
A Reliable Cryptographic Framework for Empirical Machine Unlearning Evaluation

Yiwen Tu*

University of Michigan, Ann Arbor
evantu@umich.edu

Pingbang Hu*

University of Illinois Urbana-Champaign
pbb@illinois.edu

Jiaqi W. Ma

University of Illinois Urbana-Champaign
jiaqima@illinois.edu

Abstract

Machine unlearning updates machine learning models to remove information from specific training samples, complying with data protection regulations that allow individuals to request the removal of their personal data. Despite the recent development of numerous unlearning algorithms, reliable evaluation of these algorithms remains an open research question. In this work, we focus on membership inference attack (MIA) based evaluation, one of the most common approaches for evaluating unlearning algorithms, and address various pitfalls of existing evaluation metrics lacking theoretical understanding and reliability. Specifically, by modeling the proposed evaluation process as a *cryptographic game* between unlearning algorithms and MIA adversaries, the naturally induced evaluation metric measures the data removal efficacy of unlearning algorithms and enjoys provable guarantees that existing evaluation metrics fail to satisfy. Furthermore, we propose a practical and efficient approximation of the induced evaluation metric and demonstrate its effectiveness through both theoretical analysis and empirical experiments. Overall, this work presents a novel and reliable approach to empirically evaluating unlearning algorithms, paving the way for the development of more effective unlearning techniques.

1 Introduction

Machine unlearning is an emerging research field in artificial intelligence (AI) motivated by the “Right to be Forgotten,” outlined by various data protection regulations such as the General Data Protection Regulation (GDPR) [Mantelero, 2013] and the California Consumer Privacy Act (CCPA) [CCPA, 2018]. Specifically, the Right to be Forgotten grants individuals the right to request that an organization erase their personal data from its databases, subject to certain exceptions. Consequently, when such data were used for training machine learning models, the organization may be required to update their models to “unlearn” the data to comply with the Right to be Forgotten. A naive solution is retraining the model on the remaining data after removing the requested data points, but this solution is computationally prohibitive. Recently, a plethora of unlearning algorithms have been developed to efficiently update the model without complete retraining, albeit usually at the price of removing the requested data information only approximately [Cao and Yang, 2015, Bourtole et al., 2021, Guo et al., 2020, Neel et al., 2021, Sekhari et al., 2021, Chien et al., 2023, Kurmanji et al., 2023].

*Equal contribution.

Despite the active development of unlearning algorithms, the fundamental problem of properly evaluating these methods remains an open research question, as highlighted by the Machine Unlearning Competition held at NeurIPS 2023². The unlearning literature has developed a variety of evaluation metrics for measuring the *data removal efficacy* of unlearning algorithms, i.e., to which extent the information of the requested data points are removed from the unlearned model. Existing metrics can be roughly categorized as attack-based [Graves et al., 2020, Kurmanji et al., 2023, Goel et al., 2023, Hayes et al., 2024, Sommer et al., 2020, Goel et al., 2023], theory-based [Triantafillou and Kairouz, 2023, Becker and Liebig, 2022], and retraining-based [Golatkhar et al., 2021, Wu et al., 2020, Izzo et al., 2021], respectively. Each metric has its own limitations and there is no consensus on a standard evaluation metric for unlearning. Among these metrics, the membership inference attack (MIA) based metric, which aims to determine whether specific data points were part of the original training dataset based on the unlearned model, is perhaps the most commonly seen in the literature. MIA is often considered a natural unlearning evaluation metric as it directly measures the privacy leakage of the unlearned model, which is a primary concern of unlearning algorithms.

Most existing literature directly uses MIA performance³ to measure the data removal efficacy of unlearning algorithms. However, such metrics can be unreliable as MIA performance is not a well-calibrated metric when used for unlearning evaluation, leading to counterintuitive results. For example, naively retraining the model is theoretically optimal for data removal efficacy, albeit computationally prohibitive. Nevertheless, retraining is not guaranteed to yield the lowest MIA performance compared to other approximate unlearning algorithms. This discrepancy arises because MIAs themselves are imperfect and can make mistakes in inferring data membership. Furthermore, MIA performance is also sensitive to the composition of data used to conduct MIA and the specific choice of MIA algorithm. Consequently, the results obtained using different MIAs are not directly comparable and can vary significantly, making it difficult to draw definitive conclusions about the efficacy of unlearning algorithms. These limitations render the existing MIA-based evaluation brittle and highlight the need for a more reliable and comprehensive framework for assessing the performance of unlearning algorithms.

In this work, we aim to address the challenges associated with MIA-based unlearning evaluation by introducing a game-theoretical framework named the *unlearning sample inference game*. Within this framework, we gauge the data removal efficacy through a game where, informally, the challenger (model provider) endeavors to produce an unlearned model, while the adversary (MIA adversary) seeks to exploit the unlearned model to determine the membership status of the given samples. By carefully formalizing the game, with controlled knowledge and interaction between both parties, we ensure that the success rate of the adversary in the unlearning sample inference game possesses several desirable properties, and thus can be used as an unlearning evaluation metric, circumventing the aforementioned pitfalls of MIA performance. Specifically, it ensures that the adversary's success rate towards the retrained model is precisely zero, thereby certifying retraining as the theoretically optimal unlearning method. Moreover, it provides a provable guarantee for certified machine unlearning algorithms [Guo et al., 2020], aligning the proposed metric with theoretical results in the literature. Lastly, it inherently accommodates the existence of multiple MIA adversaries, resolving the conflict between different choices of MIAs. However, the computational demands of exactly calculating the proposed metric pose a practical issue. To mitigate this, we introduce a *SWAP* test as a practical approximation, which also inherits many of the desirable properties of the exact metric. Empirically, this test proves robust to changes in random seed and dataset size, enabling model maintainers to conduct small-scale experiments to gauge the quality of their unlearning algorithms.

Finally, we highlight our contributions in this work as follows:

- We present a formalization of the *unlearning sample inference game*, establishing a novel unlearning evaluation metric for data removal efficacy.
- We demonstrate several provable properties of the proposed metric, circumventing various pitfalls of existing MIA-based metrics.
- We introduce a straightforward and effective *SWAP* test for efficient empirical analysis. Through thorough theoretical examination and empirical experiments, we show that it exhibits similar desirable properties.

²See <https://unlearning-challenge.github.io/>.

³For example, the accuracy or the area under the receiver operating characteristic curve (AUC) of the inferred membership.

In summary, this work offers a game-theoretic framework for reliable empirical evaluation of machine unlearning algorithms, tackling one of the most foundational problems in this field.

2 Related work

2.1 Machine unlearning

Machine unlearning, as initially introduced by Cao and Yang [2015], is to update machine learning models to remove the influence of selected training data samples, effectively making the models “forget” those samples. Most unlearning methods can be categorized as exact unlearning and approximate unlearning. Exact unlearning requires the unlearned models to be indistinguishable from models that were trained from scratch without the removed data samples. However, it can still be computationally expensive, especially for large datasets and complex models. On the other hand, approximate unlearning aims to remove the influence of selected data samples while accepting a certain level of deviation from the exactly unlearned model. This allows for more efficient unlearning algorithms, making approximate unlearning increasingly popular practically. While approximate unlearning is more time and space-efficient, it does not guarantee the complete removal of the influence of the removed data samples. We refer the audience to the survey on unlearning methods by Xu et al. [2023] for a more comprehensive overview.

2.2 Machine unlearning evaluation

Evaluating machine unlearning involves considerations of computational efficiency, model utility, and data removal efficacy. Computational efficiency refers to the time and space complexity of the unlearning algorithms, while model utility measures the prediction performance of the unlearned models. These two aspects can be measured relatively straightforwardly through, e.g., computation time, memory usage, or prediction accuracy. Data removal efficacy, on the other hand, assesses the extent to which the influence of the requested data points has been removed from the unlearned models, which is highly non-trivial to measure and has attracted significant research efforts recently. These efforts for evaluating or guaranteeing data removal efficacy can be categorized into several groups. We provide an overview below and refer readers to Appendix A for an in-depth review.

- **Retraining-based:** Generally, retraining-based evaluation measures the parameter or posterior difference between unlearned models and retrained models, the gold standard for data removal [Golatkhar et al., 2020, 2021, He et al., 2021, Izzo et al., 2021, Peste et al., 2021, Wu et al., 2020]. However, they are often unreliable as measures like parameter difference can be sensitive to the randomness led by the training dynamics [Cretu et al., 2023].
- **Theory-based:** Another line of work tries to characterize data removal efficacy by requiring a strict theoretical guarantee for the unlearned models [Chien et al., 2023, Guo et al., 2020, Neel et al., 2021] or turning to information-theoretic analysis [Becker and Liebig, 2022]. However, they have strong model assumptions or require inefficient white-box access to target models, thus limiting their applicability in practice.
- **Attack-based:** Since attacks are the most direct way to interpret privacy risks, attack-based evaluation is a common metric in unlearning literature [Chen et al., 2021, Goel et al., 2023, Graves et al., 2020, Hayes et al., 2024, Kurmanji et al., 2023, Sommer et al., 2020, Song and Mittal, 2021]. Our work belongs to this category and addresses the pitfalls of existing attack-based methods.

3 Proposed evaluation framework

3.1 Preliminaries

To mitigate the limitations of directly using MIA accuracy as the evaluation metric for unlearning algorithms, we draw inspiration from *cryptographic games* [Katz and Lindell, 2007], which are a fundamental tool to define and analyze the security properties of cryptographic protocols. In particular, we leverage the notion of *advantage* to form a more reliable and well-calibrated metric for evaluating the effectiveness of unlearning algorithms. We leave a more detailed description to Appendix B.1.

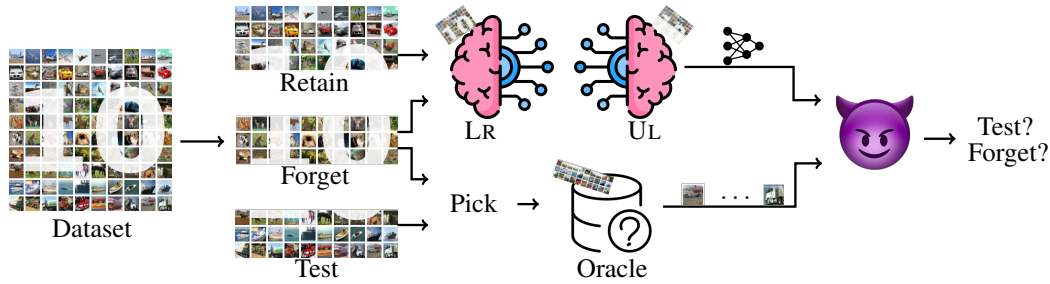


Figure 1: The unlearning sample inference game framework for our machine unlearning evaluation.

In the rest of this section, we introduce the proposed unlearning evaluation framework based on a carefully designed cryptographic game and an advantage metric associated with the game. We also provide provable guarantees for the soundness of the proposed metric.

3.2 Unlearning sample inference game

We propose the *unlearning sample inference game* $\mathcal{G} = (\text{UL}, \mathcal{A}, \mathcal{D}, \mathbb{P}_{\mathcal{D}}, \alpha)$ that characterizes the privacy risk of the unlearned models against an MIA adversary. It involves two players (a *challenger* named UL and an *adversary* $\mathcal{A}$), a finite dataset $\mathcal{D}$ with a *sensitivity distribution* $\mathbb{P}_{\mathcal{D}}$ defined over $\mathcal{D}$, and an *unlearning portion* parameter α . Intuitively, the game works as follows:

- the challenger UL performs unlearning on a “forget set” of data for a model trained on the union of the “retain set” and the “forget set,” with sizes of two subsets of $\mathcal{D}$ subject to a ratio α ;
- the adversary $\mathcal{A}$ attacks the challenger’s unlearned model by telling whether some random data points (according to $\mathbb{P}_{\mathcal{D}}$) are originally in the “forget set” or an unused set of data called “test set.”

An illustration is given in Figure 1. A detailed discussion of various design choices we made can be found in Appendix B.2. Below, we formally introduce the game, starting with the initialization phase.

Initialization. The game starts by randomly splitting the dataset $\mathcal{D}$ into three disjoint sets: a *retain set* $\mathcal{R}$, a *forget set* $\mathcal{F}$, and a *test set* $\mathcal{T}$, i.e., $\mathcal{D} = \mathcal{R} \cup \mathcal{F} \cup \mathcal{T}$, subject to the following restrictions:

- $\alpha = |\mathcal{F}|/|\mathcal{R} \cup \mathcal{F}|$: The *unlearning portion* α specifies how much data needs to be unlearned with respect to the original dataset used by the model.
- $|\mathcal{F}| = |\mathcal{T}|$: The sizes of $\mathcal{F}$ and $\mathcal{T}$ are equal to avoid potential inductive biases.

Under restrictions (a) and (b), the size of $\mathcal{R}$, $\mathcal{F}$, and $\mathcal{T}$ are determined, depending on α . We denote $\mathcal{S}_{\alpha}$ as the finite collection of all possible dataset splits satisfying restriction (a) and (b) such that $s \sim \mathcal{U}(\mathcal{S}_{\alpha})^4$ is in the form of $s = (\mathcal{R}, \mathcal{F}, \mathcal{T})$, where the tuple is ordered by the retain, forget, and test set. After splitting $\mathcal{D}$ according to s , a *random oracle* $\mathcal{O}_s(b)$ is then constructed according to s and the sensitivity distribution $\mathbb{P}_{\mathcal{D}}$, together with a secret bit $b \in \{0, 1\}$. The intuition of this random oracle is that it offers the “two scenarios” we mentioned in Section 3.1, respectively specified by $b = 0$ and $b = 1$: when the oracle $\mathcal{O}_s(b)$ is called, it emits a data point $x \sim \mathcal{O}_s(b)$ sampled from either $\mathcal{F}$ (when $b = 0$) or $\mathcal{T}$ (when $b = 1$), where the sampling probability is respect to $\mathbb{P}_{\mathcal{D}}$.

