



Anonymisation Versus Encryption: A Comprehensive Evaluation of Privacy-Enhancing Technologies for Machine Learning

Jakob Folz^{1,2(✉)} , Johann Guggumos³ , Manjitha D. Vidanalage² ,
Amar Almaini² , Dalibor Fiala^{1,4} , Michael Heigl² ,
and Martin Schramm²

¹ Department of Computer Science and Engineering, Faculty of Applied Sciences,
University of West Bohemia, Technická 8., 30100 Plzeň, Czech Republic
`{jfolz,dalfia}@kiv.zcu.cz`

² Institute ProtectIT, Faculty of Computer Science, Deggendorf Institute of
Technology, Dieter-Görlitz-Platz 1., 94469 Deggendorf, Germany
`{manjitha.dahanayaka-vidanalage,amar.almaini,
michael.heigl,martin.schramm}@th-deg.de`

³ Chair for Civil Law, Liability Law and Digital Law, University of Augsburg,
Universitätsstr. 2, 86159 Augsburg, Germany
`johann.guggumos@jura.uni-augsburg.de`

⁴ CCA Group a.s., Parková 11a, 32600 Plzeň, Czech Republic

Abstract. The increasing use of open data for training machine learning models creates a tension between data utility and data privacy, particularly under the strict requirements of the General Data Protection Regulation (GDPR). Traditional anonymisation techniques aim to reduce the risk of re-identification but often impair model performance. Homomorphic encryption (HE), by contrast, allows encrypted computation while preserving data fidelity, but introduces significant computational demands. This paper presents a comparative study of anonymisation and HE using three public tabular datasets: Adult, Titanic, and OULA. Results show that HE can preserve or even improve model accuracy, with a maximum gain of 4.9 % in the Titanic dataset and a loss of only 2.3 % in the Adult dataset. However, it increases inference time by up to 9000× and storage by up to 300×. Anonymisation maintains stable accuracy with negligible performance and memory overhead. A legal analysis confirms that HE qualifies as pseudonymisation under the GDPR, offering clear benefits for compliant and secure data sharing. These findings provide practical guidance for selecting privacy-enhancing technologies in machine learning applications, taking into account legal, technical, and operational factors.

Keywords: Anonymisation · privacy · re-identification risk · GDPR · open data · homomorphic encryption · privacy enhancing technologies

1 Introduction

The rapid development in the field of artificial intelligence (AI), particularly machine learning (ML), is inextricably linked to the availability of large amounts of data. Complex models, from neural networks to ensemble methods, only reach their full potential through training with extensive and diverse datasets [22]. However, this data-intensive nature of modern ML stands in fundamental tension with the principles of European data protection law as enshrined in the General Data Protection Regulation (GDPR). The GDPR establishes a strict legal framework for the processing of personal data and prioritises the rights and freedoms of natural persons. Core principles such as data minimisation, which requires that only data necessary for the purpose may be processed, and purpose limitation, which allows processing only for specified, explicit and legitimate purposes, pose direct challenges to the “more data is better” paradigm of ML [1]. In particular, the reuse of data for new purposes, such as training an ML model that was not originally intended, requires explicit consent from the data subject or careful compatibility assessment to avoid legal violations [1]. This tension creates significant hurdles for the use of open data repositories, which harbour enormous innovation potential but often contain sensitive or personal information. At its core, this manifests as a conflict of objectives, often described as a trade-off between data utility and data privacy [6]. High data utility, which is essential for precise predictive models and personalised services, requires detailed and unaltered data. Conversely, robust privacy measures often require obfuscation, aggregation, or reduction of data, which inevitably diminishes their analytical value. Traditional anonymisation techniques, which aim to remove or generalise identifying features, often lead to significant information loss that impairs the performance of ML models [12]. Several formal privacy models have been developed to address these concerns, including k -anonymity [28, 29, 33], which ensures each record is indistinguishable from at least $k - 1$ others; l -diversity [21], which requires diverse sensitive attribute values within equivalence classes; and t -closeness [20], which maintains similar distributions between equivalence classes and the overall dataset. This work positions itself as an investigation of approaches that do not view this trade-off as an inevitable zero-sum game, but rather seek ways to achieve both goals, high utility and strong protection, simultaneously [12]. Although anonymisation and homomorphic encryption (HE) are well understood as concepts, the scientific literature lacks comprehensive, empirical comparative studies that systematically evaluate these techniques under realistic conditions for ML applications. Existing work often focuses on isolated aspects, such as analysing accuracy loss in k -anonymity [12] or comparing their own solution with privacy-preserving techniques [16], without including cryptographic methods such as HE in their comparative analyses. An integrated analysis that compares not only model quality but also systemic overhead (computation time, memory requirements) and detailed legal classification under the GDPR represents a significant research gap. From this gap emerges the central research question of this work: *What added value can pseudonymisation, partic-*

ularly through HE, contribute to GDPR-compliant provision of open data? To answer this overarching question, the following sub-questions are investigated:

1. How does the strength of anonymisation (varying k , l and t values) affect the quality (accuracy, precision, F1, recall) and efficiency of ML models?
2. What performance overhead (time, memory) do training and inference using HE generate compared to using non-anonymised and anonymised datasets?
3. To what extent is model quality preserved when applying HE, and what role does the necessary data quantisation play?
4. How should HE be legally classified compared to anonymisation as a technical and organisational measure in the sense of pseudonymisation under the GDPR?

The contribution of this work lies in a holistic, empirically grounded analysis that combines technical benchmarks with an in-depth legal assessment. Through direct comparison of non-anonymised, anonymised, and HE training on three different datasets (OULA, Adult, Titanic), quantitative data on model quality, performance overhead, and resource consumption are generated. These empirical results are then embedded within the GDPR legal framework to enable a well-founded assessment of the “added value” of each technology. The results should provide organisations that wish to provide or use open data for ML purposes with a concrete, risk-based decision-making foundation for choosing appropriate privacy protection methods.

