
Curvature Clues: Decoding Deep Learning Privacy with Input Loss Curvature

Deepak Ravikumar Efstathia Soufleri Kaushik Roy

Department of Electrical and Computer Engineering

Purdue University

West Lafayette, IN 47907

{dravikum, esoufler, kaushik}@purdue.edu

Abstract

In this paper, we explore the properties of loss curvature with respect to input data in deep neural networks. Curvature of loss with respect to input (termed input loss curvature) is the trace of the Hessian of the loss with respect to the input. We investigate how input loss curvature varies between train and test sets, and its implications for train-test distinguishability. We develop a theoretical framework that derives an upper bound on the train-test distinguishability based on privacy and the size of the training set. This novel insight fuels the development of a new black box membership inference attack utilizing input loss curvature. We validate our theoretical findings through experiments in computer vision classification tasks, demonstrating that input loss curvature surpasses existing methods in membership inference effectiveness. Our analysis highlights how the performance of membership inference attack (MIA) methods varies with the size of the training set, showing that curvature-based MIA outperforms other methods on sufficiently large datasets. This condition is often met by real datasets, as demonstrated by our results on CIFAR10, CIFAR100, and ImageNet. These findings not only advance our understanding of deep neural network behavior but also improve the ability to test privacy-preserving techniques in machine learning.^{1 2}

1 Introduction

Deep neural networks are being increasingly trained on sensitive datasets; thus ensuring the privacy of these models is paramount. Membership inference attacks (MIA) have become the standard approach to test a model's privacy [Murakonda and Shokri, 2020]. These attacks take a trained model and aim to identify if a given example was used in its training. Recent work has linked curvature of loss with respect to input with memorization [Garg et al., 2024] and differential privacy [Dwork et al., 2006, Ravikumar et al., 2024]. Inspired by this line of research, we investigate the properties of input loss curvature and leverage our insights to develop a new membership inference attack.

Curvature of loss with respect to input (termed input loss curvature) is defined as the trace of the Hessian of loss with respect to the input [Moosavi-Dezfooli et al., 2019, Garg and Roy, 2023]. Prior works that study loss curvature have focused on two lines of research. The first line of research has focused on studying the loss curvature with respect to the weights of the deep neural net [Keskar et al., 2017, Wu et al., 2020, Jiang et al., 2020, Foret et al., 2021, Kwon et al., 2021, Andriushchenko and Flammarion, 2022] to better understand generalization. The second line of research studied the loss curvature with respect to the input (i.e. data) for gaining insight into adversarial robustness

¹Code available at <https://github.com/DeepakTatachar/Curvature-Clues>

²Project website with models and other assets is available at <https://engineering.purdue.edu/NRL/projects/curvature-clues>

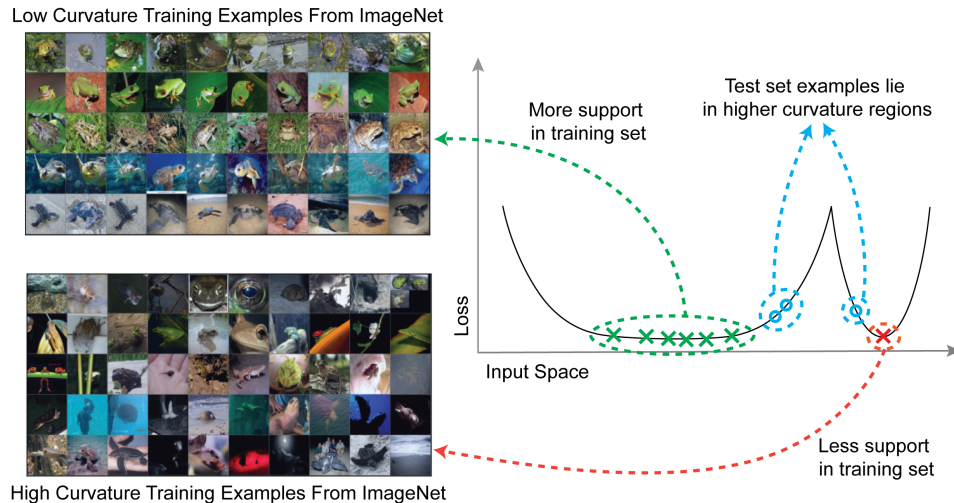


Figure 1: Visualizing low and high input curvature samples from a ResNet50 trained on ImageNet. Low input curvature training set images are prototypical and have lots of support in the trainset, while high input curvature train set examples have less support and are atypical. Test set examples lie around the training set images in higher curvature regions.

[Moosavi-Dezfooli et al., 2019, Fawzi et al., 2018], coresets [Garg and Roy, 2023] and memorization [Garg et al., 2024, Ravikumar et al., 2024], mainly focusing on its properties on the train set. Thus, there is a gap in our understanding of input loss curvature on unseen test examples.

Input loss curvature on the train set captures the prototypicality of an image. This is visualized in Figure 1 (left) which shows high curvature examples from ImageNet [Russakovsky et al., 2015] train set. Figure 1 (right) builds intuition as to why this is the case – a sample with lots of support in the dataset is more likely to lie in low curvature regions while atypical/non-prototypical examples have less support in the training set and thus lie in higher curvature regions [Garg et al., 2024, Ravikumar et al., 2024]. With this established, we make a novel observation that on average the input loss curvature scores for test set examples are higher than train set examples. This is because test samples were not optimized for, hence they lie slightly off the flat minima, in regions of higher input curvature (also visualized in Figure 1). We leverage this insight and develop a theoretical framework that focuses on the distinguishability of train-test input loss curvature scores. Studying this distinguishability leads us to understand performance bounds of membership inference attacks [Shokri et al., 2017, Sablayrolles et al., 2019, Song and Mittal, 2021, Carlini et al., 2022].

Our theoretical analysis obtains an upper bound on the KL divergence between train-test input loss curvature scores. This provides a ceiling for membership inference attack performance. The analysis also reveals that the upper bound when using input loss curvature scores is dependent on the number of training set examples and the differential privacy parameter ϵ . This insight helps us understand the conditions under which input loss curvature can be more effective at detecting train vs. test examples. Conditions such as the number of training examples and the privacy guarantees of the model.

To test our theoretical results we perform black-box (setting where adversaries have access only to model outputs) membership inference attack (MIA) using input loss curvature scores. However, input loss curvature score calculation needs access to model parameters (which are not available in a black box setting). To resolve this issue we propose using a zero-order input loss curvature estimation. Zero-order estimation can calculate input loss curvature without needing access to model parameters. Our input loss curvature-based black-box membership inference attack results show that real datasets contain enough training examples for curvature scores to outperform current state-of-the-art membership inference methods. Our results show that curvature based MIA outperforms prior state-of-the-art techniques on 10% or larger subsets of CIFAR100 dataset (about 5000 sample training set).

In summary, our contributions are as follows:

- **Theoretical Foundation:** We provide theoretical analysis to understand the train-test distinguishability with input loss curvature scores, demonstrating that input loss curvature is more

sensitive and hence better at detecting train vs. test set examples than current state-of-the-art membership inference attacks.

- **Adapting Theory To Practice:** We propose using zero order input loss curvature estimation to enable black box membership inference attack using input loss curvature scores.
- **Better Attack:** We conduct experiments to validate our theoretical results. Specifically, we show that input loss curvature enables more effective black box membership inference attacks when compared to existing state-of-the-art techniques.

2 Related Work

Membership Inference Attacks are used as a tool to test privacy [Murakonda and Shokri, 2020]. These attacks aim to identify if a particular data point was included in a model’s training dataset [Shokri et al., 2017]. Existing techniques often leverage model outputs like loss values [Shokri et al., 2017, Yeom et al., 2018, Sablayrolles et al., 2019], confidence scores [Carlini et al., 2022] or modified entropy [Song and Mittal, 2021]. Shokri et al. [2017] proposed the use of shadow models which are auxiliary models trained on subsets of the target model’s data to aid in inference. Several modifications to this approach have been proposed. One important addition is the focus on example hardness, where the authors Sablayrolles et al. [2019] proposed an attack that scaled the loss using per example hardness threshold which is estimated by training shadow models. Watson et al. [2022] proposed a similar approach but in the offline case where they calibrated the example hardness using the average loss of shadow models not trained on the target example. Ye et al. [2022] models various attacks into four categories and performs differential analysis to explain the gaps between them. Both Ye et al. [2022] and Long et al. [2020] consider the entire loss distribution of samples that are not in the training set. However, they face challenges in extrapolating to low false positive rates (FPRs). To address this issue Carlini et al. [2022] propose using a parametric model along with shadow models to improve performance. Orthogonal to these approaches Choquette-Choo et al. [2021] suggests the use of input augmentations during evaluation to improve the performance of the attacks. Similarly Jayaraman et al. [2021] propose the MERLIN attack, which queries the target model multiple times on a sample perturbed with Gaussian noise.

Input Loss Curvature is defined as the trace of the Hessian of loss with respect to the input [Moosavi-Dezfooli et al., 2019, Garg and Roy, 2023]. The aim is to measure the sensitivity of the deep neural network to a specific input. In general, loss curvature with respect to the parameters of deep neural nets has received lots of attention [Keskar et al., 2017, Wu et al., 2020, Jiang et al., 2020, Foret et al., 2021, Kwon et al., 2021, Andriushchenko and Flammarion, 2022], specifically due to its role in characterizing the sharpness of the learning objective which is closely connected to generalization. However, loss curvature with respect to the input data has received much less attention. It has been studied in the context of adversarial robustness [Fawzi et al., 2018, Moosavi-Dezfooli et al., 2019], coresets [Garg and Roy, 2023]. It has recently been linked with memorization [Garg et al., 2024, Ravikumar et al., 2024] and privacy [Ravikumar et al., 2024]. The authors in Moosavi-Dezfooli et al. [2019] showed that adversarial training decreases the input loss curvature and provided a theoretical link between robustness and curvature. In an orthogonal direction, Garg and Roy [2023] proposed the use of low input loss curvature examples as training dataset sets called coresets, which they showed to be more data-efficient. However, all of these works have focused on input loss curvature on the trainset. In this paper we focus on input loss curvature and its behavior on test or unseen examples to understand and improve the performance of membership inference attacks. Before we discuss our contributions, we present a few preliminaries, notation, and background needed for this paper

3 Notation and Background

Notation. Let us consider a supervised learning problem, where the goal is to learn a mapping from an input space $\mathcal{X} \subset \mathbb{R}^d$ to an output space $\mathcal{Y} \subset \mathbb{R}$. The learning is performed using a randomized algorithm $\mathcal{A}$ on a training set S . Note that a randomized algorithm employs a degree of randomness as a part of its logic. The training set S contains m elements denoted as $z_1, \dots, z_m$. Each element $z_i = (x_i, y_i)$ is drawn from an unknown distribution $\mathcal{D}$, where $z_i \in \mathcal{Z}, x_i \in \mathcal{X}, y_i \in \mathcal{Y}$ and $\mathcal{Z} = \mathcal{X} \times \mathcal{Y}$. Thus, we define the training set $S \in \mathcal{Z}^m$ as $S = \{z_1, \dots, z_m\}$. Another related concept is that of adjacent datasets, which are obtained when the set’s i^{th} element is removed and

defined as

$$S^{\setminus i} = \{z_1, \dots, z_{i-1}, z_{i+1}, \dots, z_m\}.$$

Additionally, the concept of adjacent (or neighboring) datasets is linked to the distance between datasets. The distance between two datasets S and S' , denoted by $\|S - S'\|_1$, measures the number of samples that differ between them. The notation $\|S\|_1$ represents the size of a dataset S . When a randomized learning algorithm $\mathcal{A}$ is applied to a dataset S , it produces a hypothesis denoted by $h_S^\phi = \mathcal{A}(\phi, S)$, where $\phi \sim \Phi$ is the random variable associated with the algorithm's randomness. A cost function $c : \mathcal{Y} \times \mathcal{Y} \rightarrow \mathbb{R}^+$ is used to measure the hypothesis's performance. The cost of the hypothesis h at a sample z_i is also referred to as the loss ℓ at z_i , defined as $\ell(h, z_i) = c(h(x_i), y_i)$. The performance of a hypothesis h is measured using risk $R(h) = \mathbb{E}_{z_i}[\ell(h, z_i)]$ and approximated using empirical risk $R_{emp}(h) = (1/m) \sum_{i=1}^m \ell(h, z_i)$.