We make some remarks on the role of the sensitivity distribution $\mathbb{P}_{\mathcal{D}}$, as it seems opaque at first glance. Intuitively, $\mathbb{P}_{\mathcal{D}}$ captures biases stemming from various origins, such that more sensitive data will have greater sampling probability, hence greater privacy risks. For instance, if the forget set comprises data that users request to delete, with some being more sensitive than others, a corresponding bias should be incorporated into the game. In particular, we tailor our random oracle $\mathcal{O}_s(b)$ to sample data according to $\mathbb{P}_{\mathcal{D}}$, so when the adversary engages with the oracle, it gains increased exposure to more sensitive data, compelling the challenger to unlearn such data more effectively, thereby necessitating a heightened level of defense.

Challenger Phase. The challenger is given the retain set $\mathcal{R}$, the forget set $\mathcal{F}$, and a learning algorithm LR (takes a dataset and outputs a learned model) and a unlearning algorithm UL (takes the

⁴Throughout the paper, $\mathcal{U}(\cdot)$ denotes the uniform distribution.

original model and a training subset to unlearn, and outputs an unlearned model). For simplicity, we denote the challenger as UL, as the unlearning algorithm is the component under evaluation. The goal of the challenger is to unlearn $\mathcal{F}$ by UL from the model trained with LR on $\mathcal{R} \cup \mathcal{F}$. Intuitively, for an ideal UL, for any $x \in \mathcal{F} \cup \mathcal{T}$, it is statistically impossible to decide whether $x \in \mathcal{F}$ or $x \in \mathcal{T}$ given accesses to the unlearned model $m := \text{UL}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F})$. As both LR and UL can be randomized, m follows a distribution $\mathbb{P}_{\mathcal{M}}(\text{UL}, s)$ depending on the split s and UL, where $\mathcal{M}$ denotes the set of all possible models. This distribution $\mathbb{P}_{\mathcal{M}}(\text{UL}, s)$ summarizes the result of the challenger.

Adversary Phase. The adversary $\mathcal{A}$ is an (efficient) algorithm that has access to the unlearned model $m = \text{UL}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F})$ and the random oracle $\mathcal{O} = \mathcal{O}_s(b)$, where both s and b are unknown to $\mathcal{A}$. The goal of the adversary is to guess $b \in \{0, 1\}$ by interacting with m and $\mathcal{O}$, i.e., after interacting with $\mathcal{O}$ and m , decide whether the data points from $\mathcal{O}$ are from $\mathcal{F}$ or $\mathcal{T}$. Notation-wise, we write $\mathcal{A}^{\mathcal{O}}(m) \mapsto \{0, 1\}$. Note that in one play of the game, $\mathcal{O}$ is fixed as either $\mathcal{O}_s(0)$ or $\mathcal{O}_s(1)$ but will not switch between $b = 0$ and $b = 1$.

3.3 Advantage and unlearning quality

By viewing the unlearning sample inference game as a cryptographic game, with the discussion in Section 3.1, the corresponding *advantage* can be defined as follows:

Definition 3.1 (Advantage). *Given an unlearning sample inference game $\mathcal{G} = (\text{UL}, \mathcal{A}, \mathcal{D}, \mathbb{P}_{\mathcal{D}}, \alpha)$, the advantage of $\mathcal{A}$ against UL is defined as*

$$\text{Adv}(\mathcal{A}, \text{UL}) = \frac{1}{|\mathcal{S}_{\alpha}|} \left| \sum_{s \in \mathcal{S}_{\alpha}} \Pr_{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s)}^{\mathcal{O} = \mathcal{O}_s(0)} (\mathcal{A}^{\mathcal{O}}(m) = 1) - \sum_{s \in \mathcal{S}_{\alpha}} \Pr_{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s)}^{\mathcal{O} = \mathcal{O}_s(1)} (\mathcal{A}^{\mathcal{O}}(m) = 1) \right|.$$

To simplify notation, we sometimes omit $m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s)$ and substitute $\mathcal{O}_s(b)$ to the superscript of $\mathcal{A}$ when it's clear, i.e., we can write $\Pr(\mathcal{A}^{\mathcal{O}_s(b)}(m) = 1)$. With the definition of advantage, measuring the quality of the challenger UL is standard by considering the worst-case guarantee:

Definition 3.2 (Unlearning Quality). *For any unlearning algorithm UL, its Unlearning Quality under an unlearning sample inference game $\mathcal{G}$ is defined as*

$$\mathcal{Q}(\text{UL}) := 1 - \sup_{\mathcal{A}} \text{Adv}(\mathcal{A}, \text{UL}),$$

where the supremum is over all efficient adversary $\mathcal{A}$.

Our definition of advantage (Unlearning Quality) has several theoretical merits as detailed below.

Zero Grounding for RETRAIN. Consider UL being the gold-standard unlearning method, i.e., the retraining method RETRAIN where $\text{RETRAIN}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) = \text{LR}(\mathcal{R})$. Since the forget set $\mathcal{F}$ and the test set $\mathcal{T}$ are all unforeseen data to retrained models trained only on $\mathcal{R}$, one should expect RETRAIN to defend any adversary $\mathcal{A}$ perfectly, leading to a zero advantage. The following Theorem 3.3 shows that our definition of advantage in Definition 3.1 indeed achieves such a desirable zero grounding property.

Theorem 3.3 (Zero Grounding). *For any adversary $\mathcal{A}$, $\text{Adv}(\mathcal{A}, \text{RETRAIN}) = 0$ where RETRAIN is the retraining method. Hence, $\mathcal{Q}(\text{RETRAIN}) = 1$.*

The proof of Theorem 3.3 can be found in Appendix B.3.

At a high level, the zero grounding property of the advantage is due to its symmetry—we measure the difference between $\mathcal{O}_s(0)$ and $\mathcal{O}_s(1)$ across all possible splits in $\mathcal{S}_{\alpha}$, such that each data point has the same chance to appear in both the forget set $\mathcal{F}$ and the test set $\mathcal{T}$. In comparison to conventional MIA-based evaluation that only measures the MIA performance on a single data split, this symmetry guarantees that all MIA adversaries have a zero advantage on RETRAIN even if the MIA is biased for certain data points, as the bias will be canceled out between symmetric splits that put these data points in $\mathcal{F}$ and $\mathcal{T}$ respectively.

Guarantee Under Certified Removal. We establish an upper bound on the advantage using the well-established notion of *certified removal* [Guo et al., 2020], which is inspired by differential privacy [Dwork, 2006]:

Definition 3.4 (Certified Removal [Guo et al., 2020]; Informal). *For a fixed dataset $\mathcal{D}$, let LR and UL be learning & unlearning algorithm respectively, and denote $\mathcal{H}$ to be the hypothesis class containing all possible models that can be produced by LR and UL. Then, for any $\epsilon, \delta > 0$, the unlearning algorithm UL is said to be (ϵ, δ) -certified removal if for any $\mathcal{W} \subset \mathcal{H}$ and for **any** disjoint $\mathcal{R}, \mathcal{F} \subseteq \mathcal{D}$*

$$\begin{cases} \Pr(\text{UL}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) \in \mathcal{W}) \leq e^\epsilon \Pr(\text{RETRAIN}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) \in \mathcal{W}) + \delta; \\ \Pr(\text{RETRAIN}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) \in \mathcal{W}) \leq e^\epsilon \Pr(\text{UL}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) \in \mathcal{W}) + \delta. \end{cases}$$

Given its root in differential privacy⁵, certified removal has been widely accepted as a rigorous measure of the goodness of approximate unlearning methods [Neel et al., 2021, Chien et al., 2023], where smaller ϵ and δ indicate better unlearning. However, in practice, it is difficult to empirically quantify (ϵ, δ) for most approximate unlearning methods.

In the following Theorem 3.5, we provide a lower bound for the proposed Unlearning Quality for an (ϵ, δ) -certified removal unlearning algorithm, showing the close theoretical connection between the proposed Unlearning Quality metric and certified removal, while the proposed metric is easier to measure empirically.

Theorem 3.5 (Guarantee Under Certified Removal). *Given an (ϵ, δ) -certified removal unlearning algorithm UL with some $\epsilon, \delta > 0$, for any adversary $\mathcal{A}$ against UL, we have $\text{Adv}(\mathcal{A}, \text{UL}) \leq 2 \cdot (1 - \frac{2-2\delta}{e^\epsilon+1})$. Hence, $\mathcal{Q}(\text{UL}) \geq \frac{4-4\delta}{e^\epsilon+1} - 1$.*

The formal definition of certified removal and the proof of Theorem 3.5 can be found in Appendix B.4.

3.4 SWAP test

Direct calculation requires enumerating dataset splits in $\mathcal{S}_\alpha$, which is computationally infeasible. Hence, we propose a simple approximation scheme named the *SWAP* test, which requires as few as two dataset splits to approximate the advantage and still preserves desirable properties as the original definition. The idea is to consider the *swap pair* between a forget set $\mathcal{F}$ and a test set $\mathcal{T}$. Specifically, pick a random split $s = (\mathcal{R}, \mathcal{F}, \mathcal{T}) \in \mathcal{S}_\alpha$ and calculate the term corresponding to s in Definition 3.1:

$$\text{Adv}_s(\mathcal{A}, \text{UL}) := \Pr_{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s)}(\mathcal{A}^{\mathcal{O}_s(0)}(m) = 1) - \Pr_{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s)}(\mathcal{A}^{\mathcal{O}_s(1)}(m) = 1).$$

Next, *swap* $\mathcal{F}$ and $\mathcal{T}$ in s to get $s' = (\mathcal{R}, \mathcal{T}, \mathcal{F})$, and calculate its corresponding term in Definition 3.1:

$$\text{Adv}_{s'}(\mathcal{A}, \text{UL}) := \Pr_{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s')}(\mathcal{A}^{\mathcal{O}_{s'}(0)}(m) = 1) - \Pr_{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s')}(\mathcal{A}^{\mathcal{O}_{s'}(1)}(m) = 1).$$

Finally, average the two advantages above and obtain

$$\overline{\text{Adv}}_{\{s, s'\}}(\mathcal{A}, \text{UL}) := \frac{|\text{Adv}_s(\mathcal{A}, \text{UL}) + \text{Adv}_{s'}(\mathcal{A}, \text{UL})|}{2}.$$

In essence, we approximate Definition 3.1 by replacing $\mathcal{S}_\alpha$ with $\{s, s'\}$. Note that the *SWAP* test relies on the restriction (b) to be valid, i.e., $|\mathcal{F}| = |\mathcal{T}|$.

SWAP Test versus Random Splits. The key insight is that the *SWAP* test reserves the symmetry in the original definition of advantage, and as shown in Proposition 3.6 (see Remark B.2 for proof), it still grounds the advantage of any adversary $\mathcal{A}$ against RETRAIN to zero, preserving the same theoretical guarantees as Theorem 3.3.

Proposition 3.6 (Zero Grounding of SWAP Test (Informal)). *For any adversary $\mathcal{A}$ and swap splits $s, s' \in \mathcal{S}_\alpha$, $\overline{\text{Adv}}_{\{s, s'\}}(\mathcal{A}, \text{RETRAIN}) = 0$.*

On the contrary, naively taking two random splits with non-empty overlap can lead to an adversary with high advantage against RETRAIN:

Proposition 3.7 (High Advantage Under Random Splits). *For any two splits $s_1, s_2 \in \mathcal{S}_\alpha$ satisfying a moderate non-degeneracy assumption, there's an efficient deterministic adversary $\mathcal{A}$ such that $\overline{\text{Adv}}_{\{s_1, s_2\}}(\mathcal{A}, \text{UL}) = 1$ for **any** unlearn method UL. Particularly, $\overline{\text{Adv}}_{\{s_1, s_2\}}(\mathcal{A}, \text{RETRAIN}) = 1$.*

⁵In fact, it has been shown that a model with differential privacy guarantees automatically enjoys certified removal guarantees for any training data point [Guo et al., 2020].

The full statement and the proof of Proposition 3.7 can be found in Appendix B.5.

Remark 3.8 (Offsetting MIA Accuracy/AUC for RETRAIN). *One may wonder whether we could achieve zero grounding by simply offsetting the MIA accuracy for RETRAIN to zero (or offsetting the MIA AUC to 0.5). In Appendix B.6, we provide a discussion on why this strategy will lead to pathological cases for measuring unlearning performance.*

3.5 Practical implementation

While the proposed *SWAP* test significantly reduces the computational cost for evaluating the advantage of an adversary, evaluating the Unlearning Quality is still challenging since: 1.) most of the state-of-the-art MIAs do not exploit the covariance between data points; 2.) it is impossible to solve the supremum in Definition 3.2 exactly. We will start by addressing the first challenge.

