various algorithms.

Algorithm	Runtime
Feature Selection	4 sec
Random Subsampling	3 sec
Personalized Random Subsampling	3 sec
Approximating the Target Utility	4 sec
Modelling the Target Utility	3 min
Evolutionary Algorithms	12 min
Data Pruning	4 sec

Table 4: Runtime on minimizing the Bank dataset for all 6 algorithms in our paper, as well as the data pruning technique added in the rebuttal. The timings are averaged over different runs with changing minimization percentages. The trends match our expectations as discussed in Appendix D of our paper, i.e., the approximating target utility technique is incredibly fast (as fast as the random baselines), albeit it comes at the cost of stricter assumptions. Modelling the target utility is slower, but with a lesser set of assumptions, while evolutionary algorithms are the slowest, but achieve the best performance across all methods. The data pruning technique also has a comparable runtime to the approximating target utility technique, potentially because both of these methods rely on calculating second-order derivatives.

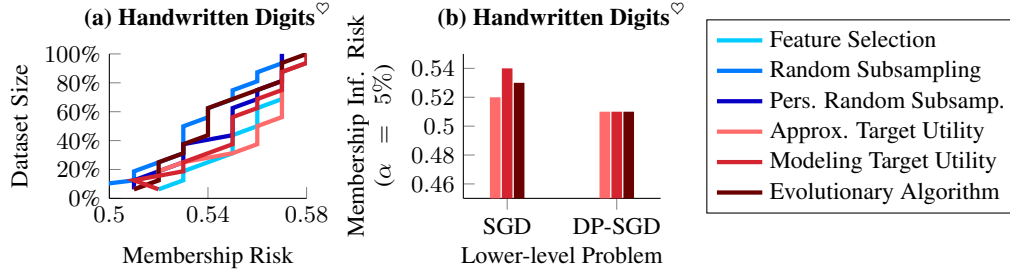


Figure 5: Membership inference risks under changing sparsity on the handwritten digits dataset. The minimization algorithms can reduce inference risks and are pushed to even better trade-offs by introducing DP-SGD in the lower-level objective.

E.3 DP-SGD to improve privacy-utility trade-off

We analyze modifications in data minimization to counter membership inference. An effective mitigation can be obtained by introducing a differentially private learning algorithm, DP-SGD [1], into the lower-level program of the bi-level optimization in equation 2. We test its compatibility with data minimization by re-evaluating membership inference in this new setting. Note, we do not introduce DP-SGD into utility calculation, i.e., once the data is minimized, the rest of the pipeline outside the “wall” remains unchanged.

Evaluating DP-SGD Modifications. By incorporating DP-SGD into the minimization optimization without altering other components, we want to assess the compatibility between these two methods. The results, collected in Figure 5(b), clearly demonstrate a reduction in membership inference risks at the same accuracy threshold, with prominent improvements for methods that were more susceptible to information leakage, such as modelling target utility algorithms. DP-SGD is indeed compatible with minimization, reinforcing the benefits of considering privacy during minimization.

F Additional experiments

F.1 Results on additional datasets

We also provide results for utility and privacy across changing sparsity on additional datasets. This includes the wine quality dataset [7], a dataset containing various attributes of 6,463 unique wines and a binary label to classify them as red or white wine, and the ACSIncome and ACSEmployment tasks of the folktables dataset [8], with census information about individuals and a label marking whether their income is above \$50,000 for ACSIncome or whether they are employed or not for ACSEmployment. Similar to 20 newsgroup dataset, we only choose a random subset of 5000 data points from both ACSIncome and ACSEmployment datasets. We also provide detailed results on the bank dataset, the handwritten digits dataset, and the 20 newsgroup dataset, in this section.

The results for utility and privacy are collected in Figure 6 and Figure 7, respectively. The trends of utility match the behaviour seen in the main text, i.e., these datasets have redundant information that can be removed without suffering any performance loss. Similarly, for privacy results, we see the trends of main text like feature selection highly misaligned with re-identification, and all methods containing high reconstruction risks even after extreme minimization, replicated in these datasets. Thus, there is a clear misalignment between data minimization and privacy expectations.