2 Legal Aspects

2.1 Legal Classification of De-identification Processing

The GDPR distinguishes between anonymisation and pseudonymisation as forms of de-identification processes. Anonymisation, according to the GDPR, is a process through which personal data is irreversibly altered in such a way that the data subject is no longer identifiable, either directly or indirectly, by any means reasonably likely to be used [17]. Data that has been successfully anonymised is considered to fall outside the material scope of the GDPR pursuant to Recital 26 GDPR. However, the qualification of data as anonymised is subject to a strict and context-sensitive assessment [7]. This can be difficult to achieve depending on the quality and quantity of the data, as well as the intended further processing. The reason for this is the absence of guidelines from legislators and uncertainties in practical application. Furthermore, technological developments such as large language models and big data analysis can circumvent anonymisation techniques. These technologies are rapidly improving and are easy to use due to generative interaction. Pseudonymisation, as defined in Art. 4 No. 5 GDPR, entails the processing of personal data in such a manner that the data can no longer be attributed to a specific data subject without the use of additional information. These reference lists, allocation tables, or cryptographic keys, necessary for re-identification, must be stored separately and protected by

technical and organisational measures, Art. 31 (1) GDPR, to prevent identification [31]. Despite the modified state of the data, pseudonymised data is still personal data within the meaning of Art. 4 No. 1 GDPR and remains fully subject to the Regulation's provisions. Unlike anonymisation, pseudonymised data requires a legal basis like the consent of the data subject for further processing. Although pseudonymisation does not fall outside the scope of the GDPR due to the removal of personal references, it is a valuable measure that enhances data protection. Its advantages over anonymisation include the ability to trace data subjects and privileges granted by the legislator.

2.2 Homomorphic Encryption as Pseudonymisation Technique

Although not mentioned in the context of de-identification processes, encryption is an important security measure for protecting personal data against unauthorised access. Furthermore, encrypting the entire dataset can fulfil the same criteria as pseudonymisation. Instead of creating a protected part, either through an assignment table or a hash value, and a readable part, the encryption of the entire data set also encrypts the readable part. At first glance, this has no added value for pseudonymisation. Subsequent processing becomes impossible as the data is no longer in a readable form. However, the strength of HE is to process non-readable data. If datasets are pseudonymised, attackers can no longer determine the data subjects via the direct identifiers. However, it is still possible to re-identify them using indirect or quasi-identifying information [31]. The requirements on anonymisation essentially apply here, because in this constellation, as in the present case of pseudonymisation, the status of anonymity was only lifted through the analysis of the quasi-identifiers. If only the direct identifiers were deleted or pseudonymised within a data set, data subjects can still be re-identified by the indirect identifiers and the combination with additional knowledge. As a result, data protection models were introduced to limit the risks of singling out, linkability and inference [24]. When choosing pseudonymisation by creating an encrypted pseudonym, it is necessary that the encryption key, which is the additional information, is stored securely and separately at the user's system to fulfil the requirements of Art. 4 No. 5, 31 GDPR. The primary purpose of encryption is to protect against unauthorised knowledge and access and to ensure the integrity and confidentiality of data processing in accordance with Art. 5 (1) lit. f GDPR. Encryption refers to a procedure that makes information unrecognisable and restores readability to the recipient of the data using a key. Cryptography offers mathematical or algorithmic methods for this purpose. It is essential under data protection law that data processing systems are encrypted at least during the transmission and storage of information. When encrypting data, the information concerning a data subject is stored in an encrypted format such that it can only be read with the aid of a decryption key. Consequently, data stored in the cloud remains classified as personal data, as it can be decrypted and made readable at any time. This decryption can be performed by the key holder or potentially by an unauthorised third party who gains access to the decryption key. HE enables calculations to be performed on

encrypted data without requiring decryption. As a result, data can be processed in its encrypted state. When the results of these calculations are decrypted, they match those that would have been obtained if the calculations had been performed on plain data [10].

2.3 Effect of Anonymisation with Pseudonymisation

Data is regularly pseudonymised and processed by the controller. However, this does not mean that the data is also pseudonymised for third parties. Anonymity can be compromised when additional information is not disclosed. This shift to anonymity, established by the Court of Justice of the European Union (CJEU) in its decision C-413/23 P of September 4, 2025, is crucial for the implications of data transfer, particularly when encrypted data is used as a pseudonym. With this assessment by considering re-identification possibilities by the data recipient, further use cases are possible because the GDPR is no longer applicable in case of anonymised data. Indirect identifiers must nevertheless be checked for re-identification possibilities. For example, if a pseudonym is combined with the date of birth or other special attributes, there is a high risk that this information will be sufficient to re-identify the data subject. Consequently, as with anonymisation, the same problem of legal certainty arises. Pseudonymisation and unreadability of the data by means of HE closes the hurdles of secondary confidentiality because no information can be extracted from the data. From a data sharing perspective like open data, there is great potential, because the highly secure, encrypted data can be stored with the user's permission on a third-party server and from there shared for other purposes, such as training data for machine learning.

2.4 Evaluation of De-identification Processes

When evaluating de-identification techniques, it is essential to weigh up the data protection of those affected and the usability of the data. The better the data protection, the less valuable the data. Anonymisation can provide a strong data protection because the personal reference is removed and those affected can no longer be re-identified. However, this could require qualitative data to be deleted, and it is difficult to achieve. The same applies to conventional pseudonymisation, except that a pseudonym replaces the direct identifiers. As shown above, encrypting the entire dataset can also be a pseudonymisation technique if the key is stored separately. If the data is used to train AI models, it can be seen that encrypting the entire data set using HE is a successful alternative to anonymisation. The loss of information caused by anonymisation techniques is avoided. As a result, more meaningful results can be generated.