Differential Privacy. Introduced by Dwork et al. [2006], differential privacy is defined as follows: A randomized algorithm $\mathcal{A}$ with domain $\mathcal{Z}^m$ is considered ϵ -differentially private if for any subset $\mathcal{R} \subset \text{Range}(\mathcal{A})$ and for any datasets $S, S' \in \mathcal{Z}^m$ differing in at most one element (i.e. $\|S - S'\|_1 \leq 1$):

$$\Pr_\phi[h_S^\phi \in \mathcal{R}] \leq e^\epsilon \Pr_\phi[h_{S'}^\phi \in \mathcal{R}], \quad (1)$$

where the probability is taken over the randomness of the algorithm $\mathcal{A}$, with $\phi \sim \Phi$.

Error Stability of a possibly randomized algorithm $\mathcal{A}$ for some $\beta > 0$ is defined as Kearns and Ron [1997]:

$$\forall i \in \{1, \dots, m\}, \quad \left| \mathbb{E}_{\phi, z}[\ell(h_S^\phi, z)] - \mathbb{E}_{\phi, z}[\ell(h_{S^{\setminus i}}^\phi, z)] \right| \leq \beta, \quad (2)$$

where $z \sim \mathcal{D}$ and $\phi \sim \Phi$.

Generalization. A randomized algorithm $\mathcal{A}$ is considered to generalize with confidence δ and at a rate of $\gamma'(m)$ if:

$$\Pr[|R_{emp}(h, S) - R(h)| \leq \gamma'(m)] \geq \delta. \quad (3)$$

Uniform Model Bias. The hypothesis h produced by algorithm $\mathcal{A}$ to learn the true conditional $h^* = \mathbb{E}[y|x]$ from a dataset $S \sim \mathcal{D}^m$ has uniform bound on model bias denoted by Δ if:

$$\forall S \sim \mathcal{D}^m, \quad \left| \mathbb{E}_\phi[R(h_S^\phi) - R(h^*)] \right| \leq \Delta. \quad (4)$$

ρ -Lipschitz Hessian. The Hessian of ℓ is Lipschitz continuous on $\mathcal{Z}$, if $\forall z_1, z_2 \in \mathcal{Z}$, and $\forall h \in \text{Range}(\mathcal{A})$, there exists some $\rho > 0$ such that:

$$\|\nabla_{z_1}^2 \ell(h, z_1) - \nabla_{z_2}^2 \ell(h, z_2)\| \leq \rho \|z_1 - z_2\|. \quad (5)$$

Input Loss Curvature. As defined by Moosavi-Dezfooli et al. [2019], Garg et al. [2024], input loss curvature is the sum of the eigenvalues of the Hessian H of the loss with respect to input z_i . This is conveniently expressed using the trace as:

$$\text{Curv}_\phi(z_i, S) = \text{tr}(H) = \text{tr}(\nabla_{z_i}^2 \ell(h_S^\phi, z_i)) \quad (6)$$

v -adjacency. A dataset S is said to contain v -adjacent (read as epsilon-adjacent) elements if it contains two elements z_i, z_j such that $z_j = z_i + \alpha$ for some $\alpha \in B_p(v)$ (read as v -Ball). This condition can be ensured through construction. Consider a dataset S' which has no z_j s.t. $z_j = z_i + \alpha; z_j, z_i \in S'$. We can construct S such that $S = \{z | z \in S'\} \cup \{z_i + \alpha\}$ for some $z_i \in S', \alpha \in B_p(v)$, thereby ensuring v -adjacency.

Membership Inference Threat Model. In a membership inference security game [Yeom et al., 2018, Jayaraman et al., 2021, Carlini et al., 2022], a challenger and an adversary interact to test the privacy of a machine learning model. The process begins with the challenger sampling a training dataset from a distribution and training a model on this data. The challenger then flips a coin to decide whether to select a fresh data point from the distribution, which is not part of the training set, or to choose a data point from the training set. This selected data point is given to the adversary, who has access to the same data distribution and the trained model. The adversary's task is to determine whether the given data point was part of the training set or not. If the adversary's guess is correct, the game indicates a successful membership inference attack. Such a game is said to be in a **black-box** setting when the adversary has access to only the challenger's output.

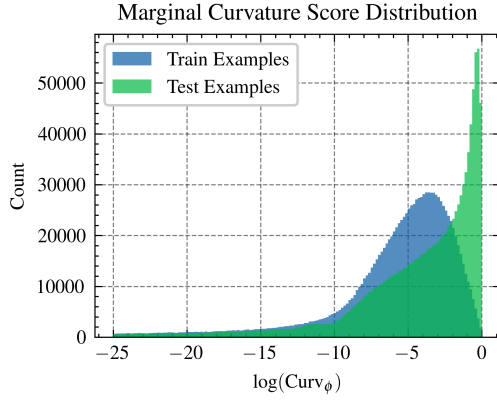


Figure 2: Marginal curvature score histogram for all images in ImageNet when samples in train set vs test set.

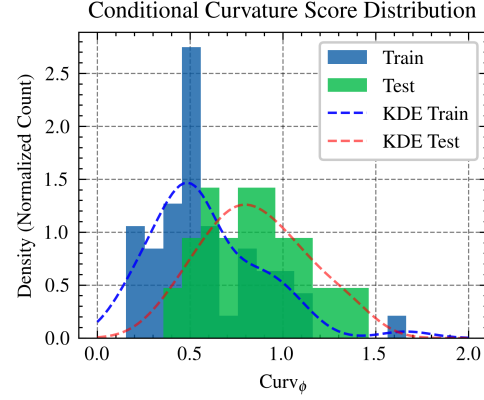


Figure 3: Conditional curvature scores for a single image from ImageNet when using 56 shadow models.

4 Theoretical Analysis

In this section, we analyze the distinguishability of train-test samples using KL divergence. We do so for two cases, the first case when using network’s output probability, and second when using input loss curvature. The importance of analyzing network output probabilities stems from its utilization in the current state-of-the-art attack, LiRA [Carlini et al., 2022]. Thus, studying both cases will let us theoretically analyze and compare their performance with input loss curvature.

Before we begin the analysis, we briefly discuss the conditional and marginal distributions of input loss curvature for train and test examples. Discussing this is important to build intuition for the analysis. Figure 2 visualizes the histogram (proxy for distribution) of input loss curvature for train and test examples from ImageNet [Russakovsky et al., 2015]. Specifically, it plots the log of the input loss curvature $\log(\text{Curv}_\phi)$ obtained on pre-trained ResNet50 [He et al., 2016] models from Feldman and Zhang [2020]. It plots a histogram of these scores, showing the frequency of specific log curvature values.

A naive membership inference attack would apply a threshold to separate these distributions. However, it is common practice to consider sample specific scores (i.e., conditioned on the target sample). This is visualized in Figure 3 which plots the distribution of curvature scores for a single ImageNet sample, indicating differences when the sample is part of the training set versus when it is a test or unseen sample. The data was generated using models from Feldman and Zhang [2020]. The figure also includes a kernel density estimate (KDE, shown as a dashed line) to better visualize the underlying distribution. Figure 3 suggests that, similar to Carlini et al. [2022] sample conditional curvature scores can be modeled using a Gaussian parametric model. If we represent the curvature score Curv_ϕ as a random variable, then $\text{Curv}_\phi \sim \mathcal{N}(\mu, \sigma)$. The probability density function of Curv_ϕ is a function of the randomness of the algorithm ϕ , the dataset S and the i^{th} sample z_i and be denoted by $p_c(\phi, S, z_i)$. Similarly, let $p(\phi, S, z_i)$ denote the probability density function of the neural network’s output probability, which is also parameterized by the randomness of the algorithm ϕ , the dataset S and the i^{th} sample z_i . With this setup we present the following theoretical results on the upper bound on the KL divergence between train and test distribution for the two cases. Theorem 4.1 presents the upper bound on the KL divergence when using the neural network’s output probability scores, and Theorem 4.2 presents the upper bound on the KL divergence when using input loss curvature.

Theorem 4.1 (Privacy bounds Train-Test KL Divergence). *Assume ϵ -differential private algorithm, then the KL divergence between train-test distributions of the neural network’s output probability is upper bound by the differential privacy parameter ϵ given by:*

$$D_{\text{KL}} \left(p(\phi, S, z_i) \parallel p(\phi, S^{\setminus i}, z_i) \right) \leq \epsilon \quad (7)$$

Sketch of Proof. Given that the algorithm $\mathcal{A}$ is ϵ -differentially private, we know that the probability of the output $h_S^\phi \in \mathcal{R}$ on dataset S is bounded by e^ϵ times the probability of the same event

on the neighboring dataset $S^{\setminus i}$, i.e. $\Pr_{\phi}[h_S^{\phi} \in \mathcal{R}] \leq e^{\epsilon} \Pr_{\phi}[h_{S^{\setminus i}}^{\phi} \in \mathcal{R}]$. From this inequality it follows that the KL divergence between the output distributions of $\mathcal{A}$ on S and $S^{\setminus i}$ is bounded by $\text{KL}(p(\phi, S, z_i) \parallel p(\phi, S^{\setminus i}, z_i)) \leq \epsilon$. The full proof for Theorem 4.1 is provided in Appendix A.3.