Weak Adversary. As the current state-of-the-art MIAs make independent decisions on each data point [Bertran et al., 2023, Shokri et al., 2017, Carlini et al., 2022] without considering their covariance, therefore, for empirical analysis, we accommodate our unlearning sample inference game by restricting the adversary’s knowledge such that it can only interact with the oracle once. We call such a adversary as *weak adversary* $\mathcal{A}_w$, which will first learn a binary classifier $f(\cdot)$ by interacting with m , and output its prediction of b as $f(x)$ by querying the oracle $\mathcal{O} = \mathcal{O}_s(b)$ *exactly once* to obtain $x \sim \mathcal{O}$, where both s and b are unknown to $\mathcal{A}_w$. In this case, its advantage can be defined as

$$\text{Adv}(\mathcal{A}_w, \text{UL}) = \frac{1}{|\mathcal{S}_\alpha|} \left| \sum_{s \in \mathcal{S}_\alpha} \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ x \sim \mathcal{O}_s(0)}} (\mathcal{A}_w(m, x) = 1) - \sum_{s \in \mathcal{S}_\alpha} \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ x \sim \mathcal{O}_s(1)}} (\mathcal{A}_w(m, x) = 1) \right|$$

and the Unlearning Quality now becomes $\mathcal{Q} := 1 - \sup_{\mathcal{A}_w} \text{Adv}(\mathcal{A}_w, \text{UL})$, analogously to Definitions 3.1 and 3.2. These new definitions are subsumed under the original paradigm since the only difference is the number of interactions with the oracle.

Approximating the Supremum. While it is impossible to solve the supremum in Definition 3.2 exactly, a plausible interpretation is that the supremum is *approximately* solved by the adversary, as most of the state-of-the-art MIA adversaries are formulated as end-to-end optimization problems [Bertran et al., 2023]. By assuming these MIA adversaries are trying to maximize the advantage when constructing the final classifier $f(\cdot)$ and that the search space is large enough to parameterize all the possible weak adversaries of our interests, we can interpret that the supremum is approximately solved. Moreover, in practice, one can refine the estimation of the supremum by selecting the most potent among multiple state-of-the-art MIA adversaries.

4 Experiment

In this section, we provide empirical evidence of the effectiveness of the proposed evaluation framework. In what follows, for brevity, we will use *SWAP* test to refer to the proposed practical approximations for calculating the proposed evaluation metric, which in reality is a combination of the *SWAP* test in Section 3.4 and other approximations discussed in Section 3.5. We further denote $\mathcal{Q}$ as the proposed metric, Unlearning Quality, calculated by the *SWAP* test. With these notations established, our goal is to validate the theoretical results, demonstrate additional observed benefits of the proposed Unlearning Quality metric, and ultimately show that it outperforms other attack-based evaluation metrics. More details can be found in Appendix C. Furthermore, due to space limit, we conduct additional experiments in Appendix C.4, where we compare different unlearning algorithms’ Unlearning Quality across different dataset sizes, unlearning portion parameters, datasets, and model architectures, attacks, and also a linear setting with small privacy budgets to verify our theory.

4.1 Experiment settings

We focus on one of the most common tasks in the machine unlearning literature, image classification, and perform experiments on the CIFAR10 dataset [Krizhevsky et al., 2009], which is licensed under CC-BY 4.0. Moreover, we opt for ResNet [He et al., 2016] as the *target model* produced by some learning algorithms LR, whose details can be found in Appendix C.2. Finally, the following is the setup of the unlearning sample inference game $\mathcal{G} = (\mathcal{A}, \text{UL}, \mathcal{D}, \mathbb{P}_{\mathcal{D}}, \alpha)$ for the evaluation experiment:

- **Initialization:** Since some MIA adversaries require training the so-called *shadow models* using data sampled from the same distribution of the training data used by the target model [Shokri et al., 2017], we start by splitting the whole dataset to accommodate the training of shadow models. In particular, we split the given dataset into two halves, one for training the target model (which we call the *target dataset*), and the other for training shadow models for some MIAs. The target dataset is what we denoted as $\mathcal{D}$ in the game. To initialize the game, we consider a uniform sensitivity distribution $\mathbb{P}_{\mathcal{D}} = \mathcal{U}(\mathcal{D})$ since we do not have any prior preference for the data. The unlearning portion parameter is set to be $\alpha = 0.1$ unless specified. This implies $\mathcal{O}_s(0) = \mathcal{U}(\mathcal{F})$ and $\mathcal{O}_s(1) = \mathcal{U}(\mathcal{T})$, where $s = (\mathcal{R}, \mathcal{F}, \mathcal{T}) \in \mathcal{S}_{\alpha}$ is the split we choose to use for the game.
- **Challenger Phase:** As mentioned at the beginning of the section, we choose the learning algorithm LR which outputs ResNet as the target model. On the other hand, the corresponding unlearning algorithms UL we select for comparison are: 1.) RETRAIN: retrain from scratch (the gold-standard); 2.) FISHER: Fisher forgetting [Golatkar et al., 2020]; 3.) FTFINAL: fine-tuning final layer [Goel et al., 2023]; 4.) RETRFINAL: retrain final layer [Goel et al., 2023]; 5.) NEGGRAD: negative gradient descent [Golatkar et al., 2020]; 6.) SALUN: saliency unlearning [Fan et al., 2024]; 7.) SSD: selective synaptic dampening [Foster et al., 2024]; 8.) NONE: identity (no unlearning, dummy baseline). Among them, RETRAIN is the gold standard for exact unlearning while NONE is a dummy baseline for reference. All other methods are approximate unlearning methods, with SSD being the most recent state-of-the-art methods.
- **Adversary Phase:** Since we’re unaware of any non-weak MIAs, we focus on the following SOTA black-box (weak) MIA adversaries $\mathcal{A}$ to approximate the advantage: 1.) shadow model-based [Shokri et al., 2017]; 2.) correctness-based, confidence-based, modified entropy [Song and Mittal, 2021].

4.2 Ablation study

We first provide some empirical evidence that our approximation is reasonable and effective through two lenses: dataset size and α .

Unlearning Quality Versus Dataset Size. We vary the dataset size to assess how the unlearning quality metric $\mathcal{Q}$ scales with data. Experiments are conducted with $\alpha = 0.1$, and results are shown in Table 1. The relative ranking of unlearning methods remains consistent across sizes, indicating that small-scale evaluations can reliably approximate large-scale performance. This scalability makes $\mathcal{Q}$ an efficient and practical metric for benchmarking unlearning algorithms. The retraining method maintains $\mathcal{Q} \approx 1$, further confirming the theoretical grounding in Theorem 3.3.

Table 1: Unlearning Quality *versus* dataset size η (in percentage). The relative ranking of different unlearning methods with added standard deviations.

UL	η			
	0.1	0.4	0.8	1.0
RETRFINAL	0.340 ± 0.017	0.586 ± 0.015	0.621 ± 0.014	0.634 ± 0.025
FTFINAL	0.131 ± 0.011	0.585 ± 0.016	0.619 ± 0.014	0.634 ± 0.024
FISHER	0.751 ± 0.024	0.679 ± 0.005	0.734 ± 0.006	0.791 ± 0.020
NEGGRAD	0.124 ± 0.010	0.564 ± 0.018	0.603 ± 0.014	0.656 ± 0.035
SALUN	0.476 ± 0.014	0.617 ± 0.016	0.689 ± 0.013	0.748 ± 0.004
SSD	0.975 ± 0.008	0.939 ± 0.025	0.929 ± 0.021	0.928 ± 0.015
RETRAIN	0.999 ± 0.000	0.997 ± 0.001	0.993 ± 0.001	0.993 ± 0.001

Unlearning Quality Versus α . We further vary the unlearning portion parameter α , defined as the ratio of the forget set to the full training set, to investigate how the unlearning quality metric $\mathcal{Q}$ changes. The results are presented in Table 2. We observe that the relative ranking of different unlearning methods remains largely consistent across varying α values. This stability suggests that our metric reliably captures the comparative performance of unlearning algorithms under different data removal proportions.

Table 2: Unlearning Quality *versus* α . The relative ranking of different unlearning methods stays mostly consistent across different α .

UL	α			
	0.1	0.25	0.4	0.67
RETRFINAL	0.513 ± 0.054	0.489 ± 0.055	0.413 ± 0.027	0.500 ± 0.024
FTFINAL	0.511 ± 0.054	0.484 ± 0.054	0.394 ± 0.041	0.479 ± 0.020
FISHER	0.904 ± 0.046	0.871 ± 0.044	0.908 ± 0.038	0.810 ± 0.050
NEGGRAD	0.637 ± 0.105	0.757 ± 0.073	0.684 ± 0.016	0.631 ± 0.088
SALUN	0.755 ± 0.017	0.762 ± 0.043	0.895 ± 0.047	0.893 ± 0.028
SSD	0.944 ± 0.038	0.960 ± 0.019	0.837 ± 0.054	0.913 ± 0.037

4.3 Validation of theoretical results

We empirically validate Theorems 3.3 and 3.5. While it is easy to verify **grounding**, i.e., $\mathcal{Q}(\text{RETRAIN}) = 1$, validating the lower-bound of $\mathcal{Q}$ for unlearning algorithms with (ϵ, δ) -certified removal guarantees is challenging since such algorithms are not known beyond convex models. However, if the model is trained with (ϵ, δ) -differential privacy (DP) guarantees, then even if we do not apply any unlearning on the model (i.e., UL = NONE), the model still automatically satisfies (ϵ, δ) -certified removal for any training data point [Guo et al., 2020]. As the DP algorithm exists for non-convex models [Abadi et al., 2016], this suggests that one can analyze the impact of the *DP privacy budget* on $\mathcal{Q}$ of an (ϵ, δ) -DP model. In particular, we fix $\delta = 10^{-5}$ and consider varying ϵ to be 50, 150, 600, and ∞ (∞ corresponds to no DP training). The corresponding Unlearning Quality results are reported in Table 3.

Firstly, as can be seen from the last row of Table 3, $\mathcal{Q}(\text{RETRAIN}) \approx 1$ with high precision for all ϵ , achieving **grounding** almost perfectly, thus validating Theorem 3.3. Furthermore, Theorem 3.5 suggests that the lower bound of $\mathcal{Q}$ should **negatively correlate** with ϵ . Indeed, empirically, we observe such a trend with high precision, again validating our theoretical findings. Moreover, Table 3 shows that the Unlearning Quality also maintains a **consistent** relative ranking between unlearning algorithms among different ϵ , proving its robustness. Finally, our metric suggests that SSD significantly outperforms other unlearning methods, which is consistent with the fact that it is the most recent state-of-the-art method among the unlearning methods we evaluate.

Remark 4.1. From Table 3, even for NONE with $\epsilon = \infty$, the Unlearning Quality $\mathcal{Q}$ is still relatively high (0.587). This is partly because the current state-of-the-art (weak) MIA adversaries are not good enough: if the weak adversary becomes better, our evaluation metric can also benefit from this.

Remark 4.2. We distinguish between A) the computation cost of the proposed unlearning evaluation metric and B) the cost of evaluating the metric itself. The latter is significantly higher due to DP training and is reported in the appendix. In real applications, only the scalable cost of A) is incurred.

4.4 Comparison to other metrics

We compare our Unlearning Quality metric $\mathcal{Q}$ to other existing attack-based evaluation metrics, demonstrating the superiority of the proposed metric. We limit our evaluation to the MIA-based evaluation, and within this category, three MIA-based evaluation metrics are most relevant to our setting [Triantafillou and Kairouz, 2023, Golatkar et al., 2021, Goel et al., 2023]. While none of them enjoy the **grounding** property, in particular, the one proposed by Triantafillou and Kairouz [2023] requires training attacks for every forget data sample, which is extremely time-consuming, and we leave it out from the comparison and focus on comparing our metric with the other two.

The two metrics we will compare are the pure MIA AUC (Area Under Curve) [Golatkar et al., 2021] and the Interclass Confusion (IC) Test [Goel et al., 2023]. The former is straightforward but falls

Table 3: Unlearning Quality *versus* DP budgets. We use † to indicate that the results' standard error of the mean is < 0.01 and use * for < 0.005 .

UL	ϵ			
	50	150	600	∞
NONE	$0.972^\dagger$	$0.960^\dagger$	0.932^*	$0.587^\dagger$
NEGGRAD	0.980^*	$0.975^\dagger$	$0.953^\dagger$	$0.628^\dagger$
RETRFINAL	$0.972^\dagger$	$0.964^\dagger$	$0.939^\dagger$	$0.576^\dagger$
FTFINAL	0.973^*	$0.963^\dagger$	$0.939^\dagger$	$0.574^\dagger$
FISHER	0.973^*	$0.967^\dagger$	0.942^*	$0.709^\dagger$
SALUN	0.979^*	$0.972^\dagger$	0.945^*	0.689^*
SSD	0.996^*	0.988^*	$0.981^\dagger$	$0.888^\dagger$
RETRAIN	0.998^*	0.996^*	0.997^*	0.993^*

Table 4: Comparison between IC score and MIA score under different DP budgets. See Table 3 for more context.