3 Related Work

The effects of k -anonymity on ML model performance have been extensively studied in the research community. Studies utilising the Adult (Census Income)

dataset consistently demonstrate that classification accuracy decreases as the k -value increases. However, these accuracy losses are often deemed “acceptable” for moderate k -values (e.g., up to 100) [12]. This body of work has established a relatively mature understanding of the k -anonymity trade-off. Notably, some investigations suggest that performance losses can remain within acceptable bounds even at relatively high k -values, making the method viable for certain applications [12]. Furthermore, studies have revealed that robustness to anonymised data varies by ML algorithm, with some models demonstrating greater resilience than others [34]. The landscape of HE-based ML research is rapidly evolving, with much of the work focusing on technical implementation challenges. Performance analyses, including those published by the developers of Concrete-ML, provide initial insights into expected overhead. They report inference times ranging from milliseconds for simple linear models to several seconds for complex ensemble models, such as XGBoost [35]. Simultaneously, recent work demonstrates that HE is being adapted for sophisticated architectures such as Kolmogorov-Arnold Networks (KANs), underscoring the technology’s increasing maturity and relevance [19]. A key focus in HE research has been the trade-off between model complexity (e.g., neural network depth) and computational overhead [13]. Earlier work often resorted to HE-friendly network architectures with approximated, simple activation functions, which inherently limited accuracy [25]. However, modern approaches, such as those implemented in Concrete-ML through programmable bootstrapping and lookup tables, enable exact emulation of more complex functions, significantly improving accuracy. The central insight is that HE can be configured to achieve accuracy nearly identical to models trained on quantised non-anonymised data. However, this preservation of accuracy comes at the cost of significant execution time overhead [32]. A complementary line of research explores architectural optimisations for HE systems. In our prior work, we proposed a scalable, distributed HE scheme tailored for high-computational complexity models [5]. The approach distributes encrypted computation across multiple nodes to overcome performance bottlenecks inherent in centralised HE systems. While that study focused on improving throughput and supporting deeper models via parallelism, it did not benchmark HE against alternative privacy approaches. In contrast, the present work emphasises comparative evaluation by analysing HE and anonymisation side by side under identical conditions and within a legal compliance context. It is crucial to emphasise that the prevailing use case for HE in machine learning is privacy-preserving inference. The typical workflow involves training a model on unencrypted plaintext data and then performing inference on encrypted data [4]. Users can encrypt their sensitive data, send it to the server, and receive an encrypted prediction result, thus preserving the privacy of their query [3]. This focus on inference stems from the fact that training complex models directly on HE data is considered extremely computationally intensive and often impractical due to performance overhead; some studies explicitly identify training as one of the remaining limitations of HE [23]. In contrast to this common paradigm, the present work deliberately investigates a more challenging yet potentially

more secure approach, where both training and inference are performed entirely on encrypted data. The availability of specialised frameworks like Concrete-ML, which explicitly offer functions for training on encrypted data, makes such empirical investigation feasible. Despite the wealth of research on both technologies individually, there are very few, if any, works that conduct a direct, multi-metric, empirical comparison of anonymisation and HE using modern, practice-relevant tools. While some survey papers mention both techniques as alternatives in the Privacy-Preserving Machine Learning spectrum, they are rarely benchmarked against each other under identical conditions and on a unified metric basis [13]. This work addresses precisely this gap by providing a holistic comparison that considers not only predictive accuracy but also the entire pipeline costs in terms of time and memory, factors that are crucial for practical implementation and informed technology decisions.

3.1 Experimental Design

All experiments were conducted on a Linux system (6.1.0-25-amd64) with a 13th Gen Intel Core i7-13700HX processor (16 cores, 24 logical processors) and 64 GB RAM. We used Python 3.10 with Concrete ML 1.x, which implements fully homomorphic encryption (FHE) using the Fast Fully Homomorphic Encryption over the Torus scheme, alongside scikit-learn 1.x for machine learning models. To ensure statistical reliability and reproducibility of results, we employed a comprehensive experimental protocol. Importantly, our objective was not to achieve optimal accuracy for each dataset, but rather to ensure fair comparability between different privacy-preserving methods within each dataset. Each experiment was repeated 10 times to obtain robust statistical measures. For random seed management, we used a base random state of 42 with run-specific offsets to ensure reproducibility while maintaining independence between runs. The SGDClassifier was trained for 15 iterations in each experiment. We utilised a 70%-30% train-test split ratio with stratified sampling to maintain class distribution consistency across splits. All FHE experiments were conducted with real encryption execution rather than simulation mode, ensuring that our measurements reflect actual computational overhead and encryption behaviour. After completing all runs, we calculated average values across the 10 repetitions for each metric.

3.2 Timing Methodology

To systematically record execution times, we established distinct measurement categories across the experimental pipeline. Processing time encompasses data loading, categorical encoding, type conversions, and train-test splitting. Normalisation time captures the MinMaxScaler fitting and transformation to the range $[-1, 1]$. Training time measures model training duration, which differs between plain mode (SGDClassifier training on clear data) and encrypted mode (SGDClassifier training with FHE). For encrypted experiments, we additionally

measure compilation time for FHE circuit compilation. Prediction time quantifies inference duration on the test set, again distinguishing between plain mode (inference on unencrypted data) and encrypted mode (inference on encrypted data). Finally, total time represents the end-to-end execution duration across all stages.