Theorem 4.2 (Dataset Size and Privacy bound Curvature KL Divergence). *Let the assumptions of error stability 2, generalization 3, and uniform model bias 4 hold. Further, assume $0 \leq \ell \leq L$. Let the conditional distribution be parameterized with variance σ . Then, any ϵ -differential private algorithm using a dataset of size m with a probability at least $1 - \delta$ satisfies:*

$$\text{D}_{\text{KL}}(p_c(\phi, S, z_i) \parallel p_c(\phi, S^{\setminus i}, z_i)) \leq \frac{[Lm(1 - e^{-\epsilon}) + c]^2}{2\sigma^2} \quad (8)$$

$$c = (4m - 1)\gamma + 2(m - 1)\Delta + \frac{\rho}{6}\mathbb{E}[\|\alpha\|^3] + L \quad (9)$$

Sketch of Proof. We begin by assuming a Gaussian model for the curvature distribution and expand the formula for the KL divergence, which leads to expressions involving the mean and standard deviation of the curvature scores. To establish an upper bound for this expression, we prove a lemma concerning the upper bound of the mean curvature, utilizing results from Ravikumar et al. [2024]. This enables us to express the upper bound on the mean curvature in terms of the privacy parameter ϵ . Subsequently, this result is employed to derive an upper bound on the KL divergence of the curvature scores. The proof for Theorem 4.2 is provided in Appendix A.5.

Discussion. Theorem 4.1 and Theorem 4.2 provide the upper bound on KL Divergence between train and test distributions of the probability and input loss curvature scores, respectively. They imply the upper limit of MIA performance, and interestingly, the result of Theorem 4.2 suggests the role of the dataset size on MIA performance. However, the relation is not as straightforward as suggested by Equation 8. This is because, in real applications, parameter σ depends on ϵ , and so does the loss bound L . In fact, the loss bound is also dependent on the number of samples m [Bousquet and Elisseeff, 2002].

Theorem 4.3 (Dataset Size and Curvature MIA Performance). *Let the assumptions of error stability 2, generalization 3, and uniform model bias 4 hold. Further, assume $0 \leq \ell \leq L$, and the bounds of Theorem 4.1 and 4.2 are tight. Then, the performance of MIA using curvature scores exceeds that of confidence scores with a probability at least $1 - \delta$ when:*

$$m > \frac{(\sqrt{2\sigma^2\epsilon}) - c}{L(1 - e^{-\epsilon})} \quad (10)$$

Theorem 4.3 suggests that the performance of curvature based MIA will exceed that of probability score based methods when the size of the dataset used to train the target model exceed a certain threshold. Indeed, this is what we observe in practice (see section 6.4).

Sketch of Proof. The proof compares the upper bounds from Theorem 4.1 and Theorem 4.2, followed by a series of algebraic manipulations. This identifies the conditions under which the performance of curvature scores exceeds that of confidence scores. The full proof can be found in Appendix A.6.

On the validity of the assumptions. Before presenting our experiments to validate our theory, we briefly discuss the validity of our assumptions in practical settings. Research by Hardt et al. [2016] shows that stochastic gradient methods, such as stochastic gradient descent, achieve small generalization error and exhibit uniform stability. Thus, the assumptions of stability (Equation 2) and generalization (Equation 3) are justified. Model bias is a property of the model, and a uniform bound across different datasets seems reasonable. Note, uniform loss bound (and L independent of m as used by Theorem 4.3) holds true for certain losses and for statistical models and is often assumed in learning theory [Wang et al., 2016]. Lastly, the v -adjacency can be ensured through construction. For large datasets this may not be needed, because the size of the ball $B_p(v)$ is unconstrained. Hence two samples from the same class that are similar may suffice. Given the size of modern datasets, this assumption is also reasonable.

5 Zero Order Input Loss Curvature MIA

To test if input loss curvature based membership inference performs better than existing methods we need an efficient technique to estimate curvature. We are interested in black-box membership

inference attacks, where one does not have access to the target network’s parameters. However, current techniques use Hutchinson’s trace estimator [Hutchinson, 1989] to measure input loss curvature such as from Garg and Roy [2023], Garg et al. [2024] or Ravikumar et al. [2024]. This approach needs to evaluate the gradient and hence requires access to model parameters. To solve this issue, we propose using a zero-order estimation technique. Zero-order curvature estimation starts with a finite-difference estimation. Consider a function $f : \mathbb{R}^n \rightarrow \mathbb{R}$, then the Hessian at a given point z_i can be estimated as follows:

$$\nabla^2 f(z_i) = n^2 \frac{f(z_i + hv + hu) - f(z_i - hv + hu) - f(z_i + hv - hu) + f(z_i - hv - hu)}{4h^2} uv^\top \quad (11)$$

where h is a small increment (a hyper parameter in our case), and u, v are vectors in $\mathbb{R}^n$. In our case, to get the input loss curvature, we have $f \leftarrow \ell(g(x_i), y_i)$, where g is the neural network, ℓ is the loss function and $z_i = (x_i, y_i)$ are the image, label pair. The pseudo-code for obtaining input loss curvature score using zero order estimation shown in Algorithm 1 in Appendix A.1. To execute membership inference attack using input loss curvature scores, we propose the following methodology. First, we begin by training shadow models, similar to Shokri et al. [2017], Carlini et al. [2022]. These shadow models are used to obtain empirical estimates of parametric model for the curvature score described in Section 4. During the inference phase, we employ a likelihood ratio between the sample being in the train set vs test set parametric models to identify the membership status of a given sample (see Appendix A.2 for pseudo-code). In addition, we perform a negative log likelihood ratio test. We denote the results of likelihood test as ‘LR’ and the results of the negative log likelihood ratio test as ‘NLL’.

6 Experiments

6.1 Experimental Setup

Datasets. To evaluate our theory, we consider the classification task using standard vision datasets. Specifically, we use the CIFAR10, CIFAR100 [Krizhevsky et al., 2009] and ImageNet [Russakovsky et al., 2015] datasets.

Architectures. For experiments on ImageNet we use the ResNet50 architecture [He et al., 2016]. For CIFAR10 and CIFAR100, we used the ResNet18 architecture. Details regarding hyperparameters are provided in Appendix A.13. To improve reproducibility, we have open sourced the code at <https://github.com/DeepakTatachar/Curvature-Clues>.

Training. For experiments using private models, we trained ResNet18 models with the Opacus library [Yousefpour et al., 2021] using DP-SGD with a maximum gradient norm of 1.0 and a privacy parameter of $\delta = 1 \times 10^{-5}$. Shadow models for CIFAR10 and CIFAR100 were trained on a 50% subset of the data for 300 epochs. For ImageNet, we used pre-trained models from Feldman and Zhang [2020], trained on a 70% subset of ImageNet. More details about training and compute resources are provided in Appendix A.13.

Metrics. To evaluate curvature scores, we use AUROC and balanced accuracy. The Receiver Operating Characteristic (ROC) is the plot of the True Positive Rate (TPR) against the False Positive Rate (FPR). The area under the ROC is called Area Under the Receiver Operating Characteristic (AUROC). AUROC of 1 denotes an ideal detection scheme, since the ideal detection algorithm results in 0 false positive and false negative samples.

6.2 Membership Inference

In this subsection, we compare the performance of the proposed input loss curvature based membership inference against prior MIA techniques.

Setup: We use CIFAR10, CIFAR100 and ImageNet datasets to test the MIA performance. We consider a **black-box** MIA setup similar to Carlini et al. [2022]. We use ResNet18 for CIFAR10 and CIFAR100, for ImageNet we use the ResNet50 architecture. For all the membership inference attacks, we compute a full ROC curve and report the results. When using shadow models we 64 for CIFAR10 and CIFAR100 and 52 for ImageNet. The AUROC plot for CIFAR100 for various

Method	ImageNet		CIFAR100		CIFAR10	
	Bal. Acc.	AUROC	Bal. Acc.	AUROC	Bal. Acc.	AUROC
Curv ZO NLL (Ours)	69.16 $\pm$ 0.08	77.45 $\pm$ 0.09	84.47 $\pm$ 0.21	93.49 $\pm$ 0.18	61.92 $\pm$ 0.87	68.82 $\pm$ 1.30
Curv ZO LR (Ours)	68.76 $\pm$ 0.04	72.28 $\pm$ 0.04	80.48 $\pm$ 0.10	90.15 $\pm$ 0.04	55.00 $\pm$ 0.17	58.89 $\pm$ 0.38
Carlini et al. [2022]	66.14 $\pm$ 0.01	73.46 $\pm$ 0.02	81.55 $\pm$ 0.13	88.89 $\pm$ 0.16	58.23 $\pm$ 0.29	61.73 $\pm$ 0.32
Yeom et al. [2018]	58.50 $\pm$ 0.02	63.23 $\pm$ 0.03	76.29 $\pm$ 0.39	82.11 $\pm$ 0.31	55.57 $\pm$ 0.52	60.44 $\pm$ 0.75
Sablayrolles et al. [2019]	66.93 $\pm$ 0.05	76.50 $\pm$ 0.04	70.22 $\pm$ 0.41	81.11 $\pm$ 0.39	56.65 $\pm$ 0.56	61.50 $\pm$ 0.79
Watson et al. [2022]	61.40 $\pm$ 0.06	69.44 $\pm$ 0.05	62.71 $\pm$ 0.31	71.66 $\pm$ 0.50	54.86 $\pm$ 0.59	58.58 $\pm$ 0.86
Ye et al. [2022]	66.16 $\pm$ 0.02	75.79 $\pm$ 0.05	80.73 $\pm$ 0.24	90.88 $\pm$ 0.19	59.62 $\pm$ 0.84	67.30 $\pm$ 1.25
Song and Mittal [2021]	57.88 $\pm$ 0.03	63.29 $\pm$ 0.03	75.58 $\pm$ 0.29	82.28 $\pm$ 0.27	55.63 $\pm$ 0.61	60.42 $\pm$ 0.85

Table 1: Comparison of the proposed curvature score based MIA with prior methods tested on ImageNet, CIFAR100, and CIFAR10 datasets. Results reported are the mean $\pm$ std obtained over 3 seeds. For CIFAR10 and CIFAR100 64 shadow models were used and 52 for ImageNet.

methods are shown in Figure 4. Table 1 reports the average balanced accuracy and AUROC values over three seeds on various MIA methods on the three datasets. The plot in Figure 4 is a log-log plot to emphasize performance of the proposed method at very low false positive rates (see the orange line y-intercept). We also studied the effect of augmentations and the results can be found in Appendix A.7, the takeaway was that adding more augmentations improved performance. Note that the results presented in Table 1 used 2 augmentations (image + mirror) for our technique and Carlini et al. [2022] for fair comparison.

Results: From Table 1, we see that the proposed method performs better than all existing MIA techniques on both ImageNet and CIFAR datasets. Apart from AUROC and balanced accuracy results, the log-log plot emphasizes the performance at really low FPR (i.e. the y intercept, high TPR at low FPR in Figure 4). Additional results at low FPR are available in Appendix A.8.

Takeaways: As predicted by Theorem 4.2, the higher KL divergence between train and test curvature score distributions results in superior MIA performance. Further, we observe that while using a negative log likelihood ratio test does better in AUROC and balanced accuracy, the parametric likelihood ratio test does better at low false positive rates as shown in Figure 4. Thus, the proposed use of curvature should be tailored based on use case. Applications that demand high AUROC can use NLL approach, while those that demand high TPR at very low FPR should use the LR technique.