F.2 Additional results for privacy-based modifications

We provide additional results for privacy-based modifications to the data minimization algorithm on other datasets, as well as the raw trends on the handwritten dataset. The results for the handwritten dataset are collected in Figure 8, for the bank dataset are collected in Figure 9, and for the employment dataset are collected in Figure 10. The results show similar improvement as in the main text for the bank dataset, however, we don’t see similar benefits on the employment dataset. We believe that’s

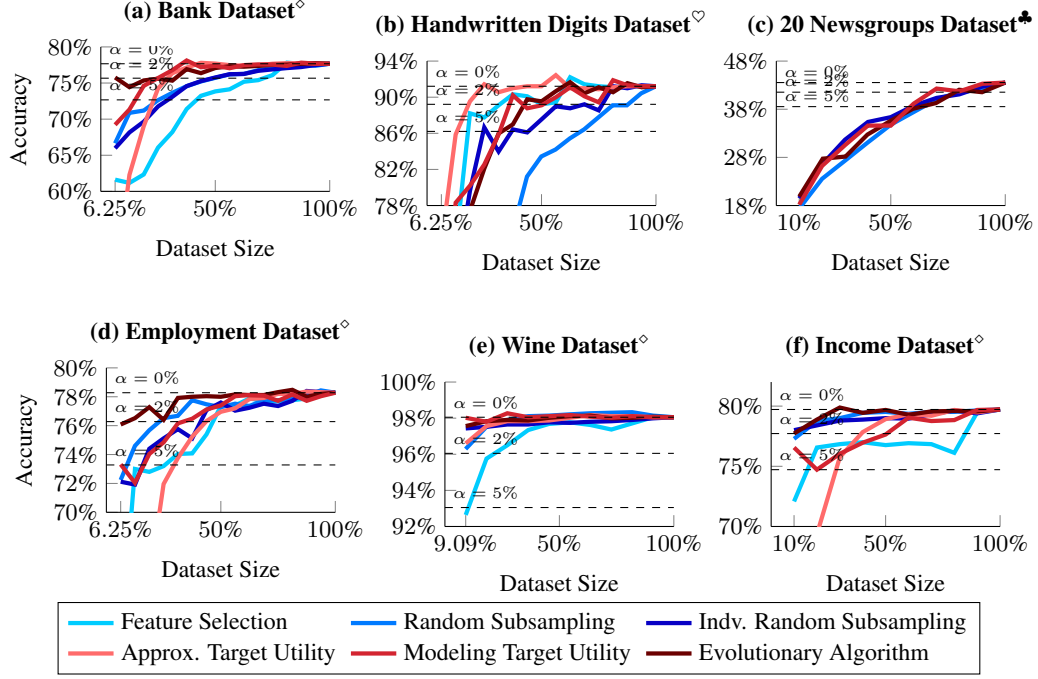


Figure 6: Utility of the minimized data across various sparsities on all datasets.