(A) IC score <i>versus</i> DP budgets.					(B) MIA score <i>versus</i> DP budgets.				
UL	ϵ				UL	ϵ			
	50	150	600	∞		50	150	600	∞
NONE	0.749 [†]	0.720 [†]	0.717 [†]	0.005 [*]	NONE	0.451 [†]	0.433 [†]	0.454 [†]	0.380 [†]
NEGGRAD	0.925 [*]	0.953 [†]	0.946 [†]	0.180 [†]	NEGGRAD	0.476 [†]	0.482 [†]	0.466 [†]	0.299 [†]
RETRFINAL	0.931 [*]	0.958 [†]	0.893 [†]	0.033 [†]	RETRFINAL	0.485 [†]	0.485 [†]	0.472 [†]	0.248 [†]
FTFINAL	0.930 [*]	0.957 [†]	0.893 [†]	0.346 [†]	FTFINAL	0.485 [†]	0.485 [†]	0.472 [†]	0.247 [†]
FISHER	0.743 [*]	0.720 [†]	0.702 [†]	0.344 [†]	FISHER	0.475 [†]	0.484 [†]	0.463 [†]	0.325 [†]
SALUN	0.866 [*]	0.887 [*]	0.841 [†]	0.081 [*]	SALUN	0.488 [*]	0.491 [*]	0.477 [†]	0.268 [*]
SSD	0.976 [*]	1.000 [*]	0.990 [†]	0.174 [†]	SSD	0.480 [*]	0.480 [*]	0.468 [†]	0.244 [*]
RETRAIN	0.962 [*]	0.972 [†]	0.976 [†]	0.975 [†]	RETRAIN	0.479 [*]	0.491 [*]	0.492 [*]	0.488 [*]

short in many aspects as discussed in the introduction; the latter, on the other hand, is a more refined metric. In brief, the IC test “confuses” a selected set S of two classes by switching their labels and training a model on the modified dataset. It then requests the unlearning algorithm to unlearn S and measures the inter-class error $\gamma \in [0, 1]$ of the unlearned model on S . Similar to the advantage, γ is the “flipped side” of the unlearning performance, which suggests defining the *IC score* by $1 - \gamma$. Similarly, for the sake of clear comparison, we report the *MIA score* defined as “ $1 - \text{MIA AUC}$ ” as the unlearning performance, where the MIA AUC is calculated on the union of the forget set and test set. We leave the details to Appendix C.3.

We conduct the comparison experiment by again analyzing the relation between the *DP privacy budget* versus the *evaluation result* of an (ϵ, δ) -DP model for the two metrics we are comparing with under the same setup as in Table 3. Specifically, we let $\delta = 10^{-5}$ and consider varying ϵ from 50 to ∞ , and we look into two aspects: 1) **negative correlation** with ϵ ; 2) **consistency** w.r.t. ϵ . The results are shown in Table 4. Firstly, we see that according to Table 4(A), the IC test **fails** to produce a **negatively correlated** evaluation result with ϵ . For instance, the IC score for NEGGRAD is notably lower at $\epsilon = 50$ than at $\epsilon = 150$, and RETRFINAL and FTFINAL also demonstrate a higher IC score at $\epsilon = 150$ than at $\epsilon = 600$. Furthermore, in terms of **consistency** w.r.t. ϵ , we again see that the IC test **fails** to satisfy this property, unlike the proposed Unlearning Quality metric $\mathcal{Q}$. For example, while NONE is better than NEGGRAD at $\epsilon = 50$, this relative ranking is not maintained at $\epsilon = \infty$. Such an inconsistency happens multiple times across Table 4(A). A similar story can be told for the MIA AUC, where from Table 4(B), we see that MIA AUC also **fails** to produce a similar trend as $\mathcal{Q}$ where the evaluation results are **negatively correlated** with ϵ . For instance, the MIA scores for NEGGRAD and FISHER are notably higher at $\epsilon = 150$ than at $\epsilon = 600$. Furthermore, in terms of **consistency** w.r.t. ϵ , we see that while NEGGRAD outperforms FISHER at $\epsilon = 50$ and $\epsilon = 600$, it performs worse than FISHER at $\epsilon = 150$, which is **inconsistent**. Overall, both the IC test and the MIA AUC fail to satisfy the properties that we have established for the Unlearning Quality, demonstrating our superiority.

5 Conclusion

In this work, we developed a game-theoretical framework named the *unlearning sample inference game* and proposed a novel metric for evaluating the data removal efficacy of approximate unlearning methods. Our approach is rooted in the concept of “advantage,” borrowed from cryptography, to quantify the success of an MIA adversary in differentiating forget data from test data given an unlearned model. This metric enjoys zero grounding for the theoretically optimal retraining method, scales with the privacy budget of certified unlearning methods, and can take advantage (as opposed to suffering from conflicts) of various MIA methods, which are desirable properties that existing MIA-based evaluation metrics fail to satisfy. We also propose a practical tool — the *SWAP* test — to efficiently approximate the proposed metric. Our empirical findings reveal that the proposed metric effectively captures the nuances of machine unlearning, demonstrating its robustness across varying dataset sizes and its adaptability to the constraints of differential privacy budgets. The ability to maintain a discernible difference and a partial order among unlearning methods, regardless of dataset size, highlights the practical utility of our approach. By bridging theoretical concepts with empirical analysis, our work lays a solid foundation for reliable empirical evaluation of machine unlearning and paves the way for the development of more effective unlearning algorithms.

References

- Martin Abadi, Andy Chu, Ian Goodfellow, H. Brendan McMahan, Ilya Mironov, Kunal Talwar, and Li Zhang. Deep learning with differential privacy. In *Proceedings of the 2016 ACM SIGSAC Conference on Computer and Communications Security*, CCS'16. ACM, October 2016. doi: 10.1145/2976749.2978318.
- Alexander Becker and Thomas Liebig. Evaluating machine unlearning via epistemic uncertainty, 2022.
- Martin Andres Bertran, Shuai Tang, Aaron Roth, Michael Kearns, Jamie Heather Morgenstern, and Steven Wu. Scalable membership inference attacks via quantile regression. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- Lucas Bourtole, Varun Chandrasekaran, Christopher A Choquette-Choo, Hengrui Jia, Adelin Travers, Baiwu Zhang, David Lie, and Nicolas Papernot. Machine unlearning. In *2021 IEEE Symposium on Security and Privacy (SP)*, pages 141–159. IEEE, 2021.
- Brennon Brimhall, Philip Mathew, Neil Fendley, Yinzhi Cao, and Matthew Green. Mirror mirror on the wall, have i forgotten it all? a new framework for evaluating machine unlearning, 2025. URL <https://arxiv.org/abs/2505.08138>.
- Yinzhi Cao and Junfeng Yang. Towards making systems forget with machine unlearning. In *2015 IEEE symposium on security and privacy*, pages 463–480. IEEE, 2015.
- Nicholas Carlini, Steve Chien, Milad Nasr, Shuang Song, Andreas Terzis, and Florian Tramer. Membership inference attacks from first principles. In *2022 IEEE Symposium on Security and Privacy (SP)*, pages 1897–1914. IEEE, 2022.
- CCPA. California consumer privacy act (ccpa), 2018.
- Min Chen, Zhikun Zhang, Tianhao Wang, Michael Backes, Mathias Humbert, and Yang Zhang. When machine unlearning jeopardizes privacy. In *Proceedings of the 2021 ACM SIGSAC conference on computer and communications security*, pages 896–911, 2021.
- Min Chen, Zhikun Zhang, Tianhao Wang, Michael Backes, Mathias Humbert, and Yang Zhang. Graph unlearning. In *Proceedings of the 2022 ACM SIGSAC Conference on Computer and Communications Security*, CCS '22. ACM, November 2022. doi: 10.1145/3548606.3559352.
- Eli Chien, Chao Pan, and Olgica Milenkovic. Efficient model updates for approximate unlearning of graph-structured data. In *International Conference on Learning Representations*, 2023.
- Rishav Chourasia and Neil Shah. Forget unlearning: Towards true data-deletion in machine learning. In *International conference on machine learning*, pages 6028–6073. PMLR, 2023.
- Ana-Maria Cretu, Daniel Jones, Yves-Alexandre de Montjoye, and Shruti Tople. Re-aligning shadow models can improve white-box membership inference attacks. *arXiv preprint arXiv:2306.05093*, 2023.
- Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 conference of the North American chapter of the association for computational linguistics: human language technologies, volume 1 (long and short papers)*, pages 4171–4186, 2019.
- Cynthia Dwork. Differential privacy. In *International colloquium on automata, languages, and programming*, pages 1–12. Springer, 2006.
- Chongyu Fan, Jiancheng Liu, Yihua Zhang, Eric Wong, Dennis Wei, and Sijia Liu. Salun: Empowering machine unlearning via gradient-based weight saliency in both image classification and generation. In *The Twelfth International Conference on Learning Representations*, 2024.
- Jack Foster, Stefan Schoepf, and Alexandra Brintrup. Fast machine unlearning without retraining through selective synaptic dampening. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 12043–12051, 2024.

- Shashwat Goel, Ameya Prabhu, Amartya Sanyal, Ser-Nam Lim, Philip Torr, and Ponnurangam Kumaraguru. Towards adversarial evaluations for inexact machine unlearning, 2023.
- A. Golatkar, A. Achille, and S. Soatto. Eternal sunshine of the spotless net: Selective forgetting in deep networks. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 9301–9309, Los Alamitos, CA, USA, jun 2020. IEEE Computer Society. doi: 10.1109/CVPR42600.2020.00932.
- Aditya Golatkar, Alessandro Achille, Avinash Ravichandran, Marzia Polito, and Stefano Soatto. Mixed-privacy forgetting in deep networks. In *2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 792–801, 2021. doi: 10.1109/CVPR46437.2021.00085.
- Laura Graves, Vineel Nagisetty, and Vijay Ganesh. Amnesiac machine learning. In *AAAI Conference on Artificial Intelligence*, 2020.
- Chuan Guo, Tom Goldstein, Awni Hannun, and Laurens Van Der Maaten. Certified data removal from machine learning models. In *International Conference on Machine Learning*, pages 3832–3842. PMLR, 2020.
- Jamie Hayes, Ilia Shumailov, Eleni Triantafillou, Amr Khalifa, and Nicolas Papernot. Inexact unlearning needs more careful evaluations to avoid a false sense of privacy. *arXiv preprint arXiv:2403.01218*, 2024.
- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- Yingzhe He, Guozhu Meng, Kai Chen, Jinwen He, and Xingbo Hu. Deepoblivate: a powerful charm for erasing data residual memory in deep neural networks. *arXiv preprint arXiv:2105.06209*, 2021.
- Zachary Izzo, Mary Anne Smart, Kamalika Chaudhuri, and James Zou. Approximate data deletion from machine learning models. In *International Conference on Artificial Intelligence and Statistics*, pages 2008–2016. PMLR, 2021.
- Jonathan Katz and Yehuda Lindell. *Introduction to modern cryptography: principles and protocols*. Chapman and hall/CRC, 2007.
- Alex Krizhevsky, Geoffrey Hinton, et al. Learning multiple layers of features from tiny images. *University of Toronto*, 2009.
- Meghdad Kurmanji, Peter Triantafillou, Jamie Hayes, and Eleni Triantafillou. Towards unbounded machine unlearning. In A. Oh, T. Neumann, A. Globerson, K. Saenko, M. Hardt, and S. Levine, editors, *Advances in Neural Information Processing Systems*, volume 36, pages 1957–1987. Curran Associates, Inc., 2023.
- Yann LeCun. The mnist database of handwritten digits. <http://yann.lecun.com/exdb/mnist/>, 1998.
- Alessandro Mantelero. The eu proposal for a general data protection regulation and the roots of the ‘right to be forgotten’. *Computer Law & Security Review*, 29(3):229–235, 2013. ISSN 0267-3649. doi: <https://doi.org/10.1016/j.clsr.2013.03.010>.
- Seth Neel, Aaron Roth, and Saeed Sharifi-Malvajerdi. Descent-to-delete: Gradient-based methods for machine unlearning. In *Algorithmic Learning Theory*, pages 931–962. PMLR, 2021.
- Alexandra Peste, Dan Alistarh, and Christoph H. Lampert. Ssse: Efficiently erasing samples from trained machine learning models, 2021.
- Youyang Qu, Xin Yuan, Ming Ding, Wei Ni, Thierry Rakotoarivelo, and David Smith. Learn to unlearn: Insights into machine unlearning. *Computer*, 57(3):79–90, 2024.
- Sebastian Ruder. An overview of gradient descent optimization algorithms. *arXiv preprint arXiv:1609.04747*, 2016.

- Ayush Sekhari, Jayadev Acharya, Gautam Kamath, and Ananda Theertha Suresh. Remember what you want to forget: Algorithms for machine unlearning. *Advances in Neural Information Processing Systems*, 34:18075–18086, 2021.
- Reza Shokri, Marco Stronati, Congzheng Song, and Vitaly Shmatikov. Membership inference attacks against machine learning models. In *2017 IEEE Symposium on Security and Privacy (SP)*, pages 3–18, 2017. doi: 10.1109/SP.2017.41.
- Richard Socher, Alex Perelygin, Jean Wu, Jason Chuang, Christopher D. Manning, Andrew Ng, and Christopher Potts. Recursive deep models for semantic compositionality over a sentiment treebank. In David Yarowsky, Timothy Baldwin, Anna Korhonen, Karen Livescu, and Steven Bethard, editors, *Proceedings of the 2013 Conference on Empirical Methods in Natural Language Processing*, pages 1631–1642, Seattle, Washington, USA, October 2013. Association for Computational Linguistics. URL <https://aclanthology.org/D13-1170/>.
- David Marco Sommer, Liwei Song, Sameer Wagh, and Prateek Mittal. Towards probabilistic verification of machine unlearning. *CoRR*, abs/2003.04247, 2020.
- Liwei Song and Prateek Mittal. Systematic evaluation of privacy risks of machine learning models. In *30th USENIX Security Symposium (USENIX Security 21)*, pages 2615–2632, 2021.
- Eleni Triantafillou and Peter Kairouz. Evaluation for the NeurIPS Machine Unlearning Competition, 2023. [Accessed 10-01-2024].
- Yinjun Wu, Edgar Dobriban, and Susan Davidson. Deltagrad: Rapid retraining of machine learning models. In *International Conference on Machine Learning*, pages 10355–10366. PMLR, 2020.
- Heng Xu, Tianqing Zhu, Lefeng Zhang, Wanlei Zhou, and Philip S. Yu. Machine unlearning: A survey, 2023.
- Ashkan Yousefpour, Igor Shilov, Alexandre Sablayrolles, Davide Testuggine, Karthik Prasad, Mani Malek, John Nguyen, Sayan Ghosh, Akash Bharadwaj, Jessica Zhao, Graham Cormode, and Ilya Mironov. Opacus: User-friendly differential privacy library in PyTorch. *arXiv preprint arXiv:2109.12298*, 2021.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [\[Yes\]](#)

Justification: The abstract and introduction clearly define the scope of both the theoretical and empirical results.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [\[Yes\]](#)

Justification: We discuss the limitations in Section 5 and in the appendix.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [\[Yes\]](#)

Justification: Given the theoretical nature of this paper, we have diligently ensured the accuracy of the theorem statements and proofs.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: Our code is included in the supplementary materials.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
 - (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: Our code is included in the supplementary materials.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: The details of the experiments are discussed in Section 4.1 and in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [Yes]

Justification: We have closely investigated the variance of our experiments by reporting our results with multiple random dataset splits in Section 4, and also comparing the variance of our proposed swapped test with that of random dataset splits in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).

- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: The information on the compute resources is reported in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: Every author of this submission has reviewed the code of ethics guidelines and confirms compliance.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. Broader impacts

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: While our work is theoretical in nature, it still has an empirical impact as our theory is designed with practical scenarios in mind. We discuss some of the potential societal impacts in the appendix.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.

- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. Safeguards

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [\[Yes\]](#)

Justification: The paper poses no such risks.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [\[Yes\]](#)

Justification: We properly cite the datasets and include their licenses in Section 4.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.

- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. **New assets**

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [NA]

Justification: The paper does not release new assets.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. **Crowdsourcing and research with human subjects**

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. **Institutional review board (IRB) approvals or equivalent for research with human subjects**

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: The paper does not involve crowdsourcing nor research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. **Declaration of LLM usage**

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [NA]

Justification: The paper does not involve LLM.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Additional related work

Machine Unlearning. Techniques like data sharding [Bourtoule et al., 2021, Chen et al., 2022] partition the training process in such a way that only a portion of the model needs to be retrained when removing a subset of the dataset, reducing the computational burden compared to retraining the entire model. For example, Guo et al. [2020], Neel et al. [2021], Chien et al. [2023] analyzed the influence of removed data on linear or convex models and proposed gradient-based updates on model parameters to remove this influence. Chourasia and Shah [2023] proposed an unlearning method that appears similar to ours but differs in key aspects, particularly in the definition of advantage, a term whose meaning varies by threat model. Their threat model leads to a criterion for unlearning effectiveness, but like theory-based approaches such as certified removal, it is hard to empirically evaluate for most approximate methods lacking guarantees. In contrast, our threat model is tailored to enable a practical and novel evaluation metric, addressing a key gap in unlearning research.

Retraining-based Evaluation. Generally, retraining-based evaluation seeks to compare unlearned models to retrained models. As introduced in the works by Golatkar et al. [2021], He et al. [2021], Golatkar et al. [2020], model accuracy on the forget set should be similar to the accuracy on the test set as if the forget set never exists in the training set. Peste et al. [2021] proposed an evaluation metric based on the normalized confusion matrix element-wise difference on selected data samples. Golatkar et al. [2021] proposed using relearn time, which is the additional time to use for unlearned models to perform comparably to retrained models. The authors also proposed to measure the ℓ_1 distance between the final activations of the scrubbed weights and the retrained model. Wu et al. [2020], Izzo et al. [2021] turned to ℓ_2 distance of weight parameters between unlearned models and retrained models. In general, beyond the need for additional implementation and the lower computational efficiency inherent in retraining-based evaluations, a more critical issue is the influence of random factors. As discussed by Cretu et al. [2023], such random factors, including the sequence of data batches and the initial configuration of models, can lead to the unaligned storage of information within models. This misalignment may foster implicit biases favoring certain retrained models.

Theory-based Evaluation. Some literature tries to characterize data removal efficacy by requiring a strict theoretical guarantee for the unlearned models. However, these methods have strong model assumptions, such as convexity or linearity, or require inefficient white-box access to target models, thus limiting their applicability in practice. For example, Guo et al. [2020], Neel et al. [2021], Chien et al. [2023] focus on the notion of the certified removal (CR), which requires that the unlearned model cannot be statistically distinguished from the retrained model. By definition, CR is parametrized by privacy parameters called *privacy budgets*, which quantify the level of statistical indistinguishability. Hence, models with CR guarantees will intrinsically satisfy an “evaluation metric” induced by the definition of CR, acting as a form of “evaluation.” On the other hand, Becker and Liebig [2022] adopted an information-theoretical perspective and turned to epistemic uncertainty to evaluate the information remaining after unlearning. A concurrent study [Brimhall et al., 2025] that appeared later than this paper proposes a framework for evaluating machine unlearning called computational unlearning, inspired by cryptographic indistinguishability games between retrained models and unlearned models. Their framework shares similar motivations to ours but has different technical assumptions.

Attack-based Evaluation. Since attacks are the most direct way to interpret privacy risks, attack-based evaluation is a common metric in unlearning literature. The classical approach is to directly calculate the MIA accuracy using various kinds of MIAs [Graves et al., 2020, Kurmanji et al., 2023]. One kind of MIA utilizes shadow models [Shokri et al., 2017], which are trained with the same model structure as the original models but on a shadow dataset sampled with the same data sampling distribution. Moreover, some MIAs calculate membership scores based on correctness and confidence [Song and Mittal, 2021]. Some evaluation metrics do move beyond the vanilla MIA accuracy. For example, Triantafillou and Kairouz [2023] leveraged hypothesis testing coupled with MIAs to compute an estimated privacy budget for each unlearning method, which gives a rather rigorous estimation of unlearning efficacy. Hayes et al. [2024] proposed a novel MIA towards machine unlearning based on Likelihood Ratio Attack and evaluated machine unlearning through a combination of the predicted membership probability and the *balanced* MIA accuracy on test and forget sets. They designed a new MIA attack with a similar attack-defense game framework. There are other evaluation metrics also based on MIAs, but with different focuses. However, as they still use MIA accuracy as the evaluation metric, the game itself doesn’t bring much for their evaluation framework other than a clear experiment procedure. Goel et al. [2023] proposed an Interclass

Confusion (IC) test that manipulates the input dataset to evaluate both model indistinguishability and property generalization. However, their metric is less direct in terms of interpreting real-life privacy risks. Lastly, For example, Chen et al. [2021] proposed a novel metric based on MIAs that know both learned and unlearned models with a focus on how much information is deleted rather than how much information is left after the unlearning process. Sommer et al. [2020] provided a backdoor verification mechanism for Machine Learning as a Service (MLaaS), which benefits an individual user valuing his/her privacy to verify the efficacy of unlearning. They focus more on user-level verification rather than model-level evaluation.

B Omitted details from Section 3

B.1 More details of advantage

In cryptographic games, there are two interacting players: a benign player named *challenger* representing the cryptographic protocol under evaluation (corresponding to the unlearning algorithm in our context), and an *adversary* attempts to compromise the security properties of the challenger. The game proceeds in several phases, including an initialization phase where the game is initialized with specific configuration parameters, a challenger phase where the challenger performs the cryptographic protocol, and an adversary phase where the adversary queries allowed by the game's rules and generates a guess of the secret protected by the challenger. Finally, the game concludes with a win or a loss for the adversary, depending on their guess. In the context of machine unlearning, the goal of the adversary is to guess whether certain given data comes from the set to be unlearned (the forget set) or the set never used in training (the test set), based on access to the unlearned model.

The notion of advantage quantifies, in probabilistic terms, how effectively an adversary can win the game when it is played repeatedly. It is often defined as the difference in the adversary's accepting rate between two distinct scenarios (e.g., with or without access to information potentially leaked by the cryptographic protocol) [Katz and Lindell, 2007]. In the context of machine unlearning, the two scenarios can refer to the data given to the adversary coming from either the forget set or the test set, respectively. The game is constructed such that, if the cryptographic protocol is perfectly secure (i.e., the unlearned model has completely erased the information of the forget set), the adversary's advantage is expected to be zero, making it a well-calibrated measure of the protocol's security.

B.2 Design choices

In this section, we justify some of our design choices when designing the unlearning sample inference game. Most of them are of practical consideration, while some are for the convenience of analysis.

Uniform Splitting. At the initialization phase, we choose the split uniformly rather than allowing sampling from an arbitrary distribution. The reason is two-fold: Firstly, since this sampling strategy corresponds to the so-called *i.i.d. unlearning* setup [Qu et al., 2024], i.e., the unlearned samples will be drawn from a distribution of $\mathcal{D}$ in an i.i.d. fashion. In this regard, uniformly splitting the dataset corresponds to a uniform distribution of $\mathcal{D}$ for the unlearned samples to be drawn from. This is the most commonly used sampling strategy when evaluating unlearning algorithms since it's difficult to estimate the probability that data will be requested to be unlearned.

Secondly, Qu et al. [2024] acknowledged the significantly greater difficulty of non-i.i.d. unlearning compared to i.i.d. unlearning empirically. A classic example of non-i.i.d. unlearning is the process of unlearning an entire class of data, where a subset of data shares the same label. Conversely, even non-uniform splitting complicates the analysis, leading to the breakdown of our theoretical results. Specifically, generalizing both Theorem 3.3 and Theorem 3.5 becomes non-trivial. Overall, non-uniform splitting presents obstacles both empirically and theoretically.

Intrinsic Learning Algorithm for Challenger. The challenger, which we denote as UL , has a learning algorithm LR in mind in our formulation. This is because the existing theory-based unlearning method, such as the certified removal algorithm [Guo et al., 2020] as defined in Definition B.3, is achieved by a combination of the learning algorithm and a corresponding unlearning method to support unlearning request with theoretical guarantees. In other words, given an arbitrary learning algorithm LR , it's unlikely to design an efficient unlearning algorithm UL with strong theoretical guarantees, at least this is not captured by the notion of certified removal. Hence, allowing the challenger to choose its learning algorithm accommodates this situation.

Black-box Adversary v.s. White-box Adversary. By default, we assume that m is given to $\mathcal{A}$ in a *black-box* fashion, i.e., $\mathcal{A}$ only has oracle access to m . However, our framework can also adapt to white-box adversaries, which requires full model parameters of m . The only difference is that the efficiency definition changes accordingly, i.e., polynomial time in the size of $|\mathcal{D}|$ for a black-box adversary or polynomial time in the number of parameters of m for a white-box adversary.

Strong adversary in practice. As discussed in Section 3.5, the current state-of-the-art MIA adversaries are all *weak*. While it is possible to formulate the unlearning sample inference game entirely with the weak adversary, we discuss the rationale behind considering the strong adversary. One of the apparent reasons is simply that the strong adversary encompasses the weak adversary, thereby enhancing the generality of our framework and theory. However, we argue that the strong adversary is more practical in many real-world scenarios, and bringing this stronger notion has further practical impacts beyond blindly generalizing our model.

Consider a scenario where an adversary conducts a membership inference attack on a large scale. We argue that in practice, it is more reasonable to aim for a high *overall membership accuracy* of a set of carefully chosen data points, rather than the individual membership status of each of them. For example, consider the case where we are interested in images sourced from the internet. In this case, it is safe to assume that images within the same webpage are either all included in the model training dataset or none are, if we assume that the training data is collected via some reasonable data mining algorithms. In such a case, the ability of an adversary to infer the common membership for a group of data points from a particular webpage becomes desirable as it is likely to enhance the overall MIA accuracy. We believe that this stronger notion of MIA adversary has more practical impacts and reflects the common practice when deploying the membership inference attack; therefore, we choose to formulate the unlearning sample inference game with it.

B.3 Proof of Theorem 3.3

We now prove Theorem 3.3. We repeat the statement for convenience.

Theorem B.1. *For any (potentially inefficient) adversary $\mathcal{A}$, its advantage against the retraining method RETRAIN in an unlearning sample inference game $\mathcal{G} = (\text{RETRAIN}, \mathcal{A}, \mathcal{D}, \mathbb{P}_{\mathcal{D}}, \alpha)$ is zero, i.e., $\text{Adv}(\mathcal{A}, \text{RETRAIN}) = 0$.*

Proof. Firstly, we may partition the collection of all the possible dataset splits $\mathcal{S}_{\alpha}$ by fixing the retain sets $\mathcal{R} \subset \mathcal{D}$. Specifically, denote the collection of dataset splits with the retain set to be $\mathcal{R}$ as $\mathcal{S}_{\alpha}[\mathcal{R}] := \{s \in \mathcal{S}_{\alpha} : s = (\mathcal{R}, \cdot, \cdot)\}$. With the usual convention, when there's no dataset split corresponds to $\mathcal{R}$, $\mathcal{S}_{\alpha}[\mathcal{R}] = \emptyset$. Observe that for any $s \in \mathcal{S}_{\alpha}[\mathcal{R}]$, we can pair it up with another dataset split that *swaps* the forget and test sets in s . In other words, for any $s = (\mathcal{R}, \mathcal{F}, \mathcal{T}) \in \mathcal{S}_{\alpha}[\mathcal{R}]$, we see that $(\mathcal{R}, \mathcal{T}, \mathcal{F})$ is also in $\mathcal{S}_{\alpha}[\mathcal{R}]$ since we assume $|\mathcal{F}| = |\mathcal{T}|$, every dataset split will be paired. In addition, since $\mathcal{R}$ is fixed in $\mathcal{S}_{\alpha}[\mathcal{R}]$, we know that $\mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) =: \mathbb{P}_{\mathcal{M}}(\mathcal{R})$ is the same for all $s \in \mathcal{S}_{\alpha}[\mathcal{R}]$ since the unlearning algorithm UL is RETRAIN, i.e., it only depends on $\mathcal{R}$. With these observations, we can then combine the paired dataset splits within the expectation.