6.3 Effect of Privacy

In this section, we study the effect of differential privacy on MIA performance and test the ϵ relation predicted by Theorem 4.2.

Setup: To study the effect of privacy on MIA attack performance, we use DP-SGD trained models. We use privacy ϵ values of 1, 2, 3, 4, 5, 10, 15, 20, 25, 30, 35, 40, 45, 50 with $\delta = 1 \times 10^{-5}$. We

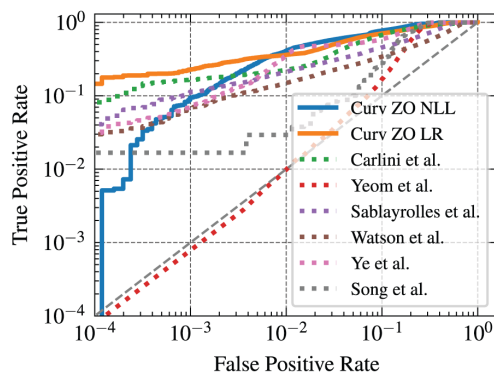


Figure 4: Comparing our method against existing techniques at low FPR. The proposed parametric Curv LR technique has the highest TPR at very low FPR.

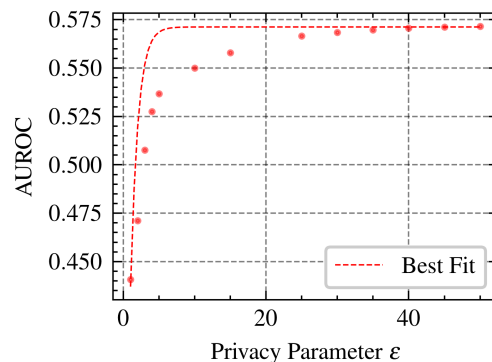


Figure 5: Validating the upper bound from Theorem 4.2 by fitting the MIA performance (AUROC of Curv LR) on DP-SGD trained models for various privacy parameters ϵ values.

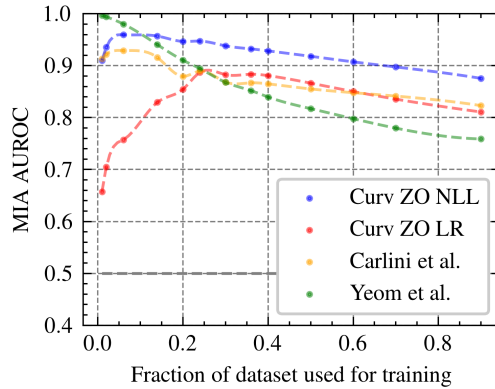


Figure 6: Visualizing MIA performance as a function of the size of the train set, which is randomly sampled.

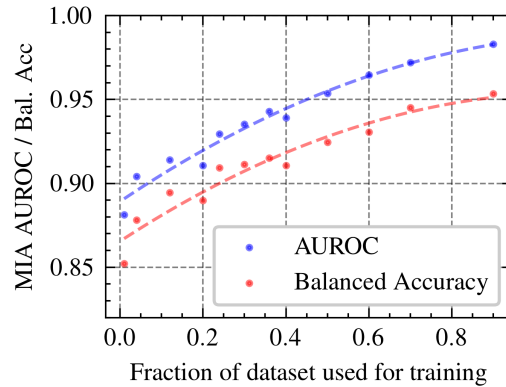


Figure 7: MIA performance as a function of the size of the train set when subsets contain the lowest curvature examples from the entire set.

plot the result of the proposed curvature based MIA in Figure 5. In Figure 5 we also plots the best fit using $s_f(L_f(1 - e^{-\epsilon}) + c_f)^2$, where s_f , L_f and c_f are fit to the data.

Results and Takeaways: We see that the theoretical prediction from Theorem 4.2 about MIA performance is well matched. Since Theorem 4.2 provides an upper bound, the results validate the theory.

6.4 Effect of Dataset Size

In this section, we study the effect of dataset size on MIA attack performance. This lets us test the relationship of MIA performance to m as predicted by Theorem 4.2 and Theorem 4.3.

Setup: For this experiment, we train models on increasing dataset size on CIFAR100. Specifically, we train multiple seeds on various subsets randomly chosen from the CIFAR100 training set. We also repeated the same by choosing the lowest curvature samples from CIFAR100 and train with the same subset sizes. Next, for each of the models we perform MIA attack and we present the results in Figure 6 (randomly chosen) and Figure 7 (lowest curvature samples).

Results: From Figure 6, we see that when samples are randomly chosen, the MIA performance decreases as we add more samples. This result is consistent with prior literature [Abadi et al., 2016]. Further as predicted by Theorem 4.3 beyond 30 – 40% subset Curv ZO LR out perform prior works and Curv ZO NLL outperforms prior works beyond 10% subset size.

To extend this analysis, we train models on curvature based coresets [Garg and Roy, 2023]. These coresets of low input loss curvature samples have been shown to memorize less [Garg et al., 2024, Ravikumar et al., 2024]. Thus we expect MIA accuracy to increase as we increase coreset size. This is exactly what happens and is shown in Figure 7 which plots the AUROC and accuracy of the NLL curvature attack as we increase curvature coreset size. However, the MIA performance is higher for comparable size in Figure 6 and 7. **This suggests that curvature based coresets result in more susceptible models**, which is also supported by results in Song et al. [2019].

Takeaways: We validate Theorem 4.3. We note that beyond a certain dataset size (of about 30 – 40% subset of the training set Curv ZO LR out perform prior works and Curv ZO NLL outperforms prior works beyond 10% subset size) curvature scores outperform probability score based MIA method. While Theorem 4.3 predicts that curvature-based MIA outperforms other methods on sufficiently large datasets. This condition is often met by real datasets as evident from the results presented.

7 Conclusion

In this paper, we explored the properties of input loss curvature on the test set. Specifically, we focus on using input loss curvature to distinguish between train and test examples. We established a theoretical framework deriving an upper bound on train-test KL Divergence based on privacy and

training set size. To extend the applicability of input loss curvature computation to a black-box setting, we propose a novel zero-order curvature estimation method. This enables the development of a new cutting-edge black-box membership inference attack (MIA) methodology that leverages input loss curvature. Through extensive experiments on the ImageNet and CIFAR datasets, we demonstrate that our input loss curvature-based MIA method outperforms existing state-of-the-art techniques. Our results corroborate our theoretical predictions regarding the relationship between MIA performance and dataset size. Specifically, we show that beyond a certain dataset size, the effectiveness of curvature scores surpasses other methods. Importantly, we observe that this dataset size condition is frequently met in practical scenarios, as evidenced by our results on the CIFAR100 dataset. These findings advance our understanding of input loss curvature in the context of privacy and help build more secure deep learning models.

Acknowledgment

This work was supported in part by, the Center for the Co-Design of Cognitive Systems (COCOSYS), a DARPA-sponsored JUMP 2.0 center, the Semiconductor Research Corporation (SRC), and the National Science Foundation.

References

- M. Abadi, A. Chu, I. Goodfellow, H. B. McMahan, I. Mironov, K. Talwar, and L. Zhang. Deep learning with differential privacy. In *Proceedings of the 2016 ACM SIGSAC conference on computer and communications security*, pages 308–318, 2016.
- M. Andriushchenko and N. Flammarion. Towards understanding sharpness-aware minimization. In *International Conference on Machine Learning*, pages 639–668. PMLR, 2022.
- J. Ansel, E. Yang, H. He, N. Gimpel, A. Jain, M. Voznesensky, B. Bao, P. Bell, D. Berard, E. Burovski, G. Chauhan, A. Chourdia, W. Constable, A. Desmaison, Z. DeVito, E. Ellison, W. Feng, J. Gong, M. Gschwind, B. Hirsh, S. Huang, K. Kalambarkar, L. Kirsch, M. Lazos, M. Lezcano, Y. Liang, J. Liang, Y. Lu, C. Luk, B. Maher, Y. Pan, C. Puhersch, M. Reso, M. Saroufim, M. Y. Siraichi, H. Suk, M. Suo, P. Tillet, E. Wang, X. Wang, W. Wen, S. Zhang, X. Zhao, K. Zhou, R. Zou, A. Mathews, G. Chanan, P. Wu, and S. Chintala. PyTorch 2: Faster Machine Learning Through Dynamic Python Bytecode Transformation and Graph Compilation. In *29th ACM International Conference on Architectural Support for Programming Languages and Operating Systems, Volume 2 (ASPLOS '24)*. ACM, Apr. 2024. doi: 10.1145/3620665.3640366. URL <https://pytorch.org/assets/pytorch2-2.pdf>.
- O. Bousquet and A. Elisseeff. Stability and generalization. *The Journal of Machine Learning Research*, 2:499–526, 2002.
- N. Carlini, S. Chien, M. Nasr, S. Song, A. Terzis, and F. Tramèr. Membership inference attacks from first principles. In *2022 IEEE Symposium on Security and Privacy (SP)*, pages 1897–1914. IEEE, 2022.
- C. A. Choquette-Choo, F. Tramèr, N. Carlini, and N. Papernot. Label-only membership inference attacks. In *International conference on machine learning*, pages 1964–1974. PMLR, 2021.
- A. Dosovitskiy, L. Beyer, A. Kolesnikov, D. Weissenborn, X. Zhai, T. Unterthiner, M. Dehghani, M. Minderer, G. Heigold, S. Gelly, J. Uszkoreit, and N. Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=YicbFdNTTy>.
- C. Dwork, F. McSherry, K. Nissim, and A. Smith. Calibrating noise to sensitivity in private data analysis. In *Theory of Cryptography: Third Theory of Cryptography Conference, TCC 2006, New York, NY, USA, March 4-7, 2006. Proceedings 3*, pages 265–284. Springer, 2006.
- A. Fawzi, S.-M. Moosavi-Dezfooli, P. Frossard, and S. Soatto. Empirical study of the topology and geometry of deep networks. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 3762–3770, 2018.