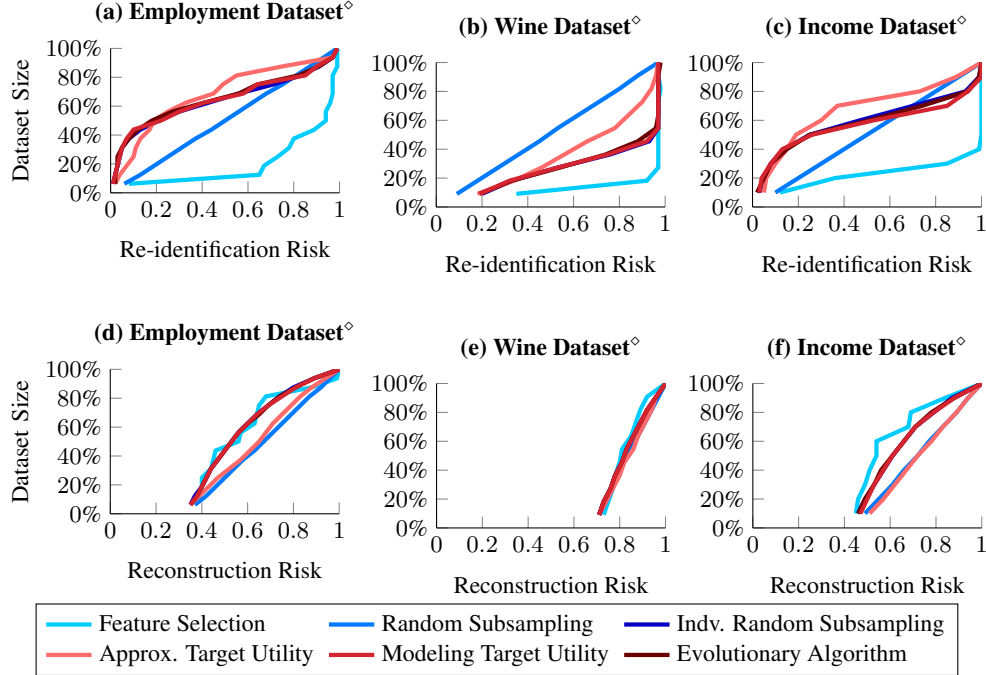


Figure 7: Re-identification and reconstruction risks under changing sparsity on additional datasets.

because these results are quite sensitive to the choice of hyperparameter β , and a more thorough search for β can show improvements for other datasets as well. Despite this discrepancy, the aim of these modifications was not to propose a novel method of incorporating privacy in minimization, but instead to highlight that minimization and privacy are compatible, and thus one can perform data

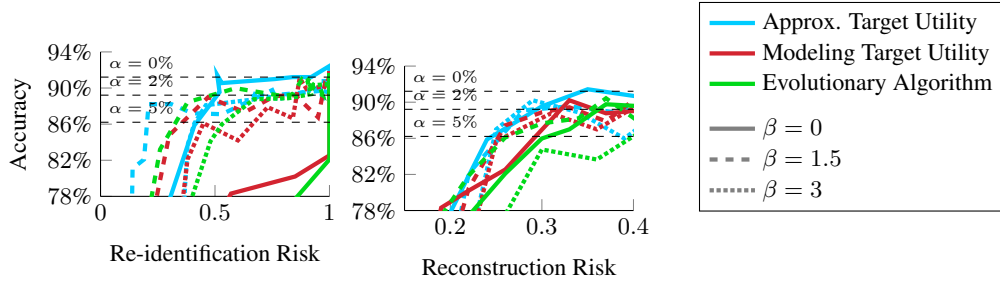


Figure 8: Re-identification and reconstruction risks (zoomed in) on handwritten digits, using feature-level privacy scores.

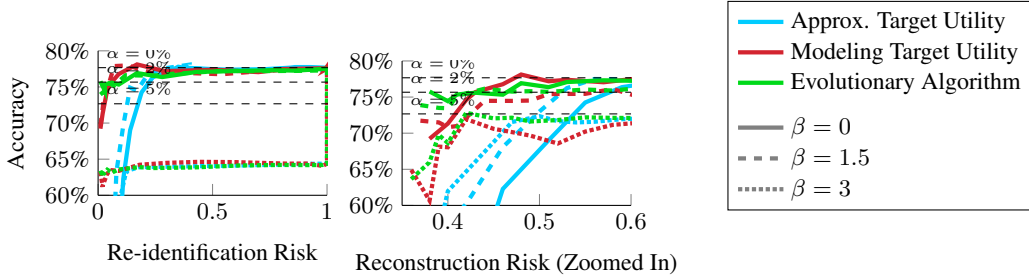


Figure 9: Re-identification and reconstruction risks (zoomed in) on the bank dataset, using feature-level privacy scores.