3.3 Classification Metrics

To evaluate and compare the performance of plain, encrypted, and anonymised models, we employed four key metrics that provide comprehensive insights into classification performance. Given the class imbalance in our datasets, we utilise weighted averaging for all metrics except accuracy. This weighting scheme ensures that each class contributes to the final metric proportionally to its frequency in the dataset. All metrics were computed using scikit-learn’s implementation with the `average='weighted'` parameter [30]. For precision, recall, and F1-score, the weighted averaging is calculated using C classes, where n_i is the number of true instances for class i , N is the total number of instances, and Metric_i represents the respective metric (Precision $_i$, Recall $_i$, or F1 $_i$) for class i :

$$\text{Metric}_{\text{weighted}} = \sum_{i=1}^C \frac{n_i}{N} \text{Metric}_i, \quad (1)$$

Accuracy measures the proportion of all correctly classified instances. While straightforward to interpret, accuracy can be misleading for imbalanced datasets as the majority class may dominate it:

$$\text{Accuracy} = \frac{\text{Number of correct predictions}}{\text{Total number of predictions}}, \quad (2)$$

Precision quantifies the proportion of positive predictions that are actually correct. This metric is crucial when false positives are costly, such as in spam detection where legitimate emails should not be marked as spam. In our context, precision helps assess whether privacy-preserving methods maintain reliable positive predictions. For each class i :

$$\text{Precision}_i = \frac{\text{TP}_i}{\text{TP}_i + \text{FP}_i}, \quad (3)$$

Recall (also known as sensitivity) measures the proportion of actual positive instances that are correctly identified. This metric is critical when false negatives have serious consequences, such as in disease diagnosis. For our evaluation, recall indicates whether encryption or anonymisation causes models to miss important positive cases:

$$\text{Recall}_i = \frac{\text{TP}_i}{\text{TP}_i + \text{FN}_i}, \quad (4)$$

F1-Score represents the harmonic mean of precision and recall. This balanced metric provides a single value that reflects both precision and recall, making it particularly valuable when comparing different privacy-preserving methods

as it captures potential trade-offs between these two aspects:

$$F1_i = \frac{2 \text{Precision}_i \text{Recall}_i}{\text{Precision}_i + \text{Recall}_i}. \quad (5)$$

The weighted averaging approach is essential for our datasets which exhibit significant class imbalance: the Adult dataset shows approximately 75% *more than 50K* versus 25% *less than 50K* income distribution; the Titanic dataset contains roughly 38% survivors versus 62% non-survivors. Without weighting, the majority class would dominate the metrics, obscuring the model’s actual performance on minority classes.

3.4 Storage Measurement Methodology

To measure storage requirements across different privacy-preserving methods, we employed two distinct approaches tailored to the data type. For unencrypted data (both non-anonymised and anonymised versions), we directly measured the actual file size on disk using `file_size = getsize(file_path)/(1024 × 1024)`, converting the result to megabytes for consistency. However, for FHE data, direct file measurement proved inadequate due to the nature of encrypted computations. Instead, we utilised Concrete ML’s circuit statistics to accurately estimate the encrypted data size. This approach calculates the output size per sample from circuit statistics (`output_size_per_sample = circuit.statistics['size_of_outputs']`), then multiplies by the total number of samples across both training and test sets (`total_samples = len(xtrain) + len(xtest)`). The resulting total size is then converted to megabytes (`total_size/(1024 × 1024)`) to maintain consistency with our unencrypted measurements.

4 Data Preparation

4.1 Dataset Introduction

For the practical evaluation, we utilised publicly available tabular datasets, specifically the Adult dataset [9], Open University Learning Analytics Dataset (OULA) [18], and the Titanic Dataset [11]. The Adult dataset, a person-level dataset sourced from the UC Irvine Machine Learning Repository, was extracted from the 1994 U.S. Census Bureau database by Ronny Kohavi and Barry Becker. This dataset contains 48,842 records and is widely utilised for binary classification tasks to predict whether an individual earns more than \$50,000 annually. The *income* attribute serves as the target variable, with two classes: *more than 50K* and *less than 50K*. The OULA dataset, provided by The Open University, which has the largest number of Bachelor’s students in the UK, includes data from seven modules of the university’s online learning platform. The dataset encompasses student demographics, assessment scores, and interactions with the Virtual Learning Environment, amounting to 32,593 records across 12 attributes. The attribute *final_result* is designated as the target for the

evaluation. The Titanic dataset is a classic machine learning dataset containing information about passengers aboard the RMS Titanic, which sank on April 15, 1912. The dataset comprises 1,309 passengers (rows) with 12 attributes. The goal is to predict the binary variable *Survived*, which indicates whether a passenger survived the disaster (1) or not (0).

4.2 Feature Selection

Adult dataset: All original features were retained despite *age*’s notably high importance. This decision recognised the fundamental distinction between legitimate predictive features and data leakage, where *age* represents demographic information available prior to income determination with natural causal relationships through accumulated work experience and career progression.

OULA: Two problematic features were removed to eliminate data leakage. The *score* feature directly determines the *final.result* target variable, whilst *date_unregistration* captures post-failure behaviour where students withdraw after poor performance, indirectly revealing outcomes.

Titanic: The most extensive reduction from 21 to 10 features (52% reduction) eliminated multiple categories of problematic features. Clear data leakage indicators included *boat* (62.9% missing), which directly indicates survival through lifeboat access, and *body* (90.8% missing), representing post-disaster information. Complex features like *name* were parsed into separate columns (*first name*, *last name*, *middle name*, *husband name*) but provided no predictive value. Similarly, *home.dest* was split into geographic components (*city_1*, *city_2*, *state_1*, *state_2*) but suffered from high missing rates (43.1–88.5%) with minimal predictive benefit.