- V. Feldman and C. Zhang. What neural networks memorize and why: Discovering the long tail via influence estimation. *Advances in Neural Information Processing Systems*, 33:2881–2891, 2020.
- P. Foret, A. Kleiner, H. Mobahi, and B. Neyshabur. Sharpness-aware minimization for efficiently improving generalization. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=6Tm1mposlrm>.
- I. Garg and K. Roy. Samples with low loss curvature improve data efficiency. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 20290–20300, 2023.
- I. Garg, D. Ravikumar, and K. Roy. Memorization through the lens of curvature of loss function around samples. In *Forty-first International Conference on Machine Learning*, 2024. URL <https://openreview.net/forum?id=WQbDS9RyY>.
- M. Hardt, B. Recht, and Y. Singer. Train faster, generalize better: Stability of stochastic gradient descent. In *International conference on machine learning*, pages 1225–1234. PMLR, 2016.
- K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 770–778, 2016.
- M. F. Hutchinson. A stochastic estimator of the trace of the influence matrix for laplacian smoothing splines. *Communications in Statistics-Simulation and Computation*, 18(3):1059–1076, 1989.
- B. Jayaraman, L. Wang, K. Knipmeyer, Q. Gu, and D. Evans. Revisiting membership inference under realistic assumptions. *Proceedings on Privacy Enhancing Technologies*, 2021(2), 2021.
- Y. Jiang, B. Neyshabur, H. Mobahi, D. Krishnan, and S. Bengio. Fantastic generalization measures and where to find them. In *International Conference on Learning Representations*, 2020. URL <https://openreview.net/forum?id=SJgIPJBFvH>.
- M. Kearns and D. Ron. Algorithmic stability and sanity-check bounds for leave-one-out cross-validation. In *Proceedings of the tenth annual conference on Computational learning theory*, pages 152–162, 1997.
- N. S. Keskar, D. Mudigere, J. Nocedal, M. Smelyanskiy, and P. T. P. Tang. On large-batch training for deep learning: Generalization gap and sharp minima. In *International Conference on Learning Representations*, 2017. URL <https://openreview.net/forum?id=H1oyRlYgg>.
- A. Krizhevsky, G. Hinton, et al. Learning multiple layers of features from tiny images, 2009.
- J. Kwon, J. Kim, H. Park, and I. K. Choi. Asam: Adaptive sharpness-aware minimization for scale-invariant learning of deep neural networks. In *International Conference on Machine Learning*, pages 5905–5914. PMLR, 2021.
- Y. Long, L. Wang, D. Bu, V. Bindschaedler, X. Wang, H. Tang, C. A. Gunter, and K. Chen. A pragmatic approach to membership inferences on machine learning models. In *2020 IEEE European Symposium on Security and Privacy (EuroS&P)*, pages 521–534. IEEE, 2020.
- H. Mehta, A. Thakurta, A. Kurakin, and A. Cutkosky. Large scale transfer learning for differentially private image classification. *arXiv preprint arXiv:2205.02973*, 2022.
- S.-M. Moosavi-Dezfooli, A. Fawzi, J. Uesato, and P. Frossard. Robustness via curvature regularization, and vice versa. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 9078–9086, 2019.
- S. K. Murakonda and R. Shokri. MI privacy meter: Aiding regulatory compliance by quantifying the privacy risks of machine learning. *arXiv preprint arXiv:2007.09339*, 2020.
- Y. Nesterov and B. T. Polyak. Cubic regularization of newton method and its global performance. *Mathematical Programming*, 108(1):177–205, 2006.
- D. Ravikumar, E. Soufleri, A. Hashemi, and K. Roy. Unveiling privacy, memorization, and input curvature links. In *Forty-first International Conference on Machine Learning*, 2024. URL <https://openreview.net/forum?id=4dxR7aw05n>.

- O. Russakovsky, J. Deng, H. Su, J. Krause, S. Satheesh, S. Ma, Z. Huang, A. Karpathy, A. Khosla, M. Bernstein, A. C. Berg, and L. Fei-Fei. ImageNet Large Scale Visual Recognition Challenge. *International Journal of Computer Vision (IJCV)*, 115(3):211–252, 2015. doi: 10.1007/s11263-015-0816-y.
- A. Sablayrolles, M. Douze, C. Schmid, Y. Ollivier, and H. Jégou. White-box vs black-box: Bayes optimal strategies for membership inference. In *International Conference on Machine Learning*, pages 5558–5567. PMLR, 2019.
- R. Shokri, M. Stronati, C. Song, and V. Shmatikov. Membership inference attacks against machine learning models. In *2017 IEEE symposium on security and privacy (SP)*, pages 3–18. IEEE, 2017.
- K. Simonyan and A. Zisserman. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556*, 2014.
- L. Song and P. Mittal. Systematic evaluation of privacy risks of machine learning models. In *30th USENIX Security Symposium (USENIX Security 21)*, pages 2615–2632, 2021.
- L. Song, R. Shokri, and P. Mittal. Privacy risks of securing machine learning models against adversarial examples. In *Proceedings of the 2019 ACM SIGSAC Conference on Computer and Communications Security*, pages 241–257, 2019.
- C. Szegedy, W. Liu, Y. Jia, P. Sermanet, S. Reed, D. Anguelov, D. Erhan, V. Vanhoucke, and A. Rabinovich. Going deeper with convolutions. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 1–9, 2015.
- A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. u. Kaiser, and I. Polosukhin. Attention is all you need. In I. Guyon, U. V. Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017. URL https://proceedings.neurips.cc/paper_files/paper/2017/file/3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf.
- Y.-X. Wang, J. Lei, and S. E. Fienberg. Learning with differential privacy: Stability, learnability and the sufficiency and necessity of ermm principle. *Journal of Machine Learning Research*, 17(183): 1–40, 2016.
- L. Watson, C. Guo, G. Cormode, and A. Sablayrolles. On the importance of difficulty calibration in membership inference attacks. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=3eIrli0TwQ>.
- D. Wu, S.-T. Xia, and Y. Wang. Adversarial weight perturbation helps robust generalization. *Advances in Neural Information Processing Systems*, 33:2958–2969, 2020.
- J. Ye, A. Maddi, S. K. Murakonda, V. Bindschaedler, and R. Shokri. Enhanced membership inference attacks against machine learning models. In *Proceedings of the 2022 ACM SIGSAC Conference on Computer and Communications Security*, pages 3093–3106, 2022.
- S. Yeom, I. Giacomelli, M. Fredrikson, and S. Jha. Privacy risk in machine learning: Analyzing the connection to overfitting. In *2018 IEEE 31st computer security foundations symposium (CSF)*, pages 268–282. IEEE, 2018.
- A. Yousefpour, I. Shilov, A. Sablayrolles, D. Testuggine, K. Prasad, M. Malek, J. Nguyen, S. Ghosh, A. Bharadwaj, J. Zhao, G. Cormode, and I. Mironov. Opacus: User-friendly differential privacy library in PyTorch. *arXiv preprint arXiv:2109.12298*, 2021.

A Appendix

A.1 Zero Order Curvature Estimation

We present the pseudo-code for zero-order curvature score calculation below. Please note that Algorithm 1 shows the input loss curvature calculation for one example; however, it can be easily implemented for batch data.

Algorithm 1 Pseudo-code for Zero Order Input Loss Curvature Estimation

```
1: Input:  $z_i = (x_i, y_i)$  (image, label data),  $\ell$  (loss function),  $n_{iter}$  (hyperparameter: number of iterations),  $h$  (hyperparameter: small increment),  $g$  (neural network)
2: Output:  $curv\_estimate$  (zero order input loss curvature estimate of  $z_i$ )
3:  $curv \leftarrow 0$ 
4: for  $i \leftarrow 1$  to  $n_{iter}$  do
5:    $v \sim \text{Rademacher}$ 
6:    $u \sim \text{Rademacher}$ 
7:    $f \leftarrow \ell(g(x_i), y_i)$ 
8:    $\nabla^2 f(z_i) \leftarrow \text{Use Equation 11}$ 
9:    $curv \leftarrow curv + \text{tr}(\nabla^2 f(z_i))$ 
10: end for
11:  $curv\_estimate \leftarrow curv/n_{iter}$ 
12: return  $curv\_estimate$ 
```

A.2 Curvature based MIA

In this section, we present the detailed step by step description and pseudo code for the curvature based MIA (refer to Algorithm 2)

1. Initialization (Lines 3-4): Two empty sets, $curv_{in}$ and $curv_{out}$, are initialized. These sets will store the curvature scores for the shadow models trained with and without the target example z_i .
2. Shadow Model Training (Lines 5-12): For a specified number of iterations N , the algorithm performs the following steps:
 - Sampling Shadow Dataset (Line 6): A shadow dataset $\mathcal{D}_{attack}$ is sampled from the data distribution $\mathcal{D}$.
 - Training IN Shadow Model (Line 7): A shadow model g_{in}^n is trained on the shadow dataset augmented with the target example, $\mathcal{D}_{attack} \cup \{z_i\}$. The curvature score for this model, denoted as $curv_{in}^n$, is computed using Algorithm 1.
 - Collecting IN Curvature Scores (Line 8): The curvature score for the training set $\mathcal{D}_{attack} \cup \{z_i\}$ of this model $curv_{in}^n$ is added to the set $curv_{in}$.
 - Training OUT Shadow Model (Line 9): Another shadow model g_{out}^n is trained on the shadow dataset excluding the target example, $\mathcal{D}_{attack} \setminus \{z_i\}$. The curvature score for this model, denoted as $curv_{out}^n$, is also computed using Algorithm 1.
 - Collecting OUT Curvature Scores (Line 10): The curvature score for the training set $\mathcal{D}_{attack} \setminus \{z_i\}$ of this model $curv_{out}^n$ is added to the set $curv_{out}$.
3. Model Parameter Estimation (Lines 13-16): After training the shadow models, the algorithm calculates the mean (μ_{in}, μ_{out}) and variance ($\sigma_{in}^2, \sigma_{out}^2$) of the collected curvature scores for both the 'IN' and 'OUT' models.
4. Target Model Query (Line 17): The curvature score for the target model g_{target} with respect to the target example z_i , denoted as $curv_{target}$, is computed using Algorithm 1.
5. Likelihood Ratio Test (Line 18-19): The likelihood ratio P_{in} is calculated. This ratio compares the probability of the observed curvature score under the 'IN' distribution against the 'OUT' distribution and is returned. For the 'NLL' version P_{in} us given by:

$$P_{in}^{NLL} = \log [P (curv_{target} | \mathcal{N}(\mu_{out}, \sigma_{out}^2))] - \log [P (curv_{target} | \mathcal{N}(\mu_{in}, \sigma_{in}^2))] \quad (12)$$

Algorithm 2 Membership Inference Attack using Input Loss Curvature Scores

Require: Target model g_{target} , target example $z_i = (x_i, y_i)$, data distribution $\mathcal{D}$

```
1:  $curv_{in} \leftarrow \{\}$ 
2:  $curv_{out} \leftarrow \{\}$ 
3: for  $n$  in  $N$  do
4:    $\mathcal{D}_{attack} \leftarrow \mathcal{D}$  ▷ Sample a shadow dataset
5:    $g_{in}^n \leftarrow \text{Train on } \mathcal{D}_{attack} \cup \{z_i\}$ . ▷ Train in shadow model
6:    $curv_{in}^n \leftarrow \text{Use Algorithm 1}$ 
7:    $curv_{in} \leftarrow curv_{in} \cup curv_{in}^n$ 
8:    $g_{out}^n \leftarrow \text{Train on } \mathcal{D}_{attack} \setminus \{z_i\}$ . ▷ Train out shadow model
9:    $curv_{out}^n \leftarrow \text{Use Algorithm 1}$ 
10:   $curv_{out} \leftarrow curv_{out} \cup curv_{out}^n$ 
11: end for
12:  $\mu_{in} \leftarrow \text{mean}(curv_{in})$ 
13:  $\mu_{out} \leftarrow \text{mean}(curv_{out})$ 
14:  $\sigma_{in}^2 \leftarrow \text{var}(curv_{in})$ 
15:  $\sigma_{out}^2 \leftarrow \text{var}(curv_{out})$ 
16:  $curv_{target} \leftarrow \text{Use Algorithm 1 with } g_{target} \text{ and } z_i$ .
17:  $P_{in} = \frac{P(curv_{target} | \mathcal{N}(\mu_{in}, \sigma_{in}^2))}{P(curv_{target} | \mathcal{N}(\mu_{out}, \sigma_{out}^2))}$ 
18: return  $P_{in}$ 
```