Specifically, for any $s \in \mathcal{S}_{\alpha}[\mathcal{R}]$, let $s' \in \mathcal{S}_{\alpha}[\mathcal{R}]$ to be s 's pair, i.e., if $s = (\mathcal{R}, \mathcal{F}, \mathcal{T})$ for some $\mathcal{F}$ and $\mathcal{T}$, then $s' = (\mathcal{R}, \mathcal{T}, \mathcal{F})$. Finally, for a given $\mathcal{R}$, let's denote the collection of all such pairs as

$$\mathcal{P}_{\mathcal{R}} := \{\{s, s'\} \subset \mathcal{S}_{\alpha}[\mathcal{R}] : s = (\mathcal{R}, \mathcal{F}, \mathcal{T}), s' = (\mathcal{R}, \mathcal{T}, \mathcal{F}) \text{ for some } \mathcal{F}, \mathcal{T}\},$$

where we use a set rather than an ordered list for the pair s, s' since we do not want to deal with repetitions. Observe that $\mathcal{O}_s(0) = \mathcal{O}_{s'}(1)$ and $\mathcal{O}_s(1) = \mathcal{O}_{s'}(0)$ since the oracles are constructed

with respect to the same preference distribution for all data splits. Hence, we have

$$\begin{aligned}
& \text{Adv}(\mathcal{A}, \text{RETRAIN}) \\
&= \frac{1}{|\mathcal{S}_\alpha|} \left| \sum_{s \in \mathcal{S}_\alpha} \left(\Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}}(\mathcal{A}^\mathcal{O}(m) = 1) - \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}}(\mathcal{A}^\mathcal{O}(m) = 1) \right) \right| \\
&= \frac{1}{|\mathcal{S}_\alpha|} \left| \sum_{\mathcal{R} \subset \mathcal{D}} \sum_{\{s, s'\} \in \mathcal{P}_{\mathcal{R}}} \left(\Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}}(\mathcal{A}^\mathcal{O}(m) = 1) - \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}}(\mathcal{A}^\mathcal{O}(m) = 1) \right. \right. \\
&\quad \left. \left. + \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s') \\ \mathcal{O} = \mathcal{O}_{s'}(0)}}(\mathcal{A}^\mathcal{O}(m) = 1) - \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s') \\ \mathcal{O} = \mathcal{O}_{s'}(1)}}(\mathcal{A}^\mathcal{O}(m) = 1) \right) \right| \\
&= \frac{1}{|\mathcal{S}_\alpha|} \left| \sum_{\mathcal{R} \subset \mathcal{D}} \sum_{\{s, s'\} \in \mathcal{P}_{\mathcal{R}}} \left(\Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\mathcal{R}) \\ \mathcal{O} = \mathcal{O}_s(0)}}(\mathcal{A}^\mathcal{O}(m) = 1) - \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\mathcal{R}) \\ \mathcal{O} = \mathcal{O}_s(1)}}(\mathcal{A}^\mathcal{O}(m) = 1) \right. \right. \\
&\quad \left. \left. + \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\mathcal{R}) \\ \mathcal{O} = \mathcal{O}_{s'}(1)}}(\mathcal{A}^\mathcal{O}(m) = 1) - \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\mathcal{R}) \\ \mathcal{O} = \mathcal{O}_{s'}(0)}}(\mathcal{A}^\mathcal{O}(m) = 1) \right) \right| = 0.
\end{aligned}$$

□

Before we end this section, we remark that the above proof implies that the SWAP test also grounds the advantage to zero:

Remark B.2. Consider a pair of swapped splits s and s' . Observe that for RETRAIN, $\mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) = \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s') =: \mathbb{P}_{\mathcal{M}}(\mathcal{R})$ since this probability only depends on $\mathcal{R}$, which is the same for s and s' . With $\mathcal{O}_s(0) = \mathcal{O}_{s'}(1)$ and $\mathcal{O}_s(1) = \mathcal{O}_{s'}(0)$, we have

$$\begin{aligned}
\overline{\text{Adv}}_{\{s, s'\}}(\mathcal{A}, \text{RETRAIN}) &= \frac{1}{2} \left| \Pr_{m \sim \mathbb{P}_{\mathcal{M}}(\mathcal{R})}(\mathcal{A}^{\mathcal{O}_s(0)}(m) = 1) - \Pr_{m \sim \mathbb{P}_{\mathcal{M}}(\mathcal{R})}(\mathcal{A}^{\mathcal{O}_s(1)}(m) = 1) \right. \\
&\quad \left. + \Pr_{m \sim \mathbb{P}_{\mathcal{M}}(\mathcal{R})}(\mathcal{A}^{\mathcal{O}_{s'}(1)}(m) = 1) - \Pr_{m \sim \mathbb{P}_{\mathcal{M}}(\mathcal{R})}(\mathcal{A}^{\mathcal{O}_{s'}(0)}(m) = 1) \right| = 0.
\end{aligned}$$

B.4 Proof of Theorem 3.5

We prove Theorem 3.5 in this section. Before this, we formally introduce the notion of *certified removal*.

Definition B.3 (Certified Removal [Guo et al., 2020]). For a fixed dataset $\mathcal{D}$, let LR and UL be a learning and an unlearning algorithm respectively, and denote $\mathcal{H} := \text{im}(\text{LR}) \cup \text{im}(\text{UL})$ ⁶ to be the hypothesis class containing all possible models that can be produced by LR and UL. Then, for any $\epsilon, \delta > 0$, the unlearning algorithm UL is said to be (ϵ, δ) -certified removal if for any $\mathcal{W} \subset \mathcal{H}$ and for any disjoint $\mathcal{R}, \mathcal{F} \subseteq \mathcal{D}$ (do not need to satisfy restriction (a) and (b)),

$$\begin{aligned}
\Pr(\text{UL}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) \in \mathcal{W}) &\leq e^\epsilon \Pr(\text{RETRAIN}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) \in \mathcal{W}) + \delta; \\
\Pr(\text{RETRAIN}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) \in \mathcal{W}) &\leq e^\epsilon \Pr(\text{UL}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) \in \mathcal{W}) + \delta.
\end{aligned}$$

We note that as $\text{RETRAIN}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) = \text{LR}(\mathcal{R})$, the above can be simplified to

$$\begin{aligned}
\Pr(\text{UL}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) \in \mathcal{W}) &\leq e^\epsilon \Pr(\text{LR}(\mathcal{R}) \in \mathcal{W}) + \delta \\
\Pr(\text{LR}(\mathcal{R}) \in \mathcal{W}) &\leq e^\epsilon \Pr(\text{UL}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}) \in \mathcal{W}) + \delta.
\end{aligned}$$

Now, we restate Theorem 3.5 for convenience.

⁶Here, $\text{im}(\cdot)$ denote the image of a function.

Theorem B.4. Given an (ϵ, δ) -certified removal unlearning algorithm UL with some $\epsilon, \delta > 0$, for any (potentially inefficient) adversary $\mathcal{A}$ against UL in an unlearning sample inference game $\mathcal{G}$, we have

$$\text{Adv}(\mathcal{A}, \text{UL}) \leq 2 \cdot \left(1 - \frac{2 - 2\delta}{e^\epsilon + 1}\right).$$

Proof. We start by considering an attack as differentiating between the following two hypotheses: the unlearning and the retraining. In particular, given a specific dataset split $s = (\mathcal{R}, \mathcal{F}, \mathcal{T}) \in \mathcal{S}_\alpha$ and an model m , consider

$$H_1: m = \text{UL}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}), \text{ and } H_2: m = \text{LR}(\mathcal{R}) = \text{RETRAIN}(\text{LR}(\mathcal{R} \cup \mathcal{F}), \mathcal{F}).$$

Alternatively, by writing the distribution of the unlearned models and the retrained models as $\mathbb{P}_{\mathcal{M}}(\text{UL}, s)$ and $\mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s)$, respectively, we may instead write

$$H_1: m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s), \text{ and } H_2: m \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s).$$

It turns out that by looking at the *type-I error* α and *type-II error* β , we can control the advantage of the adversary in this game easily. Firstly, denote the model produced under H_1 as m_1 , then under H_1 , the accuracy of the adversary is $\Pr_{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s)}(\mathcal{A}(m_1) = 1) = 1 - \alpha$. Similarly, by denoting the model produced under H_2 as m_2 , we have $\Pr_{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s)}(\mathcal{A}(m_2) = 1) = \beta$. Therefore, for this specific dataset split s , let's define the advantage of this adversary $\mathcal{A}$ for this attack as⁷

$$\widehat{\text{Adv}}_s(\mathcal{A}) := |1 - \alpha - \beta|.$$

The upshot is that since UL is an (ϵ, δ) -certified removal unlearning algorithm (Definition B.3), it is possible to control α and β , hence $\widehat{\text{Adv}}_s(\mathcal{A})$. To achieve this, since from the definition of certified removal, we're dealing with sub-collections of models, it helps to write α and β differently.

Let $\mathcal{H} := \text{supp}(\mathbb{P}_{\mathcal{M}}(\text{UL}, s)) \cup \text{supp}(\mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s))$ be the collection of all possible models, and denote $\mathcal{B} \subseteq \mathcal{H}$ to be the collection of models that the adversary $\mathcal{A}$ accepts, and $\mathcal{B}^c$ to denote its complement, i.e., the collection of models that the adversary $\mathcal{A}$ rejects. We can then re-write the type-I and type-II errors as

- Type-I error α : probability of rejecting H_1 when H_1 is true, i.e., $\alpha = \Pr(m_1 \in \mathcal{B}^c \mid s) = 1 - \Pr(m_1 \in \mathcal{B} \mid s)$.
- Type-II error β : probability of accepting H_2 when H_2 is false, i.e., $\beta = \Pr(m_2 \in \mathcal{B} \mid s)$.

With this interpretation and the fact that UL is (ϵ, δ) -certified removal, we know that

- $\Pr(m_1 \in \mathcal{B} \mid s) \leq e^\epsilon \Pr(m_2 \in \mathcal{B} \mid s) + \delta$, and
- $\Pr(m_2 \in \mathcal{B}^c \mid s) \leq e^\epsilon \Pr(m_1 \in \mathcal{B}^c \mid s) + \delta$.

Combining the above, we have $1 - \alpha \leq e^\epsilon \beta + \delta$ and $1 - \beta \leq e^\epsilon \alpha + \delta$. Hence,

$$\beta \geq \max\{0, 1 - \delta - e^\epsilon \alpha, e^{-\epsilon}(1 - \delta - \alpha)\}.$$

We then seek to get the minimum of $\alpha + \beta$, we have

$$\alpha + \beta \geq \max\{\alpha, 1 - \delta - e^\epsilon \alpha + \alpha, e^{-\epsilon}(1 - \delta - \alpha) + \alpha\}.$$

To get a lower bound, consider the minimum among the last two, i.e., consider solving α when $1 - \delta - e^\epsilon \alpha + \alpha = e^{-\epsilon}(1 - \delta - \alpha) + \alpha$, leading to

$$(e^{-\epsilon} - e^\epsilon)\alpha = e^{-\epsilon}(1 - \delta) - (1 - \delta) = (e^{-\epsilon} - 1)(1 - \delta) \implies \alpha = \frac{(e^{-\epsilon} - 1)(1 - \delta)}{e^{-\epsilon} - e^\epsilon}.$$

Hence, we have

$$\alpha + \beta \geq 1 - \delta + \alpha(1 - e^\epsilon) = 1 - \delta + \frac{(e^{-\epsilon} - 1)(1 - \delta)}{e^{-\epsilon} - e^\epsilon}(1 - e^\epsilon) = (1 - \delta) \frac{2e^{-\epsilon} - 2}{e^{-\epsilon} - e^\epsilon} = (1 - \delta) \frac{2 - 2e^\epsilon}{1 - e^{2\epsilon}},$$

⁷Note that this is different from the advantage we defined before since the attack is different, hence we use a different notation.

with the elementary identity $1 - e^{2\epsilon} = (1 + e^\epsilon)(1 - e^\epsilon)$, we finally get

$$\alpha + \beta \geq \frac{2 - 2\delta}{e^\epsilon + 1}.$$

On the other hand, considering the “dual attack” that predicts the opposite as the original attack, that is, we flip $\mathcal{B}$ and $\mathcal{B}^c$. In this case, the type-I error and the type-II error become α and $1 - \beta$, respectively. Following the same procedures, we’ll have $\alpha + \beta \leq 2 - \frac{2-2\delta}{e^\epsilon+1}$.

Note that the definition of (ϵ, δ) -certified removal is independent of the dataset split s , hence, the above derivation works for all s . In particular, the advantage of any adversary differentiating H_1 and H_2 for any s is upper bounded by

$$\widehat{\text{Adv}}_s(\mathcal{A}) = |1 - \alpha - \beta| \leq 1 - \frac{2 - 2\delta}{e^\epsilon + 1} =: \tau,$$

where we denote τ to be the upper bound of advantage. This means that **any** adversaries trying to differentiate between retrain models and certified unlearned models are upper bounded in terms of their advantage, and an explicit upper bound is given by τ .

We now show a reduction from the unlearning sample inference game to the above. Firstly, we construct two attacks based on the adversary $\mathcal{A}$ in the unlearning sample inference game, which tries to differentiate between the data point x is sampled from the forget set $\mathcal{F}$ or the test set $\mathcal{T}$. This can be formulated through the following hypothesis testing:

$$H_3: x \in \mathcal{F}, \text{ and } H_4: x \in \mathcal{T}.$$

In this viewpoint, the unlearning sample inference game can be thought of as deciding between H_3 and H_4 . Since the upper bound we get for differentiating between H_1 and H_2 holds for any efficient adversaries, therefore, we can construct an attack for deciding between H_1 and H_2 using adversaries for deciding between H_3 and H_4 . This allows us to upper bound the advantage for the latter adversaries.