4.3 Anonymisation

The datasets were anonymised using the ARX Data Anonymisation Tool [26, 27], following established procedures for tabular data anonymization as synthesized in [8]. Privacy parameters were configured to ensure $k > 10$ and $t \leq 0.5$. The ARX framework generates multiple anonymised dataset variants through the systematic application of generalisation hierarchies to quasi-identifiers (attributes that can potentially identify individuals when combined). From these variants, three distinct anonymised datasets were selected based on ascending k -anonymity values, with the selection prioritising the preservation of feature importance rankings derived from pre-computed feature importance analysis. This approach ensures that attributes with higher predictive value undergo minimal generalisation while maintaining requisite privacy guarantees. The anonymisation process implements a dual privacy criterion framework: k -anonymity provides protection against record linkage attacks by ensuring each equivalence class contains at least k records, while t -closeness mitigates attribute disclosure risks, including homogeneity attacks, background knowledge attacks, skewness attacks, and similarity attacks [14] by ensuring the sensitive value distribution of the equivalence class similar to the dataset. This multi-layered

approach addresses both identity disclosure and attribute disclosure vulnerabilities inherent in microdata publication. For the Adult Census dataset, the quasi-identifiers comprised *age*, *class of worker*, *education*, *marital stat*, *race*, *sex*, *full or part time employment stat*, *has_children*, *has_grandchildren*, *has_relatives*, *has_married*, *is_wife*, *is_child*, *has_citizenship*, and *is_native*, with *label* as the sensitive attribute. A comprehensive privacy metric evaluation was conducted for both non-anonymised and anonymised datasets, measuring k -anonymity, l -diversity, and t -closeness parameters. The complete privacy assessment results are presented in Table 1. The OULA dataset anonymisation utilised *gender*, *region*, *highest_education*, *age_band*, and *disability* as quasi-identifiers, with *final_result* serving as the sensitive attribute. For the anonymisation process of the Titanic dataset, the quasi-identifiers selected were *title*, *sex*, and *age*, with *survived* as the sensitive attribute.

Table 1. Privacy metrics of the selected datasets

Dataset	Version	k -anonymity	l -diversity	t -closeness
Adult	Non-anonymised	1	1	0.972
	Anonymised-1	12	2	0.176
	Anonymised-2	20	2	0.088
	Anonymised-3	80	2	0.283
OULA	Non-anonymised	3	1	0.5
	Anonymised-1	12	2	0.5
	Anonymised-2	22	2	0.5
	Anonymised-3	42	2	0.5
Titanic	Non-anonymised	1	1	0.618
	Anonymised-1	13	2	0.228
	Anonymised-2	18	2	0.451
	Anonymised-3	58	2	0.054

5 Evaluation

5.1 Accuracy

The left half of Table 2 presents performance metrics across privacy-preserving techniques using a blue gradient heatmap. FHE encryption shows varied performance impacts: Adult dataset demonstrates resilience with minimal degradation (94.6% vs 96.8% non-anonymised, -2.28%), whilst OULA experiences substantial loss (52.8% vs 62.3% non-anonymised, -15.32%). Surprisingly, Titanic improves under encryption (74.4% vs 71.0% non-anonymised, $+4.88\%$), suggesting quantisation acts as regularisation for this small dataset. Anonymised

Table 2. Performance and timing comparison across privacy-preserving techniques

Version	Accuracy	Precision	Recall	F1	Training (s)	Prediction (s)
Adult dataset						
Non-anonymised	0.968	0.969	0.968	0.967	0.5000	0.0110
Encrypted	0.946	0.945	0.946	0.944	744.0	1008.0
Anonymised-1	0.972	0.963	0.972	0.958	0.5100	0.0110
Anonymised-2	0.972	0.946	0.972	0.958	0.5100	0.0110
Anonymised-3	0.971	0.960	0.971	0.959	0.5100	0.0110
OULA dataset						
Non-anonymised	0.623	0.598	0.623	0.590	0.6000	0.0260
Encrypted	0.528	0.544	0.528	0.463	438.0	582.0
Anonymised-1	0.624	0.613	0.624	0.556	0.6100	0.0260
Anonymised-2	0.623	0.607	0.623	0.578	0.6100	0.0260
Anonymised-3	0.621	0.604	0.621	0.571	0.6100	0.0260
Titanic dataset						
Non-anonymised	0.710	0.749	0.710	0.699	0.0040	0.0002
Encrypted	0.744	0.762	0.744	0.741	263.1	3.16
Anonymised-1	0.596	0.610	0.596	0.594	0.0040	0.0002
Anonymised-2	0.759	0.771	0.759	0.746	0.0040	0.0002
Anonymised-3	0.634	0.665	0.634	0.596	0.0040	0.0002

variants occasionally outperform non-anonymised datasets through generalisation effects that reduce overfitting. Adult anonymised datasets (97.1–97.2%) slightly exceed non-anonymised performance, whilst OULA maintains comparable results (62.1–62.4%). Titanic shows high variance across k values (59.6–75.9%), highlighting small dataset sensitivity to anonymisation parameters. The results demonstrate that privacy-preserving techniques can achieve suitable performance levels, with technique selection depending on dataset characteristics and accuracy requirements. Large datasets favour FHE with acceptable degradation, whilst anonymisation provides consistent performance with minimal computational overhead.

5.2 Time

The right half of Table 2 presents execution times using logarithmic colour scaling to accommodate vast differences between plain and encrypted operations. FHE encryption incurs drastic computational overhead: training times increase by 3–5 orders of magnitude (Adult: 1,488 $\times$, OULA: 730 $\times$, Titanic: 65,827 $\times$), whilst prediction times show even greater penalties (Adult: 91,436 $\times$, OULA: 22,440 $\times$, Titanic: 15,800 $\times$). Anonymisation demonstrates minimal timing overhead, with processing and training times remaining within 2% of non-anonymised

execution across all datasets. This efficiency stems from anonymisation being a preprocessing step that preserves the original computational model structure. The results reveal a fundamental trade-off between computational efficiency and privacy guarantees. Whilst anonymisation offers near-zero overhead suitable for real-time applications, FHE provides stronger cryptographic privacy at the cost of 4–5 orders of magnitude slowdown, making it suitable only for offline processing where privacy is paramount.