A.3 Proof of Theorem 4.1

Consider an ϵ -DP algorithm and $S, S^{\setminus i}$ such that $\|S - S^{\setminus i}\| = 1$. Next let $\mathcal{R} \subset \text{Range}(\mathcal{A})$ such that $\mathcal{R} = \{h \mid h(x_i) = y_i\}$. Since $\mathcal{A}$ is ϵ -differentially private, then it follows from the definition of differential privacy in Equation 1 that

$$\Pr_{\phi}[h_S^{\phi} \in \mathcal{R}] \leq e^{\epsilon} \Pr_{\phi}[h_{S^{\setminus i}}^{\phi} \in \mathcal{R}] \quad (13)$$

$$\implies p(\phi, S, z_i) \leq e^{\epsilon} p(\phi, S^{\setminus i}, z_i) \quad (14)$$

Since $p(\phi, S, z_i) = \Pr_{\phi}[h_S^{\phi} \in \mathcal{R}]$ and denotes the probability density function of the neural network's output probability. Now, we use the definition of KL divergence

$$\begin{aligned} D_{\text{KL}}\left(p(\phi, S, z_i) \parallel p(\phi, S^{\setminus i}, z_i)\right) &= \mathbb{E}_{\phi} \left[\log \left(\frac{p(\phi, S, z_i)}{p(\phi, S^{\setminus i}, z_i)} \right) \right] \\ &\leq \mathbb{E}_{\phi} [\log e^{\epsilon}] \quad \text{From Equation 14} \\ D_{\text{KL}}\left(p(\phi, S, z_i) \parallel p(\phi, S^{\setminus i}, z_i)\right) &\leq \epsilon \quad \blacksquare \end{aligned}$$

A.4 Lemma A.1

Lemma A.1 (Upper bound on Input Loss Curvature). *Let $\mathcal{A}$ be a randomized algorithm which is ϵ -differentially private and the assumptions of error stability 2, generalization 3, and uniform model bias 4 hold. Further, assume $0 \leq \ell \leq L$. Then for two adjacent datasets $S, S^{\setminus i} \sim \mathcal{D}$ with a probability at least $1 - \delta$ we have:*

$$\mathbb{E}_{\phi}[\text{Curv}_{\phi}(z_i, S^{\setminus i})] \leq L(m(1 - e^{-\epsilon}) + 1) + c_1 \quad (15)$$

$$c_1 = (4m - 1)\gamma + 2(m - 1)\Delta + \frac{\rho}{6} \mathbb{E}[\|\alpha\|^3] \quad (16)$$

Proof of Lemma A.1 For this proof, we use results from Nesterov and Polyak [2006], we restate it here for convenience

Lemma A.2. *If Lipschitz assumption 5 on the Hessian of ℓ holds from Nesterov and Polyak [2006] we have*

$$|\ell(h, z_1) - \ell(h, z_2) - \langle \nabla \ell(h, z_2), z_1 - z_2 \rangle - \langle \nabla^2 \ell(h, z_2)(z_1 - z_2), z_1 - z_2 \rangle| \leq \frac{\rho}{6} |z_1 - z_2|^3 \quad (17)$$

From Lemma A.2 we have

$$-\frac{\rho}{6}|z_1 - z_2|^3 \leq \ell(h, z_1) - \ell(h, z_2) - \langle \nabla \ell(h, z_2), z_1 - z_2 \rangle - \langle \nabla^2 \ell(h, z_2)(z_1 - z_2), z_1 - z_2 \rangle$$

$$\ell(h, z_1) - \ell(h, z_2) - \langle \nabla \ell(h, z_2), z_1 - z_2 \rangle - \langle \nabla^2 \ell(h, z_2)(z_1 - z_2), z_1 - z_2 \rangle \leq \frac{\rho}{6}|z_1 - z_2|^3$$

This gives us a lower bound on $\ell(h, z_1)$

$$-\frac{\rho}{6}|z_1 - z_2|^3 + \ell(h, z_2) + \langle \nabla \ell(h, z_2), z_1 - z_2 \rangle + \langle \nabla^2 \ell(h, z_2)(z_1 - z_2), z_1 - z_2 \rangle \leq \ell(h, z_1) \quad (18)$$

Consider $z_j \in S$ such that $z_i = z_j + \alpha$ for some $j \neq i$ where $\alpha \in B_p(v)$ such that $\mathbb{E}[\alpha] = 0$ and $\mathbb{E}[\alpha^T \alpha] = 1$. Using the lower bound in Lemma A.2 from Ravikumar et al. [2024] with $z_1 = z_i, z_2 = z_j$ we get

$$\mathbb{E}_\phi[\ell(h_{S \setminus i}^\phi, z_j)] - \mathbb{E}_\phi[\ell(h_S^\phi, z_i)] \leq m\beta + (4m - 1)\gamma + 2(m - 1)\Delta \quad (19)$$

Using the result from Lemma A.2

$$-\frac{\rho}{6}\|\alpha\|^3 + \mathbb{E}_\phi[\ell(h_{S \setminus i}^\phi, z_i)] + \mathbb{E}_\phi[\langle \nabla \ell(h_{S \setminus i}^\phi, z_i), \alpha \rangle]$$

$$+ \mathbb{E}_\phi[\langle \nabla^2 \ell(h_{S \setminus i}^\phi, z_i)\alpha, \alpha \rangle] - \mathbb{E}_\phi[\ell(h_S^\phi, z_i)] \leq m\beta + (4m - 1)\gamma + 2(m - 1)\Delta \quad (20)$$

Taking Expectation over α and since the mean of $\alpha = 0$ we have

$$\mathbb{E}_\phi[\ell(h_{S \setminus i}^\phi, z_i)] + \mathbb{E}_{\alpha, \phi}[\langle \nabla \ell(h_{S \setminus i}^\phi, z_i), \alpha \rangle] + \mathbb{E}_{\alpha, \phi}[\langle \nabla^2 \ell(h_{S \setminus i}^\phi, z_i)\alpha, \alpha \rangle]$$

$$- \mathbb{E}_\phi[\ell(h_S^\phi, z_i)] \leq m\beta + (4m - 1)\gamma + 2(m - 1)\Delta + \frac{\rho}{6}\|\alpha\|^3$$

Note that we can change the order of expectation using Fubini's theorem

$$\mathbb{E}_\phi[\ell(h_{S \setminus i}^\phi, z_i)] + \mathbb{E}_{\phi, \alpha}[\langle \nabla^2 \ell(h_{S \setminus i}^\phi, z_i)\alpha, \alpha \rangle] - \mathbb{E}_\phi[\ell(h_S^\phi, z_i)] \leq m\beta + (4m - 1)\gamma + 2(m - 1)\Delta + \frac{\rho}{6}\|\alpha\|^3$$

$$\mathbb{E}_\phi[\ell(h_{S \setminus i}^\phi, z_i)] + \mathbb{E}_\phi[\text{tr}(\nabla^2 \ell(h_{S \setminus i}^\phi, z_i))] - \mathbb{E}_\phi[\ell(h_S^\phi, z_i)] \leq m\beta + (4m - 1)\gamma + 2(m - 1)\Delta + \frac{\rho}{6}\|\alpha\|^3$$

$$\mathbb{E}_\phi[\text{tr}(\nabla^2 \ell(h_{S \setminus i}^\phi, z_i))] \leq m\beta + (4m - 1)\gamma + 2(m - 1)\Delta + \frac{\rho}{6}\mathbb{E}[\|\alpha\|^3] + \mathbb{E}_\phi[\ell(h_S^\phi, z_i)] - \mathbb{E}_\phi[\ell(h_{S \setminus i}^\phi, z_j)]$$

$$\mathbb{E}_\phi[\text{tr}(\nabla^2 \ell(h_{S \setminus i}^\phi, z_i))] \leq m\beta + (4m - 1)\gamma + 2(m - 1)\Delta + \frac{\rho}{6}\mathbb{E}[\|\alpha\|^3] + L$$

From Ravikumar et al. [2024] Lemma 5.2 we have $\beta \leq L(1 - e^{-\epsilon})$. Thus we have the upper bound:

$$\mathbb{E}_\phi[\text{Curv}_\phi(z, S^{\setminus i})] \leq L(m(1 - e^{-\epsilon}) + 1) + c_1 \quad (21)$$

$$c_1 = (4m - 1)\gamma + 2(m - 1)\Delta + \frac{\rho}{6}\mathbb{E}[\|\alpha\|^3] \quad \blacksquare \quad (22)$$

A.5 Proof of Theorem 4.2

This proof uses results from Lemma A.1 provided above. For the proof of Theorem 4.2 let $p_c(\phi, S, z_i)$ be the curvature probability mass function, then using the Gaussian model discussed in Section 4 we model it as a Gaussian. Thus, we can write

$$p_c(\phi, S, z_i) = \frac{1}{\sigma_S \sqrt{2\pi}} e^{-\frac{(\phi - \mu_S)^2}{2\sigma_S^2}} \quad (23)$$

$$p_c(\phi, S^{\setminus i}, z_i) = \frac{1}{\sigma_{S \setminus i} \sqrt{2\pi}} e^{-\frac{(\phi - \mu_{S \setminus i})^2}{2\sigma_{S \setminus i}^2}} \quad (24)$$

Now consider the $D_{\text{KL}}(p_c(\phi, S, z_i) \parallel p_c(\phi, S^{\setminus i}, z_i))$

$$D_{\text{KL}}(p_c(\phi, S, z_i) \parallel p_c(\phi, S^{\setminus i}, z_i)) = \mathbb{E}_\phi \left[\log \left(\frac{p_c(\phi, S, z_i)}{p_c(\phi, S^{\setminus i}, z_i)} \right) \right] \quad (25)$$

$$= \mathbb{E}_\phi \left[-\frac{(\phi - \mu_S)^2}{2\sigma^2} + \frac{(\phi - \mu_{S^{\setminus i}})^2}{2\sigma^2} \right] \quad \text{Assume } \sigma = \sigma_{S^{\setminus i}} = \sigma_S \quad (26)$$

$$= \mathbb{E}_\phi \left[\frac{-\phi^2 - \mu_S^2 + 2\phi\mu_S + \phi^2 + \mu_{S^{\setminus i}}^2 - 2\phi\mu_{S^{\setminus i}}}{2\sigma^2} \right] \quad (27)$$

$$= \mathbb{E}_\phi \left[\frac{\mu_{S^{\setminus i}}^2 - \mu_S^2}{2\sigma^2} \right] \quad (28)$$

$$= \frac{\mu_{S^{\setminus i}}^2 - \mu_S^2}{2\sigma^2} \quad (29)$$

To upper bound Equation 29, we need an upper bound on $\mu_{S^{\setminus i}}$ and a lower bound on μ_S . The lower bound on $\mu_S^2 = 0$