Given any adversaries $\mathcal{A}$ for differentiating H_3 and H_4 , i.e., any adversaries in the unlearning sample inference game, we start by constructing our first adversary $\mathcal{A}_1$ for differentiating H_1 and H_2 as follows:

- In the left world (H_1), feed the certified unlearned model to $\mathcal{A}$; in the right world H_2 , feed the retrained model to $\mathcal{A}$.
- We create a random oracle $\mathcal{O}_s(0)$ for $\mathcal{A}$, i.e., we let the adversary $\mathcal{A}$ decide on $\mathcal{F}$. We then let $\mathcal{A}_1$ output as $\mathcal{A}$.

We note that $\mathcal{A}$ is deciding on $\mathcal{F}$, the advantage of $\mathcal{A}_1$ is

$$\widehat{\text{Adv}}_s(\mathcal{A}_1) = \left| \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^\mathcal{O}(m_1) = 1) - \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^\mathcal{O}(m_2) = 1) \right|.$$

We can also induce the average of the advantage over all dataset splits is upper bounded by the maximal advantage taken over all dataset splits:

$$\begin{aligned} & \frac{1}{|\mathcal{S}_\alpha|} \left| \sum_{s \in \mathcal{S}_\alpha} \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^\mathcal{O}(m_1) = 1) - \sum_{s \in \mathcal{S}_\alpha} \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^\mathcal{O}(m_2) = 1) \right| \\ & \leq \frac{1}{|\mathcal{S}_\alpha|} \sum_{s \in \mathcal{S}_\alpha} \left| \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^\mathcal{O}(m_1) = 1) - \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^\mathcal{O}(m_2) = 1) \right| \\ & \leq \max_{s \in \mathcal{S}_\alpha} \widehat{\text{Adv}}_s(\mathcal{A}_1). \end{aligned}$$

Similarly, we can construct a second adversary $\mathcal{A}_2$ for differentiating H_1 and H_2 as follows:

- In the left world (H_1), feed the certified unlearned model to $\mathcal{A}$. In the right world (H_2), feed retrained model to $\mathcal{A}$.
- We create a random oracle $\mathcal{O}_s(1)$ for $\mathcal{A}$, i.e., we let the adversary $\mathcal{A}$ decide on $\mathcal{T}$. We then let $\mathcal{A}_2$ outputs as the $\mathcal{A}$.

Since $\mathcal{A}$ is deciding on $\mathcal{T}$, the advantage of $\mathcal{A}_2$ is

$$\widehat{\text{Adv}}_s(\mathcal{A}_2) = \left| \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathbb{1}(\mathcal{A}^{\mathcal{O}}(m_1) = 1)) - \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathcal{A}^{\mathcal{O}}(m_2) = 1) \right|.$$

Similar to the previous calculation for $\mathcal{A}_1$, the average of the advantage is also upper bounded by the maximal advantage,

$$\begin{aligned} & \frac{1}{|\mathcal{S}_{\alpha}|} \left| \sum_{s \in \mathcal{S}_{\alpha}} \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathcal{A}^{\mathcal{O}}(m_1) = 1) - \sum_{s \in \mathcal{S}_{\alpha}} \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathcal{A}^{\mathcal{O}}(m_2) = 1) \right| \\ & \leq \frac{1}{|\mathcal{S}_{\alpha}|} \sum_{s \in \mathcal{S}_{\alpha}} \left| \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathcal{A}^{\mathcal{O}}(m_1) = 1) - \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathcal{A}^{\mathcal{O}}(m_2) = 1) \right| \\ & \leq \max_{s \in \mathcal{S}_{\alpha}} \widehat{\text{Adv}}_s(\mathcal{A}_2). \end{aligned}$$

Given the above calculation, we can now bound the advantage of $\mathcal{A}$. Firstly, let UL be any certified unlearning method. Then the advantage $\text{Adv}(\mathcal{A}, \text{UL})$ of $\mathcal{A}$ in the unlearning sample inference game (i.e., differentiating between H_3 and H_4) against UL is

$$\frac{1}{|\mathcal{S}_{\alpha}|} \left| \sum_{s \in \mathcal{S}_{\alpha}} \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^{\mathcal{O}}(m_1) = 1) - \sum_{s \in \mathcal{S}_{\alpha}} \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathcal{A}^{\mathcal{O}}(m_1) = 1) \right|.$$

On the other hand, the advantage $\text{Adv}(\mathcal{A}, \text{RETRAIN})$ of $\mathcal{A}$ against the retraining method RETRAIN can be written as

$$\frac{1}{|\mathcal{S}_{\alpha}|} \left| \sum_{s \in \mathcal{S}_{\alpha}} \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^{\mathcal{O}}(m_2) = 1) - \sum_{s \in \mathcal{S}_{\alpha}} \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathcal{A}^{\mathcal{O}}(m_2) = 1) \right|,$$

which is indeed 0 from Theorem 3.3. Combine this with the calculations above, from the reverse triangle inequality,

$$\begin{aligned}
& \text{Adv}(\mathcal{A}, \text{UL}) \\
&= |\text{Adv}(\mathcal{A}, \text{UL}) - \text{Adv}(\mathcal{A}, \text{RETRAIN})| \\
&\leq \frac{1}{|\mathcal{S}_\alpha|} \left| \sum_{s \in \mathcal{S}_\alpha} \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^\mathcal{O}(m_2) = 1) - \sum_{s \in \mathcal{S}_\alpha} \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathcal{A}^\mathcal{O}(m_2) = 1) \right. \\
&\quad \left. - \sum_{s \in \mathcal{S}_\alpha} \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^\mathcal{O}(m_1) = 1) + \sum_{s \in \mathcal{S}_\alpha} \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathcal{A}^\mathcal{O}(m_1) = 1) \right| \\
&\leq \frac{1}{|\mathcal{S}_\alpha|} \sum_{s \in \mathcal{S}_\alpha} \left| \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^\mathcal{O}(m_1) = 1) - \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(0)}} (\mathcal{A}^\mathcal{O}(m_2) = 1) \right| \\
&\quad + \frac{1}{|\mathcal{S}_\alpha|} \sum_{s \in \mathcal{S}_\alpha} \left| \Pr_{\substack{m_1 \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathcal{A}^\mathcal{O}(m_1) = 1) - \Pr_{\substack{m_2 \sim \mathbb{P}_{\mathcal{M}}(\text{RETRAIN}, s) \\ \mathcal{O} = \mathcal{O}_s(1)}} (\mathcal{A}^\mathcal{O}(m_2) = 1) \right| \\
&\leq \max_{s \in \mathcal{S}_\alpha} \widehat{\text{Adv}}_s(\mathcal{A}_1) + \max_{s \in \mathcal{S}_\alpha} \widehat{\text{Adv}}_s(\mathcal{A}_2) \\
&\leq 2\tau,
\end{aligned}$$

i.e., the advantage of any adversary against any certified unlearning method is bounded by 2τ . $\square$

B.5 Proof of Proposition 3.7

We now prove Proposition 3.7. We repeat the statement for convenience.

Proposition B.5. *For any two dataset splits $s_1, s_2 \in \mathcal{S}_\alpha$ satisfying non-degeneracy assumption, i.e., both $\mathbb{P}_{\mathcal{D}}|_{\mathcal{F}_i}(\mathcal{F}_1 \cap \mathcal{F}_2)$ and $\mathbb{P}_{\mathcal{D}}|_{\mathcal{T}_i}(\mathcal{T}_1 \cap \mathcal{T}_2)$ do not vanish polynomially faster in $|\mathcal{D}|$, then there exists a deterministic and efficient adversary $\mathcal{A}$ such that $\overline{\text{Adv}}_{\{s_1, s_2\}}(\mathcal{A}, \text{UL}) = 1$ for any unlearning method UL. In particular, $\overline{\text{Adv}}_{\{s_1, s_2\}}(\mathcal{A}, \text{RETRAIN}) = 1$.*

Proof. Consider any unlearning method UL, and design a random oracle $\mathcal{O}_{s_i}$ based on the split s_i for $i = 1, 2$ and a sensitivity distribution $\mathbb{P}_{\mathcal{D}}$ (which for simplicity, assume to have full support across $\mathcal{D}$), we see that

$$\begin{aligned}
\overline{\text{Adv}}_{\{s_1, s_2\}}(\mathcal{A}, \text{UL}) &= \frac{1}{2} \left| \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s_1) \\ \mathcal{O} = \mathcal{O}_{s_1}(0)}} (\mathcal{A}^\mathcal{O}(m) = 1) - \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s_1) \\ \mathcal{O} = \mathcal{O}_{s_1}(1)}} (\mathcal{A}^\mathcal{O}(m) = 1) \right. \\
&\quad \left. + \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s_2) \\ \mathcal{O} = \mathcal{O}_{s_2}(0)}} (\mathcal{A}^\mathcal{O}(m) = 1) - \Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s_2) \\ \mathcal{O} = \mathcal{O}_{s_2}(1)}} (\mathcal{A}^\mathcal{O}(m) = 1) \right|.
\end{aligned}$$

Consider a hard-coded adversary $\mathcal{A}$ which has a look-up table T , defined as

$$T(x) = \begin{cases} 1, & \text{if } x \in \mathcal{T}_1 \cap \mathcal{T}_2; \\ 0, & \text{if } x \in \mathcal{F}_1 \cap \mathcal{F}_2; \\ \perp, & \text{otherwise,} \end{cases}$$

where we use $\perp$ to denote an undefined output. Then, $\mathcal{A}$ predicts the bit b used by the oracle as follows: We see that Algorithm 1 with a hard-coded look-up table T has several properties:

- Since it neglects m entirely,

$$\Pr_{\substack{m \sim \mathbb{P}_{\mathcal{M}}(\text{UL}, s_i) \\ \mathcal{O} = \mathcal{O}_{s_i}(b)}} (\mathcal{A}^\mathcal{O}(m) = 1) = \Pr_{\mathcal{O} = \mathcal{O}_{s_i}(b)} (\mathcal{A}^\mathcal{O} = 1)$$

Algorithm 1: Dummy adversary $\mathcal{A}$ against a random 2-sets evaluation

Data: An unlearned model m , a random oracle $\mathcal{O}$

Result: A one bit prediction b

```
 $b \leftarrow \perp$   
while  $b = \perp$  do  
   $x \sim \mathcal{O}$   
   $b \leftarrow T(x)$   
return  $b$ 
```

- Under the non-degeneracy assumptions, $\mathcal{A}$ will terminate in polynomial time in $|\mathcal{D}|$. This is because the expected terminating time is inversely proportional to $\mathbb{P}_{\mathcal{D}|\mathcal{F}_i}(\mathcal{F}_1 \cap \mathcal{F}_2)$ and $\mathbb{P}_{\mathcal{D}|\mathcal{T}_i}(\mathcal{T}_1 \cap \mathcal{T}_2)$, hence if these two probabilities does not vanish polynomially faster in $|\mathcal{D}|$ (i.e., the non-degeneracy assumption), then it'll terminate in polynomial time in $|\mathcal{D}|$.
- Whenever $\mathcal{A}$ terminates and outputs an answer, it will be correct, i.e.,

$$\overline{\text{Adv}}_{\{s_1, s_2\}}(\mathcal{A}, \text{UL}) = 1.$$

Since the above argument works for every UL, hence even for the retraining method RETRAIN, we will have $\overline{\text{Adv}}_{\{s_1, s_2\}}(\mathcal{A}, \text{RETRAIN}) = 1$. Intuitively, such a pathological case can happen since there exists some $\mathcal{A}$ which interpolates the “correct answer” for a few splits. Though adversaries may not have access to specific dataset splits, learning-based attacks could still undesirably learn towards this scenario if evaluated only on a few splits. Thus, we should penalize hard-coded adversaries in evaluation. $\square$

B.6 Discussion on a naive strategy offsetting MIA accuracy

In this section, we consider a toy example to illustrate the difference between the proposed advantage-based metric and a naive strategy that simply offsets the MIA accuracy for RETRAIN to zero.

Let us assume there are 6 data points $\{A, B, C, D, E, F\}$, and we equally split them into the forget set $\mathcal{F}$ and the test set $\mathcal{T}$. Running MIA against RETRAIN on a retrained model m^* trained on the retain set $\mathcal{R}$ independent of the above 6 data points, the MIA predicts the probability that each data point belongs to the forget set is:

$$\begin{aligned} \text{MIA}(m^*, A) &= 0.7, & \text{MIA}(m^*, B) &= 0.4, & \text{MIA}(m^*, C) &= 0.3, \\ \text{MIA}(m^*, D) &= 0.1, & \text{MIA}(m^*, E) &= 0.6, & \text{MIA}(m^*, F) &= 0.8. \end{aligned}$$

Assume that the MIA adversary $\mathcal{A}$ chooses the cutoff of predicting a data point belonging to the forget set to be ≥ 0.5 , then the prediction will be

$$\mathcal{A}(A) = 1, \quad \mathcal{A}(B) = 0, \quad \mathcal{A}(C) = 0, \quad \mathcal{A}(D) = 0, \quad \mathcal{A}(E) = 1, \quad \mathcal{A}(F) = 1,$$

where 1 refers to the forget set while 0 refers to the test set. Since the retrained model will be the same regardless of the split of the forget set and the test set, the MIA prediction will be the same regardless of the split as well.