5.3 Storage

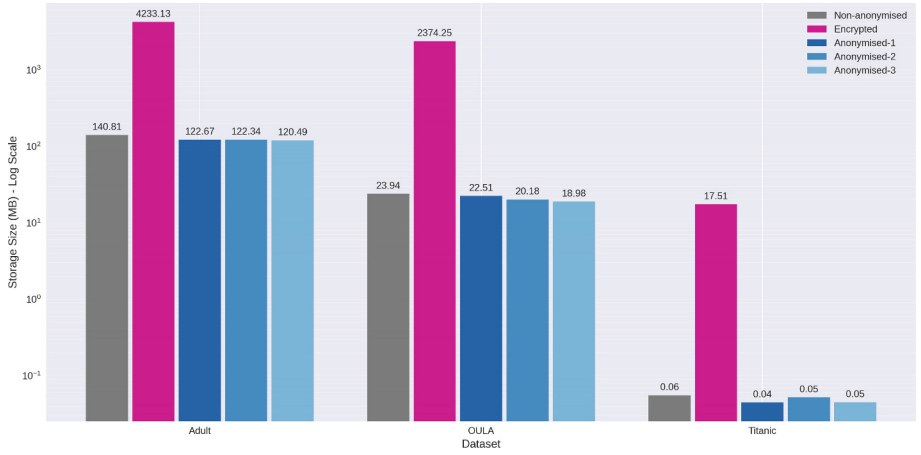


Fig. 1. Storage costs across privacy-preserving techniques

Figure 1 presents storage requirements across privacy-preserving techniques using a logarithmic scale due to extreme variations (0.06 MB to 4.2 GB). FHE Encryption demonstrates substantial storage overhead with expansion factors of $30.1\times$ (Adult), $99.2\times$ (OULA), and $291.8\times$ (Titanic). The inverse relationship between dataset size and expansion ratio suggests that FHE overhead includes significant fixed components that disproportionately affect smaller datasets. FHE inherently transforms each data point into large ciphertexts (polynomials with thousands of coefficients) and requires fixed-size cryptographic keys for operations, creating substantial storage overhead independent of the actual data volume. Anonymisation shows minimal storage impact, maintaining requirements nearly identical to non-anonymised data across all datasets. Whilst this efficiency makes k-anonymisation attractive for storage-constrained environments, it provides weaker privacy guarantees than FHE’s cryptographic security. The results highlight the fundamental trade-off between privacy strength and storage efficiency in privacy-preserving ML deployments.

5.4 Privacy Guarantees

The privacy guarantees of anonymisation and HE differ fundamentally in their approach and strength, providing distinct protection levels for data subjects under the GDPR framework. Our implementation combines k -anonymity with t -closeness to address multiple privacy attack vectors. k -anonymity ensures each record is indistinguishable from at least $k - 1$ others, creating equivalence classes that protect against record linkage attacks (matching records across datasets using quasi-identifiers). t -closeness complements this by ensuring sensitive attribute distributions within each equivalence class closely match the overall distribution, mitigating homogeneity attacks (all records in a class sharing identical sensitive values), background knowledge attacks (exploiting external information about individuals), skewness attacks (using non-uniform sensitive attribute distributions), and similarity attacks (exploiting semantic relationships between attribute values). However, anonymisation provides probabilistic rather than absolute protection. The techniques remain vulnerable to composition attacks combining multiple datasets and future advances in re-identification methods. The irreversible nature means protection levels cannot be adjusted post-publication. HE provides cryptographic security based on computational infeasibility, with our Concrete ML 1.x TFHE implementation. Unlike anonymisation, which reveals generalised values, HE ensures no individual record information is disclosed, as encrypted data appears as cryptographically indistinguishable pseudorandom noise. Our encryption offers a 128-bit security level [36], which is infeasible for brute force attacks. Its privacy guarantee is anchored in the Ring Learning With Errors problem, a post-quantum secure mathematical foundation [2, 15] that requires super-exponential computational complexity to break and maintains resistance against both classical and quantum adversaries. The cryptographic barrier prevents all statistical disclosure attacks that affect anonymisation techniques, including differential attacks, linkage attacks, and inference attacks that exploit data correlations. HE provides inherent protection against composition attacks, as encrypted datasets cannot be meaningfully combined without decryption keys, whilst maintaining 128-bit post-quantum security guarantees throughout complex computations. Anonymisation provides privacy through statistical obscurity, which makes re-identification difficult but not impossible. HE provides privacy through cryptographic certainty, rendering unauthorised access computationally infeasible rather than merely improbable. This represents a fundamental security advantage: whilst anonymisation remains vulnerable to probabilistic privacy attacks that exploit statistical patterns and correlations, HE renders such attacks entirely infeasible through cryptographic protection. For GDPR compliance, this distinction motivates the adoption of HE for high-sensitivity contexts where absolute privacy protection is required, despite the computational overhead. Anonymisation removes data from regulatory scope when successful but risks future compromise, whilst HE maintains stronger protection for long-term data utilisation, whilst keeping data within regulatory scope.

5.5 Implications for Open Data Provision

The fundamental question for open data provision is what happens after training on encrypted data. While training can proceed using only evaluation keys, accessing the trained model’s insights inevitably requires the private key at some point in the pipeline. This necessity creates two distinct deployment pathways. First, the decrypted model approach, where the private key is used immediately after training to decrypt the model weights. This yields a regular plaintext model that functions like any conventional ML model and can make predictions on standard unencrypted data. Second, the fully encrypted pipeline: Here, the private key is deferred until the very end, while the trained model remains encrypted, accepts encrypted inputs, and produces encrypted outputs. Only when final predictions need to be interpreted does the private key decrypt the results. This maintains end-to-end encryption but requires all users to operate within the same encryption domain. For open data initiatives, this represents an important architectural consideration. HE datasets offer a different paradigm than traditional open access. The private key serves as an essential safeguard, ensuring that only authorised parties can access the final sensitive results, such as the trained model parameters or individual predictions. This shifts the open data concept from releasing (often inadequately anonymised) raw data toward providing secure, encrypted analytical environments where computations can be performed while maintaining cryptographic protection. The final insights remain under the control and protection of the key holder, offering a controlled computation model rather than unrestricted data access. This approach complements existing methods like k-anonymisation, each serving different use cases within the spectrum of privacy-preserving data sharing.