For ease of notation let's define $\beta_{\max} := L(m(1 - e^{-\epsilon}) + 1)$. Notice that the upper bound of per sample $\mu_{S^{\setminus i}} = \mathbb{E}_\phi[\text{Curv}_\phi(z, S^{\setminus i})]$. Using the result from Lemma A.1 in Equation 29 we have

$$D_{\text{KL}}(p_c(\phi, S, z_i) \parallel p_c(\phi, S^{\setminus i}, z_i)) \leq \frac{(\beta_{\max} + c_1)^2 - 0^2}{2\sigma^2} \quad (30)$$

$$\leq \frac{(\beta_{\max} + c_1)^2}{2\sigma^2} \quad (31)$$

$$\leq \frac{[L(m(1 - e^{-\epsilon}) + 1) + c_1]^2}{2\sigma^2} \quad \blacksquare \quad (32)$$

A.6 Proof of Theorem 4.3

Assuming the bounds of Theorem 4.1 and Theorem 4.2 are tight, then using curvature score will outperform neural network probability scores if:

$$D_{\text{KL}}(p_c(\phi, S, z_i) \parallel p_c(\phi, S^{\setminus i}, z_i)) > D_{\text{KL}}(p(\phi, S, z_i) \parallel p(\phi, S^{\setminus i}, z_i)) \quad (33)$$

$$\frac{[L(m(1 - e^{-\epsilon}) + 1) + c_1]^2}{2\sigma^2} > \epsilon$$

$$\frac{[Lm(1 - e^{-\epsilon}) + c]^2}{2\sigma^2} > \epsilon \quad ; c = c_1 + L$$

$$L^2(1 - e^{-\epsilon})^2 m^2 + 2Lc(1 - e^{-\epsilon})m + c^2 - 2\sigma^2\epsilon > 0$$

Using the quadratic formula we have the roots given by:

$$\begin{aligned} &= \frac{-2Lc(1 - e^{-\epsilon}) \pm \sqrt{4L^2c^2(1 - e^{-\epsilon})^2 - 4L^2(1 - e^{-\epsilon})^2(c^2 - 2\sigma^2\epsilon)}}{2L^2(1 - e^{-\epsilon})^2} \\ &= \frac{-2Lc \pm \sqrt{4L^2c^2 - 4L^2(c^2 - 2\sigma^2\epsilon)}}{2L^2(1 - e^{-\epsilon})} \\ &= \frac{-Lc \pm \sqrt{L^2c^2 - L^2(c^2 - 2\sigma^2\epsilon)}}{L^2(1 - e^{-\epsilon})} \\ &= \frac{-c \pm \sqrt{c^2 - (c^2 - 2\sigma^2\epsilon)}}{L(1 - e^{-\epsilon})} \\ &= \frac{-c \pm \sqrt{2\sigma^2\epsilon}}{L(1 - e^{-\epsilon})} \end{aligned}$$

Assuming both roots are real, the larger root is $\frac{\sqrt{2\sigma^2\epsilon} - c}{L(1 - e^{-\epsilon})}$. Thus when $m > \frac{\sqrt{2\sigma^2\epsilon} - c}{L(1 - e^{-\epsilon})}$ Equation 33 is satisfied. ■

A.7 Input Augmentations Ablation

In this section we present the results of using 1 augmentation (only the image) vs 2 augmentations (image + mirror) in Tables 2, 3 and 4 for ImageNet, CIFAR100 and CIFAR10 respectively. For these results, we used 64 shadow models for CIFAR10 and CIFAR100 and 52 shadow models for ImageNet.

Method	ImageNet			
	1 Augmentation		2 Augmentations	
	Bal. Acc.	AUROC	Bal. Acc.	AUROC
Curv ZO NLL (Ours)	68.09 $\pm$ 0.07	75.84 $\pm$ 0.08	69.16 $\pm$ 0.08	77.45 $\pm$ 0.09
Curv ZO LR (Ours)	67.59 $\pm$ 0.05	70.62 $\pm$ 0.04	68.76 $\pm$ 0.04	72.28 $\pm$ 0.04
Carlini et al. [2022]	65.11 $\pm$ 0.03	71.96 $\pm$ 0.04	66.14 $\pm$ 0.01	73.46 $\pm$ 0.02

Table 2: Comparison of the proposed curvature score-based MIA with prior methods on ImageNet dataset with 1 and 2 augmentations. Results reported are the mean $\pm$ std obtained over 3 seeds.

Method	CIFAR100			
	1 Augmentation		2 Augmentations	
	Bal. Acc.	AUROC	Bal. Acc.	AUROC
Curv ZO NLL (Ours)	83.97 $\pm$ 0.23	92.72 $\pm$ 0.23	84.47 $\pm$ 0.21	93.49 $\pm$ 0.18
Curv ZO LR (Ours)	77.88 $\pm$ 0.21	88.10 $\pm$ 0.12	80.48 $\pm$ 0.10	90.15 $\pm$ 0.04
Carlini et al. [2022]	80.73 $\pm$ 0.16	87.69 $\pm$ 0.18	81.55 $\pm$ 0.13	88.89 $\pm$ 0.16

Table 3: Comparison of the proposed curvature score-based MIA with prior methods on CIFAR100 dataset with 1 and 2 augmentations. Results reported are the mean $\pm$ std obtained over 3 seeds.

Method	CIFAR10			
	1 Augmentation		2 Augmentations	
	Bal. Acc.	AUROC	Bal. Acc.	AUROC
Curv ZO NLL (Ours)	61.03 $\pm$ 0.84	67.56 $\pm$ 1.23	61.92 $\pm$ 0.87	68.82 $\pm$ 1.30
Curv ZO LR (Ours)	54.54 $\pm$ 0.19	58.06 $\pm$ 0.41	55.00 $\pm$ 0.17	58.89 $\pm$ 0.38
Carlini et al. [2022]	56.96 $\pm$ 0.27	60.05 $\pm$ 0.34	58.23 $\pm$ 0.29	61.73 $\pm$ 0.32

Table 4: Comparison of the proposed curvature score-based MIA with prior methods on CIFAR10 dataset with 1 and 2 augmentations. Results reported are the mean $\pm$ std obtained over 3 seeds.

A.8 Low FPR Results

In this section, we present the performance of curvature based MIA compared with prior works using TPR (true positive rate) at very low FPR percentages. The results are presented when using 1 augmentation (only the image) and 2 augmentations (image + mirror) in Table 5 and Figure 8. For these results, we use 64 shadow models for all the results in Table 5 (for methods that use shadow models). The results for Carlini et al. [2022] are a little lower than the one reported by Carlini et al. [2022] in Table I of their paper. However, note that the authors of Carlini et al. [2022] report the results with 256 shadow models while we report it using 64 shadow models. This highlights the benefits of curvature based approach, it achieves similar performance to Carlini et al. [2022] with 256 shadow models but with only 64 shadow models thus needs significantly less shadow models to achieve similar performance.

Takeaways: The parametric curvature LR method has the best TPR at very low FPR and makes efficient use of shadow models.

Method	1 Augmentation		2 Augmentations	
	TPR @ 0.1% FPR	TPR @ 0.01% FPR	TPR @ 0.1% FPR	TPR @ 0.01% FPR
Curv ZO LR (Ours)	21.07 $\pm$ 0.80	15.74 $\pm$ 3.12	23.92 $\pm$ 0.92	17.17 $\pm$ 1.87
Curv ZO NLL (Ours)	6.42 $\pm$ 0.62	0.05 $\pm$ 0.04	8.29 $\pm$ 1.52	0.10 $\pm$ 0.14
Carlini et al. [2022]	15.52 $\pm$ 0.54	5.02 $\pm$ 1.70	15.80 $\pm$ 0.49	6.56 $\pm$ 1.13
Sablayrolles et al. [2019]	10.50 $\pm$ 0.59	4.30 $\pm$ 0.53	10.50 $\pm$ 0.59	4.30 $\pm$ 0.53
Watson et al. [2022]	6.62 $\pm$ 0.23	2.95 $\pm$ 0.18	6.62 $\pm$ 0.23	2.95 $\pm$ 0.18
Ye et al. [2022]	6.63 $\pm$ 0.52	2.86 $\pm$ 0.42	6.63 $\pm$ 0.52	2.86 $\pm$ 0.42
Song and Mittal [2021]	1.21 $\pm$ 0.37	1.21 $\pm$ 0.37	1.21 $\pm$ 0.37	1.21 $\pm$ 0.37
Yeom et al. [2018]	0.07 $\pm$ 0.01	0.01 $\pm$ 0.00	0.07 $\pm$ 0.01	0.01 $\pm$ 0.00

Table 5: Comparison of TPR @ 0.1% FPR and TPR @ 0.01% FPR for CIFAR100 dataset with 1 and 2 augmentations. Results reported are the mean $\pm$ std obtained over 3 seeds and using 64 shadow models.

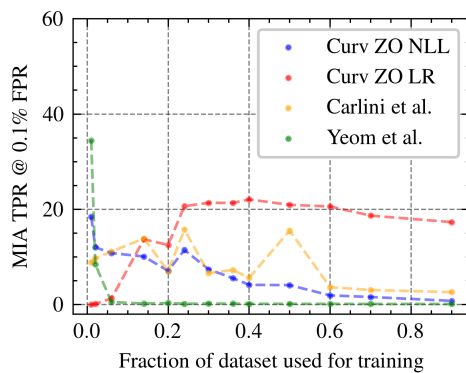


Figure 8: TPR at low FPR visualizing MIA performance as a function of the size of the train set, which is randomly sampled.

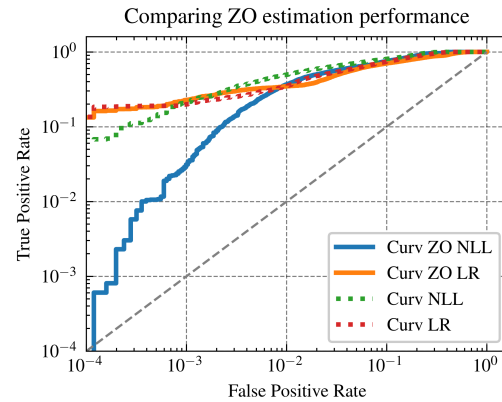


Figure 9: Low TPR performance comparison on CIFAR100 of ZO estimation (Curv ZO, black-box) vs. Hutchinson (Curv, white-box).

A.9 ZO Estimation Performance

In this section, we present the results of comparing the proposed zero-order curvature estimation with Hutchinson’s trace estimator-based curvature calculation. It is important to note that the trace estimator proposed by Garg et al. [2024] requires access to the model’s parameters, thus it is a white-box attack. In contrast, the proposed zero-order (ZO) estimation operates as a black-box attack.