Assuming that we have two imperfect unlearning algorithms, UL_1 and UL_2 . For simplicity, we could assume running MIA on the unlearned model m_1 by UL_1 will increase the predicted probability on the forget set by 0.1, while that (denoted as m_2) by UL_2 will increase it by 0.2. For example, if we set $\mathcal{F} = \{A, B, C\}$ while $\mathcal{T} = \{D, E, F\}$, then the MIA on UL_1 will predict the probability as

$$\begin{aligned} \text{MIA}(m_1, A) &= 0.8, & \text{MIA}(m_1, B) &= 0.5, & \text{MIA}(m_1, C) &= 0.4, \\ \text{MIA}(m_1, D) &= 0.1, & \text{MIA}(m_1, E) &= 0.6, & \text{MIA}(m_1, F) &= 0.8, \end{aligned}$$

while the MIA on UL_2 will predict the probability as

$$\begin{aligned} \text{MIA}(m_2, A) &= 0.9, & \text{MIA}(m_2, B) &= 0.6, & \text{MIA}(m_2, C) &= 0.5, \\ \text{MIA}(m_2, D) &= 0.1, & \text{MIA}(m_2, E) &= 0.6, & \text{MIA}(m_2, F) &= 0.8, \end{aligned}$$

Intuitively, UL_2 is worse than UL_1 in terms of unlearning quality. In this simplified setup, we claim the following.

Claim B.6. *The advantage calculated by one SWAP test over RETRAIN is always 0, while the advantage calculated by one SWAP test over UL_1 and UL_2 are respectively $1/6$ and $1/3$, all regardless of the exact split of the data points. Hence, the Unlearning Quality for RETRAIN, UL_1 , and UL_2 are respectively 1, $5/6$, and $2/3$, faithfully reflecting the unlearning algorithms' performances.*

On the other hand, the “offset MIA accuracy” is dependent on the split of the data. Specifically, when we assign $\{D, E, F\}$ as the forget set and $\{A, B, C\}$ as the test set, the MIA accuracies for all three methods are the same, making the “offset MIAs” all equal to 0.5, failing to capture the unlearning quality.

Proof. Given a split s , denote the predicted MIA result as $\hat{Y}_i^s \in \{0, 1\}$, and the actual membership as $Y_i^s = \mathbb{1}_{i \in \mathcal{F}}$ for $i \in \{A, B, \dots, F\}$. Then, consider a simple adversary $\mathcal{A}$: after getting the predicted MIA probability, i.e., $\Pr(\hat{Y}_i^s = 1)$, we update $\hat{Y}_i^s$ as 1 if $\Pr(\hat{Y}_i^s = 1) \geq 1/2$, and 0 otherwise. Then, the advantage of a particular split s is $\Pr_m(\mathcal{A}^{\mathcal{O}_s^{(0)}}(m) = 1) - \Pr_m(\mathcal{A}^{\mathcal{O}_s^{(1)}}(m) = 1) = \Pr_i(\hat{Y}_i^s = 1 \mid Y_i^s = 0) - \Pr_i(\hat{Y}_i^s = 1 \mid Y_i^s = 1)$ for $i \in \{A, B, \dots, F\}$. These probabilities are essentially an average of indicator variables: for example, $\Pr_i(\hat{Y}_i^s = 1 \mid Y_i^s = 0) = \sum_{i: Y_i^s = 0} \mathbb{1}_{\hat{Y}_i^s = 1} / |\{i: Y_i^s = 0\}|$, and since we're considering equal-size splits, $|\{i: Y_i^s = 0\}| = |\{i: Y_i^s = 1\}| = 6/2 = 3$ in our example.

To see how the calculation works out, we note that for a SWAP split (s, s') , each i appears in the forget set and the test set exactly once when calculating the SWAP advantage, i.e., $Y_i^s = 1$ and $Y_i^{s'} = 0$ or $Y_i^s = 0$ and $Y_i^{s'} = 1$. Furthermore, suppose $\hat{Y}_i^s$ and $\hat{Y}_i^{s'}$ are the same, e.g., when considering RETRAIN, then the indicators $\mathbb{1}_{\hat{Y}_i^s = 1}$ and $\mathbb{1}_{\hat{Y}_i^{s'} = 1}$ will be the same and appear in pair (specifically, with opposite sign, one in Adv_s and another in $\text{Adv}_{s'}$), and hence cancel each other out. This is why the advantage is 0 for RETRAIN. This pairing of indicators for i under splits s and s' happens for imperfect unlearning algorithms, but the indicator might change: $\hat{Y}_i^s$ can swap from 0 to 1 due to imperfect unlearning. If this happens, i contributes $1/3$ to the denominator of the SWAP advantage formula, hence $1/6$ in total (divided by 2 at the end). With this observation, we see that for UL_1 , only B 's prediction will be flipped from 0 to 1, hence UL_1 's advantage is $1/6$ in the SWAP test. The same argument applies for UL_2 where both B and C 's predictions are flipped, hence the advantage is $2/6 = 1/3$. $\square$

C Omitted details from Section 4

C.1 Computational resource and complexity

We conduct our experiment on Intel(R) Xeon(R) Gold 6338 CPU @ 2.00GHz with 4 A40 NVIDIA GPUs. It takes approximately 6 days to reproduce the experiment of standard deviation comparison between SWAP test and random dataset splitting. It takes approximately 1 day to reproduce the experiment of dataset size and random seeds. Furthermore, it takes approximately 4 days to reproduce the experiment of differential private testing.

C.2 Details of training

For target model training without differential privacy (DP) guarantees, we consider using the ResNet-20 [He et al., 2016] as our target model and train it with Stochastic Gradient Descent (SGD) [Ruder, 2016] optimizer with a MultiStepLR learning rate scheduler with milestones [100, 150] and an initial learning rate of 0.1, momentum 0.9, weight decay 10^{-5} . Moreover, we train the model with 200 epochs, and we empirically observe that this guarantees convergence. For a given dataset split, we average 3 models to approximate the randomness induced in training and unlearning procedures.

For training DP models, we use DP-SGD [Abadi et al., 2016] to provide DP guarantees. Specifically, we adopt the OPACUS implementation [Yousefpour et al., 2021] and use ResNet-18 [He et al., 2016] as our target model. The model is trained with the RMSProp optimizer using a learning rate of 0.01 and of 20 epochs. This ensures convergence as we empirically observe that 20 epochs suffice to yield a comparable model accuracy. Considering the dataset size, we use $\delta = 10^{-5}$ and tune the max gradient norm individually.

C.3 IC score and MIA score

One of the two metrics we choose to compare against is the Interclass Confusion (IC) Test [Goel et al., 2023]. In brief, the IC test “confuses” a selected set of two classes by switching their labels. This is implemented by picking two classes and randomly selecting half of the data from each class for confusion. Then the IC test proceeds to train the corresponding target models on the new datasets and perform unlearning on the selected set using the unlearning algorithm being tested, and finally measures the inter-class error of the unlearned models on the selected set, which we called the *memorization score* γ . Similar to the advantage, the memorization score is between $[0, 1]$, and the lower, the better since ideally, the unlearned model should have no memorization of the confusion. Given this, to compare the IC test with the Unlearning Quality $\mathcal{Q}$, we consider $1 - \gamma$, and refer to this new score as the *IC score*.

On the other hand, the MIA AUC is a popular MIA-based metric to measure the performance of the unlearning. It measures how MIA performs by calculating the AUC (Area Under Curve) of MIA on the union of the test set and the forget set. We note that AUC is a widely used evaluation metric in terms of classification models since compared to directly measuring the accuracy, AUC tends to measure how well the model can discriminate against each class. Finally, as defined in Section 4, we let the *MIA score* be $1 - \text{MIA AUC}$ to have a fair comparison.

Model Accuracy versus DP Budgets. We also report the classification accuracy of the original model trained with various DP budgets in Table 5. As can be seen, the classification accuracy increases as the ϵ is relaxed to a larger value, showing the inherent trade-off between DP and utility. For experiments about Unlearning Quality in Table 3 and MIA score in Table 4(B), the original models are shared and thus have the same results (Table 5(B)). We note that measuring the IC scores requires dataset modifications, so the model accuracy in the experiments of IC score (Table 5(A)) differs slightly from that in experiments of Unlearning Quality and MIA score (Table 5(B)).

Table 5: Model accuracy *versus* DP Budgets. See Table 3 for more context.

(A) Results for experiments in Table 4(A).					(B) Results for experiments in Tables 3 and 4(B).				
ϵ	50	150	600	∞	ϵ	50	150	600	∞
Accuracy	0.442*	0.506*	0.540*	0.639*	Accuracy	0.485*	0.520*	0.571*	0.660*

Remark C.1. We would like to clarify the performance difference compared to the common literature, which can be attributed to the dataset split. The original dataset is split evenly into the target and shadow datasets for the purpose of implementing MIA. Within the target dataset, further partitioning is performed to create the retain, forget, and test sets. As a result, only about 30% of the full dataset remains available for training the model, significantly reducing the effective training data. Note that the data split is necessary for our experiments, so we cannot get significantly more training data.

We experimented with training on the full dataset and applied data augmentation while keeping all other configurations unchanged. With the full dataset, the model achieved an accuracy of 85.34%. After incorporating data augmentation, the accuracy further improved to 91.13%, aligning with the past literature. This simple ablation study validates that the performance difference mainly comes from the difference in data size and the omission of data augmentation.

We select a large ϵ in our differential privacy experiments due to the significant drop in accuracy observed when ϵ is small, which stems from the same dataset size limitation. We include an additional experiment with a smaller ϵ in the linear setting, as presented at the end of Appendix C.4.

C.4 Additional experiments

In this section, we provide additional ablation experiments on our proposed Unlearning Quality metric by considering varying various parameters and settings.

Unlearning Quality Versus Model Architecture. We provide additional experimental results under different model architectures. The experiment is conducted with the CIFAR10 dataset and $\alpha = 0.1$. The results are shown in Table 6. Interestingly, we observe once again that the relative ranking of different unlearning methods stays mostly consistent across different architectures.

Table 6: Unlearning Quality *versus* different model architectures. The relative ranking of different unlearning methods stays mostly consistent under different architectures.

UL	ResNet44	ResNet56	ResNet110
RETRFINAL	0.497 ± 0.040	0.473 ± 0.010	0.476 ± 0.036
FTFINAL	0.495 ± 0.041	0.471 ± 0.011	0.477 ± 0.039
FISHER	0.847 ± 0.051	0.832 ± 0.032	0.895 ± 0.020
NEGGRAD	0.562 ± 0.025	0.537 ± 0.016	0.520 ± 0.042
SALUN	0.716 ± 0.008	0.692 ± 0.013	0.672 ± 0.033
SSD	0.939 ± 0.053	0.935 ± 0.056	0.968 ± 0.017

Unlearning Quality Versus Dataset. We provide additional experiments on vision datasets CIFAR100 [Krizhevsky et al., 2009] and MNIST [LeCun, 1998], and natural language dataset SST5 [Socher et al., 2013]. The experiment is conducted with the ResNet20 model architecture and $\alpha = 0.1$ on CIFAR100 and CIFAR10. For SST5 dataset, the experiment is conducted on the BERT model [Devlin et al., 2019] and $\alpha = 0.1$. We note that CIFAR100 has 100 classes and 50000 training images, SST5 has 5 classes and 11855 training sentences while MNIST has 10 classes and 60000 training images. CIFAR100 is considered more challenging than CIFAR10, while MNIST is considered easier than CIFAR10. The results are shown in Table 7.

In this experiment, besides the consistency we have observed throughout this section, we in addition observe that the Unlearning Quality reflects the *level of difficulties* of unlearning on different datasets. Specifically, the Unlearning Quality of most unlearning methods is higher on MNIST while lower on CIFAR100, in comparison to those on CIFAR10 and SST5.

Table 7: Unlearning Quality *versus* different datasets. In addition to the consistency of the Unlearning Quality across unlearning methods, the Unlearning Quality scores are higher on MNIST and SST5 while lower on CIFAR100, in comparison to those on CIFAR10, reflecting the level of difficulties on different datasets.

UL	CIFAR100	MNIST	SST5
RETRFINAL	0.464 ± 0.027	0.976 ± 0.020	0.404 ± 0.013
FTFINAL	0.462 ± 0.028	0.977 ± 0.021	0.404 ± 0.013
FISHER	0.606 ± 0.008	0.990 ± 0.002	0.552 ± 0.014
NEGGRAD	0.669 ± 0.016	0.980 ± 0.017	0.376 ± 0.017
SALUN	0.697 ± 0.082	0.995 ± 0.002	0.456 ± 0.015
SSD	0.923 ± 0.058	0.998 ± 0.001	0.592 ± 0.011

Validation of Theorem 3.5 With Linear Models We experimented with a method from Guo et al. [2020], which is an unlearning algorithm for linear models with (ϵ, δ) -certified removal guarantees. We followed the experimental setup in Guo et al. [2020], training a linear model on part of the MNIST dataset for a binary classification task distinguishing class 3 from class 8.

In their algorithm, the parameter ϵ controls a budget indicating the extent of data that can be unlearned. During the iterative unlearning, when the accumulated gradient residual norm is beyond the unlearning budget, the unlearning guarantee is broken and retraining will kick in. So ϵ cannot be made arbitrarily small. Below, we report the Advantage metric for their unlearning algorithm with different ϵ (δ is fixed as $1e - 4$), as well as the Retrain method as reference:

Table 8: Change of advantage in linear models with respect to decreasing ϵ .

ϵ	0.8	0.6	0.4	0.3	RETRAIN
Advantage	0.010	0.009	0.005	0.003	0.002

We can see that the Advantage monotonically decreases as ϵ decreases, which aligns with our Theorem 3.5.