6 Conclusion and Future Work

This paper presented a comprehensive evaluation of anonymisation and HE as privacy-enhancing techniques for GDPR compliant machine learning on tabular open datasets. Using three public datasets (Adult, OULA, and Titanic), we assessed their impact on model accuracy, execution time, and storage requirements within a unified experimental setup. The results demonstrate that anonymisation preserves high model quality while introducing minimal computational and storage overhead, making it a practical choice for many applications. HE achieves comparable or even improved accuracy, with a gain of up to 4.9 percent in the Titanic dataset, but incurs significant computational and storage costs, with inference times increasing by up to ninety thousand times and storage requirements by up to three hundred times. However, regarding our central question about HE’s value for open data provision, our analysis reveals that the private key requirement fundamentally limits true open data scenarios, shifting instead toward controlled computation environments. A legal analysis confirmed that HE qualifies as a form of pseudonymisation under the GDPR, enabling secure data sharing while maintaining regulatory compliance. Importantly, pseudonymisation can have an anonymising effect when the data

recipient lacks access to the additional information required for re-identification. This strengthens HE's value in high-security contexts, despite the computational burden. Furthermore, HE processes raw data without information loss, unlike anonymisation which requires generalisation. In scenarios where sufficient k -groups cannot be formed to guarantee anonymity, HE presents a promising alternative for secure data sharing, as it maintains data utility while providing cryptographic protection. Future work will expand this comparative framework to include additional privacy-enhancing technologies such as synthetic data generation and differential privacy, providing a more comprehensive landscape of available techniques. We envision developing a decision support system that recommends the optimal privacy-preserving method based on specific use case requirements, considering factors such as data utility needs, computational constraints, regulatory requirements, and security guarantees. This work serves as a foundation for such a system by establishing benchmarking methodologies and evaluation criteria across multiple dimensions. Further legal exploration is also needed to incorporate emerging interpretations of anonymisation standards within the EU. Finally, broader usability studies and real-world deployments could help identify practical integration challenges and drive adoption of privacy technologies in machine learning workflows. Further legal exploration is needed to comprehensively assess how different privacy-enhancing technologies align with GDPR requirements, examining not only emerging interpretations of anonymisation standards but also the legal classification of synthetic data, differential privacy, and other techniques under EU data protection law. This legal analysis would form an integral part of the decision support framework, ensuring that recommendations consider both technical performance and regulatory compliance. Finally, broader usability studies and real-world deployments could help identify practical integration challenges and drive adoption of privacy technologies in machine learning workflows.

Acknowledgments. The research project EAsyAnon (“Verbundprojekt: Empfehlungs- und Auditsystem zur Anonymisierung.”, funding indicator: 16KISA128K) received funding from the German Federal Ministry of Research, Technology and Space (BMFTR) under the umbrella of the funding guideline “Forschungsnetzwerk Anonymisierung für eine sichere Datennutzung” and is financed by the package NextGeneration EU.

References

1. Ademokun, O.V.: AI and the GDPR: Understanding the foundations of compliance. <https://techgdp.com/blog/ai-and-the-gdpr-understanding-the-foundations-of-compliance/>
2. Ahmaddunnisa, S., Mathe, S.E.: Multi-LFSR architectures for BRLWE-based post quantum cryptography. *IEEE Access* (2024)
3. Akram, A., Khan, F., Tahir, S., Iqbal, A., Shah, S.A., Baz, A.: Privacy preserving inference for deep neural networks: optimizing homomorphic encryption for efficient and secure classification. *IEEE Access* **12**, 15684–15695 (2024)