The comparative results are presented in Table 6. Additionally, Figure 9 illustrates the True Positive Rate (TPR) vs. False Positive Rate (FPR) for the ZO estimation and Hutchinson’s trace estimator, utilizing a log-log plot to emphasize performance at low FPR.

Takeaways. From Table 6 we see that zero-order approach to curvature estimation results in notable differences in performance. Specifically, when examining AUROC and balanced accuracy, particularly on the CIFAR-10 dataset, a gap of several percentage points can be observed between the two methods. From Figure 9, we observe that the proposed LR technique performs similarly between Hutchinson’s trace estimator [Garg et al., 2024] and the proposed ZO estimation. However, the NLL method exhibits much greater sensitivity to estimation errors, as demonstrated by a significant performance drop at low FPR levels.

Dataset	Method	AUROC	Balanced Accuracy
CIFAR100	Curv ZO NLL	93.49 ± 0.18	84.47 ± 0.21
	Curv NLL	95.01 ± 0.19	85.66 ± 0.30
	Curv ZO LR	90.15 ± 0.04	80.48 ± 0.10
	Curv LR	92.87 ± 0.02	83.75 ± 0.16
CIFAR10	Curv ZO NLL	68.82 ± 1.30	61.92 ± 0.87
	Curv NLL	74.20 ± 1.42	64.96 ± 1.05
	Curv ZO LR	58.89 ± 0.38	55.00 ± 0.17
	Curv LR	68.53 ± 0.27	61.82 ± 0.11

Table 6: Performance Metrics for CIFAR100 and CIFAR10 of zero order estimation (Curv ZO, black-box) vs. Hutchinson trace estimation (Curv, white-box) [Garg et al., 2024].

A.10 Performance Across Architectures

In this section, we present the results of evaluating the proposed zero-order MIA on different neural network architectures. We selected four architectures: Inception [Szegedy et al., 2015] (specifically, the smaller Inception architecture used by Feldman and Zhang [2020]), VGG11 [Simonyan and Zisserman, 2014], VGG11 with batch normalization (denoted as VGG11BN), and finally, a ViT-B16 vision transformer [Vaswani et al., 2017, Dosovitskiy et al., 2021]. The results are provided in Table Table 7. Please note that for all of the results we used ResNet18 shadow models and 2 augmentations.

Takeaways. The proposed techniques outperform the previous state-of-the-art method by Carlini et al. [2022] across all architectures. Additionally, the lower performance on the Vision Transformer (ViT) model can be attributed to transfer learning. Since ViTs do not perform well when trained from scratch on CIFAR-100, we used an ImageNet-pretrained model, which was fine-tuned on CIFAR-100. This led to significant privacy benefits, which aligns with the findings of Mehta et al. [2022].

Arch.	Method	AUROC	Bal. Acc.	TPR @ 1% FPR
Inception	Curv ZO NLL	87.77 ± 0.61	79.04 ± 0.65	15.33 ± 1.38
	Curv ZO LR	81.94 ± 0.79	73.46 ± 0.88	19.46 ± 0.81
	Carlini et al. [2022]	83.77 ± 0.63	76.70 ± 0.66	16.66 ± 0.82
VGG11	Curv ZO NLL	73.68 ± 2.99	66.53 ± 3.35	6.19 ± 1.15
	Curv ZO LR	68.29 ± 1.38	64.30 ± 1.13	8.14 ± 1.85
	Carlini et al. [2022]	62.85 ± 0.71	61.08 ± 0.87	2.07 ± 0.19
VGG11BN	Curv ZO NLL	87.76 ± 0.17	79.12 ± 0.11	15.56 ± 1.34
	Curv ZO LR	76.17 ± 0.65	70.98 ± 0.26	22.80 ± 0.40
	Carlini et al. [2022]	70.57 ± 0.39	67.13 ± 0.43	3.22 ± 0.34
ViT-B16	Curv ZO NLL	58.10 ± 0.20	55.40 ± 0.18	1.60 ± 0.01
	Curv ZO LR	53.98 ± 0.30	52.55 ± 0.18	1.91 ± 0.25
	Carlini et al. [2022]	54.59 ± 0.05	53.22 ± 0.10	1.44 ± 0.14

Table 7: MIA performance metrics for different architectures and methods on CIFAR100 dataset.

A.11 Broader Impact

The development and analysis of input loss curvature in deep neural networks presented in this work have significant implications for both the academic and practical fields of machine learning and privacy. By introducing a novel black-box membership inference attack that leverages input loss curvature, this research advances the state-of-the-art in privacy testing for machine learning models.

From a societal perspective, the ability to detect and mitigate membership inference attacks is crucial for maintaining user privacy in machine learning applications. This work helps pave the way for more robust privacy-preserving techniques, ensuring that sensitive information is protected against unauthorized inference. Furthermore, the insights gained from the relationship between input loss curvature and memorization can guide the development of more secure machine learning models, which is particularly important as these models are increasingly deployed in sensitive domains such as healthcare, finance, and personal data processing.

Additionally, this research highlights the potential of using subsets of training data as a defense mechanism against membership inference attacks. By identifying that training on certain sized subsets can improve resistance to these attacks, this work offers practical guidance for model training practices that can enhance privacy without significantly compromising performance.

A.12 Limitations

In this paper we presented the a theoretical analysis and experimental evidence for improved MIA using input loss curvature. The theoretical and empirical evidence also shows that certain sized subsets of the training set may provide defense against membership inference attacks. However, as mentioned in the paper, the MIA performance improvement occurs only above a certain training dataset size below which the non-parametric model of ‘Curv NLL’ works better and is not explained by the theoretical analysis. Similar to techniques that use shadow models such as by Shokri et al. [2017], Carlini et al. [2022] we need to train shadow models, which can be computationally expensive. Further, the method described requires more queries at minimum $4\times$ more than Carlini et al. [2022]. We believe these limitations can be addressed by follow up research.

A.13 Reproducibility Details

In this section we present additional details for reproducing our results. Additionally we have released our code here <https://github.com/DeepakTatachar/Curvature-Clues>.

Training. For experiments that use private models, we use the Opacus library [Yousefpour et al., 2021] to train ResNet18 models for 20 epochs till the privacy budget is reached. We use DP-SGD [Abadi et al., 2016] with the maximum gradient norm set to 1.0 and privacy parameter $\delta = 1 \times 10^{-5}$. The initial learning rate was set to 0.001. The learning rate is decreased by 10 at epochs 12 and 16 with a batch size of 128.

For shadow model training on CIFAR10 and CIFAR100 we trained on 50% randomly sampled subset of the data for 300 epochs with a batch size of 512 for CIFAR100 and 256 for CIFAR10. We used SGD optimizer with the initial learning rate set to 0.1, weight decay of 1×10^{-4} and momemtum of 0.9. The learning rate was decayed by 0.1 at 180th and 240th epoch. For ImageNet we used pre-trained models from Feldman and Zhang [2020] as shadow models which were trained on a 70% subset of ImageNet. For both CIFAR10 and CIFAR100 datasets, we used the following sequence of data augmentations for training: resize (32×32), random crop, and random horizontal flip, this is followed by normalization.

Testing. During testing we used resize followed by normalization. We used two augmentations the original image and its mirror. The number of augmentations used are specified in the corresponding experiment section. When using pre-trained models from Feldman and Zhang [2020] we validated the accuracy of the models before performing experiments.

Compute Resources. All of the experiments were performed on a heterogeneous compute cluster consisting of 9 1080Ti's, 6 2080Ti's and 4 A40 NVIDIA GPUs, with a total of 100 CPU cores and a combined 1.2 TB of main system memory. However, the results can be replicated with a single GPU with 11GB of VRAM.

Hyperparameters. Our code uses 2 hyper parameters for zero-order input loss curvature estimation. The n_{iter} and h in Algorithm 1. We used $n_{iter} = 10$ and $h = 0.001$. To improve reproducibility, we have provided the code in the supplementary material.

A.14 Licenses for Assets Used

For each of the assets used we present the licenses below we also have provided the correct citation in the main paper as well as here for convenience.

1. ImageNet [Russakovsky et al., 2015]: Terms of access available at <https://image-net.org/download.php>
2. CIFAR10 [Krizhevsky et al., 2009]: Unknown / No license provided
3. CIFAR100 [Krizhevsky et al., 2009]: Unknown / No license provided
4. Pre-trained ImageNet models and code: We used pre-trained ImageNet models from Feldman and Zhang [2020] which is licensed under Apache 2.0 <https://github.com/google-research/heldout-influence-estimation/blob/master/LICENSE>.
5. Baseline methods: We re-implemented the baseline methods hence is provided with along with our code which is distributed under the MIT License.
6. Opacus [Yousefpour et al., 2021]: Licensed under Apache 2.0, <https://github.com/pytorch/opacus/blob/main/LICENSE>.
7. Pytorch [Ansel et al., 2024]: Custom BSD-style license available at <https://github.com/pytorch/pytorch/blob/main/LICENSE>.
8. ResNet Model Architecture [He et al., 2016]: MIT license available at <https://github.com/kuangliu/pytorch-cifar/blob/master/LICENSE>

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 clearly outlines the exploration of input loss curvature in deep neural networks and the development of a theoretical framework for train-test distinguishability. It introduces a novel black-box membership inference attack (MIA) utilizing input loss curvature and the use of subsets of training data as a mechanism to defend against shadow model based MIA. These claims are substantiated through theoretical analysis and experimental validation, as presented in the paper.

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 of the work in Appendix A.12, where we note the increased computation requirement imposed by zero-order curvature calculation and other limitations of the proposed methods.

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: All the assumptions are stated in the Theorems and discussed prior to their utilization in the main paper. The complete proof has also been provided in the Appendix of the paper, specifically Appendix A.3, A.5 and A.6.

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: We include all the training details for the models in Appendix A.13. We also provide the pseudo-code for the proposed technique in Appendix A.2 along with all the hyper parameter details (see Appendix A.13). Further, we also provide the code in the supplementary material.

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: The datasets used are publicly available and are the standard datasets used for vision tasks. Our code is included in the supplementary material and will be released publicly after the conference deadline. The code replicates the results of our experiments.

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, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: The training details are reported in Section 6.1 of the paper and more details are provided in Appendix A.13.

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: All reported results are obtained by averaging over 3 seeds and reported with mean and standard deviation values.

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: For the experiments, we used a heterogeneous cluster whose details are mentioned in Appendix A.13.

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: The paper conforms to the code of ethics, we discuss potential impacts of this research paper in Appendix A.11.

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: We provide a detailed discussion about the broader impact of the research presented in the paper in Appendix A.11.

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: [NA]

Justification: We do not release data or models thus this paper does not pose 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 provide the detailed list of assets used and the terms under which they were licensed in Appendix A.14. We have respected the licensing terms of the original owners.

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: [\[Yes\]](#)

Justification: We release the code for reproducing the results under the MIT license and details of its usage are documented in Appendix A.13 and the supplementary material README file.

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 or 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 or 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.