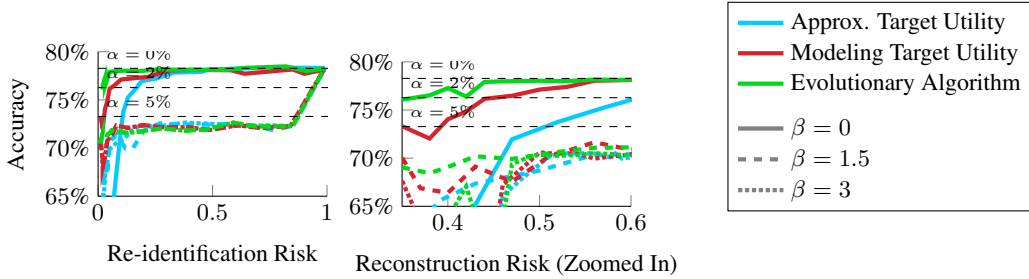


Figure 10: Re-identification and reconstruction risks (zoomed in) on ACSEmployment, using feature-level privacy scores.

minimization in line with the regulations while making sure they respect individual privacy in the dataset.

F.3 Additional comparison with Data Pruning [26]

While some existing algorithms like data pruning [26] also reduce the dataset size, they were not originally designed to operationalize data minimization, and thus might provide a guarantee of compliance within our framework, i.e., they might not be able to guarantee that the pruned data is indeed ‘minimized’. Despite this, a comparison with such algorithms can be enriching.

We provide an additional set of results after performing data pruning on the Bank dataset in Figure 11. In regards to efficacy, the data pruning method does well compared to other baselines under a stricter threshold, i.e., it is the best method to minimize the first few data points. However, the method does very poorly on looser threshold bounds, i.e., the performance drops quickly under extreme minimization. This can be attributed to the approximation of influence score calculations not being

valid under such extreme minimization. As for various privacy attacks, the trends of Data Pruning overlap with the trends of Random Subsampling, which is expected given that both of these methods are depth-based minimization algorithms, i.e., they remove complete data points. Thus, as data points are completely removed, their re-identification risk drops while data points still present have all their information and can thus be easily re-identified, hence the linear trend. The linear trend of reconstruction can be interpreted similarly.

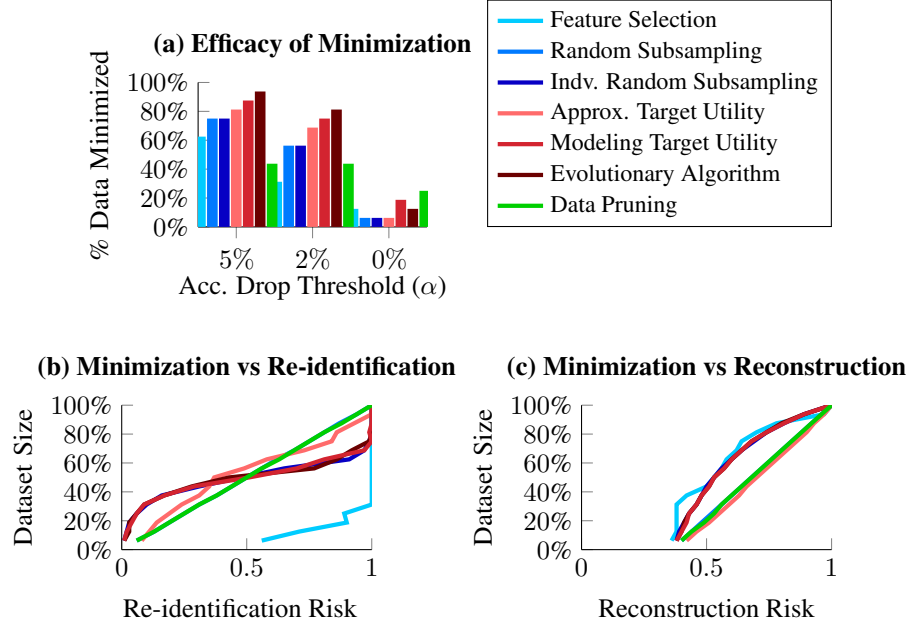


Figure 11: Additional set of results after including data pruning on the Bank dataset.