4. Badawi, A., Faizal Bin Yusof, A., M.: Private pathological assessment via machine learning and homomorphic encryption. *BioData Mining* **17**(1), 33 (2024)
5. Almaini, A., Folz, J., Woelfl, D., Al-Dubai, A., Schramm, M., Heigl, M.: A new scalable distributed homomorphic encryption scheme for high computational complexity models. In: 2023 International Wireless Communications and Mobile Computing (IWCMC), pp. 890–897 (2023). <https://doi.org/10.1109/IWCMC58020.2023.10183131>
6. Ammara: Data privacy vs. data utility: the tug of war in data science. <https://medium.com/@160623747003/data-privacy-vs-data-utility-the-tug-of-war-in-data-science-00d719be2344>
7. Auer-Reinsdorff, A., Conrad, I. (eds.) *Handbuch IT- und Datenschutzrecht*, 3, auflage C.H. Beck, München (2019)
8. Aufschläger, R., Folz, J., März, E., Guggumos, J., Heigl, M., Buchner, B., Schramm, M.: Anonymization procedures for tabular data: an explanatory technical and legal synthesis. *Information* **14**(9), 487 (2023)
9. Barry Becker, R.K.: Adult 10.24432/C5XW20. <https://archive.ics.uci.edu/dataset/2>
10. Bitkom: Anonymisierung und pseudonymisierung von daten für projekte des maschinellen lernens (2020)
11. Dawson, R.J.M.: The 'unusual episode' data revisited (1995). <https://jse.amstat.org/v3n3/datasets.dawson.html>. Dataset obtained from extended compilation available at: <https://www.kaggle.com/datasets/vinicius150987/titanic3>
12. Díaz, J.S.P., García, Á.L.: Comparison of machine learning models applied on anonymized data with different techniques. In: 2023 IEEE International Conference on Cyber Security and Resilience (CSR), pp. 618–623. IEEE (2023)
13. Fang, H., Qian, Q.: Privacy preserving machine learning with homomorphic encryption and federated learning. *Future Internet* **13**(4), 94 (2021)
14. Folz, J., Vidanalage, M.D., Aufschläger, R., Almaini, A., Heigl, M., Fiala, D., Schramm, M.: Scoring system for quantifying the privacy in re-identification of tabular datasets. *IEEE Access* (2025)
15. He, P., Guin, U., Xie, J.: Novel low-complexity polynomial multiplication over hybrid fields for efficient implementation of binary ring-lwe post-quantum cryptography. *IEEE J. Emerg. Selected Top. Circ. Syst.* **11**(2), 383–394 (2021)
16. Krishnamoorthy, M.V.: Data obfuscation through latent space projection (LSP) for privacy-preserving AI governance: case studies in medical diagnosis and finance fraud detection 10.48550/arXiv. 2410.17459. <http://arxiv.org/abs/2410.17459>. Version: 1
17. Kühling, J., Buchner, B. (eds.): *Datenschutz-Grundverordnung BDSG: Kommentar*, 4, auflage C.H.Beck, München (2024)
18. Kuzilek, J., Hlosta, M., Zdráhal, Z.: Open university learning analytics dataset (2017)
19. Lai, Z., Zhou, Y., Zheng, P., Chen, L.: Efficient privacy-preserving KAN inference using homomorphic encryption (2024). arXiv preprint
20. Li, N., Li, T., Venkatasubramanian, S.: t-closeness: Privacy beyond k-anonymity and l-diversity. In: 2007 IEEE 23rd International Conference on Data Engineering, pp. 106–115. IEEE (2006)
21. Machanavajjhala, A., Kifer, D., Gehrke, J., Venkatasubramanian, M.: l-diversity: Privacy beyond k-anonymity. *ACM Trans. Knowl. Discovery Data (TKDD)* **1**(1), 3-es (2007)

22. Office of the Victorian Information Commissioner: Artificial intelligence and privacy - issues and challenges. <https://ovic.vic.gov.au/privacy/resources-for-organisations/artificial-intelligence-and-privacy-issues-and-challenges/>
23. Onoufriou, G., Hanheide, M., Leontidis, G.: EDLaaS: Fully homomorphic encryption over neural network graphs for vision and private strawberry yield forecasting 10.48550/arXiv. 2110.13638. <http://arxiv.org/abs/2110.13638>
24. Party, A.D.P.W.: Opinion 5/2014 on anonymisation techniques, wp216. Adopted on 10 April 2014
25. Popescu, A.B., Taca, I.A., Nita, C.I., Vizitiu, A., Demeter, R., Suciu, C., Itu, L.M.: Privacy preserving classification of EEG data using machine learning and homomorphic encryption. *Appl. Sci.* **11**(16), 7360 (2021)
26. Prasser, F., Kohlmayer, F., Lautenschläger, R., Kuhn, K.A.: Arx-a comprehensive tool for anonymizing biomedical data. In: *AMIA Annual Symposium Proceedings*, vol. 2014, p. 984. American Medical Informatics Association (2014)
27. Prasser, F., Kohlmayer, F., Lautenschläger, R., Kuhn, K.A.: Arx—a comprehensive tool for anonymizing biomedical data. In: *AMIA Annual Symposium Proceedings*, vol. 2014, p. 984. American Medical Informatics Association (2014). <https://arx.deidentifier.org/>
28. Samarati, P.: Protecting respondents identities in microdata release. *IEEE Trans. Knowl. Data Eng.* **13**(6), 1010–1027 (2001). <https://doi.org/10.1109/69.971193>
29. Samarati, P., Sweeney, L.: Protecting privacy when disclosing information: k-anonymity and its enforcement through generalization and suppression (1998)
30. Scikit-Learn Developers: scikit-learn: Metrics and scoring—classification metrics. https://scikit-learn.org/stable/modules/model_evaluation.html#classification-metrics (2024). Accessed 16 Jan 2025
31. Simitis S, Hornung G (2025) Spiecker genannt. In: Döhmman I (ed) *Datenschutzrecht: DSGVO mit BDSG*, 2. auflage edn, Nomos (2025)
32. Stoian, A., Frery, J., Bredehoft, R., Montero, L., Kherfallah, C., Chevallier-Mames, B.: Deep neural networks for encrypted inference with TFHE (2023). 10.48550/arXiv. 2302.10906. <http://arxiv.org/abs/2302.10906>
33. Sweeney, L.: k-anonymity: a model for protecting privacy. *Internat. J. Uncertain. Fuzziness Knowl. Based Syst.* **10**(05), 557–570 (2002)
34. Wimmer, H., Powell, L.: A comparison of the effects of k-anonymity on machine learning algorithms. In: *Proceedings of the Conference for Information Systems Applied Research* ISSN, vol. 2167, p. 1508 (2014)
35. Zama: Concrete-ml documentation: ML model implementation examples. https://github.com/zama-ai/concrete-ml/blob/main/docs/tutorials/ml_examples.md
36. Zama: Concrete framework security documentation (2025). <https://docs.zama.ai/concrete/explanations/security>. Technical documentation for TFHE-based homomorphic encryption security

Open Access This chapter is licensed under the terms of the Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License (<http://creativecommons.org/licenses/by-nc-nd/4.0/>), which permits any noncommercial use, sharing, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license and indicate if you modified the licensed material. You do not have permission under this license to share adapted material derived from this chapter or parts of it.

The images or other third party material in this chapter are included in the chapter's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the chapter's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

