

Robust and Large-Payload DNN Watermarking via Fixed, Distribution-Optimized, Weights

Benedetta Tondi, *Member, IEEE*, Andrea Costanzo, Mauro Barni, *Fellow, IEEE*

Abstract—The design of an effective multi-bit watermarking algorithm hinges upon finding a good trade-off between the three fundamental requirements forming the watermarking trade-off triangle, namely, robustness against network modifications, payload, and unobtrusiveness, ensuring minimal impact on the performance of the watermarked network. In this paper, we first revisit the nature of the watermarking trade-off triangle for the DNN case, then we exploit our findings to propose a white-box, multi-bit watermarking method achieving very large payload and strong robustness against network modification. In the proposed system, the weights hosting the watermark are set prior to training, making sure that their amplitude is large enough to bear the target payload and survive network modifications, notably retraining, and are left unchanged throughout the training process. The distribution of the weights carrying the watermark is theoretically optimised to ensure the secrecy of the watermark and make sure that the watermarked weights are indistinguishable from the non-watermarked ones. The proposed method can achieve outstanding performance, with no significant impact on network accuracy, including robustness against network modifications, retraining and transfer learning, while ensuring a payload which is out of reach of state of the art methods achieving a lower - or at most comparable - robustness.

Index Terms—DNN IPR protection, Deep Learning security, DNN watermarking, White box watermarking, Robust DNN watermarking



1 INTRODUCTION

Neural network watermarking has been proposed as a solution to protect the Intellectual Property Rights (IPR) associated to DNN (Deep Neural Networks) models. In particular, the presence of a proprietary watermark within a DNN model has been proposed as a way to prove the ownership of the model, thus making it possible to trace possible abuses and resolve ownership disputes [1]. In addition to ownership verification, DNN watermarking can be used to trace back the history of DNN models, for fingerprinting and traitor tracing [2], to verify the integrity of a model [3], and even to determine the source of contents generated by generative models [4]. For most applications it is required that the watermark be robust against DNN modifications, like pruning, quantisation, fine tuning and even transfer learning.

As it is well known from previous studies on digital media watermarking [5], and as it has been restated recently in the framework of DNN watermarking [6], the mere presence of a watermark within a DNN model is not enough to doubtlessly prove the ownership of the model. Such a possibility, in fact, depends on the ownership verification protocol wherein the watermark is used, and the threat model associated to it (see for instance [1], [7]). In some cases, the watermark needs to be authenticated by a certification authority, while in other cases additional properties such as non-invertibility must be satisfied, or the possibility of embedding a minimum payload guaranteed. Similar arguments hold for applications other than ownership verification [8].

In order to investigate the very essence of the watermarking process, that is the process whereby a certain

information is indissolubly tied to a DNN model, here we avoid to link our study to a specific application, rather we consider watermarking as a primitive defined by its black-box characteristics and by certain general properties like robustness, payload and unobtrusiveness. The way the watermark is used within a protocol, like for instance an ownership verification protocol, and the possible additional requirements stemming from the protocol, are left to other researches.

In the above framework a first distinction must be made between zero-bit and multi-bit watermarking. With zero-bit watermarking, the watermark extractor, now called detector, is only asked to decide whether the inspected model contains a given known watermark or not. With multi-bit watermarking, instead, the watermark extractor, more properly referred to as decoder, extracts the watermark bits without knowing them in advance. In both cases, the extraction of the watermark requires the knowledge of a secret key¹, whose presence is crucial to ensure the secrecy of the watermark, and avoid that non-authorised users can access the watermarking channel².

A key-difference between multimedia and DNN watermarking derives from the observation that a DNN model is not a static object but a function defined by the way it maps the input samples into the output space [11]. In the following, we indicate such a mapping as $y = \Phi(x, \mathbf{w})$, where $x \in \chi$ indicates the input of the model, y the

1. In some zero-bit watermarking schemes, the key corresponds to the watermark itself, however, the watermark and the key play a different role and it is advisable to keep them separate.

2. For sake of simplicity, here we broadly define the watermarking channel as the mechanism whereby the watermark payload is associated to the host model. At the same time, access to the watermarking channel is defined as the possibility to write, read, or erase the watermark bits from the host network (see [9], [10] for a more rigorous and comprehensive description of these concepts).

corresponding output, and $\mathbf{w}$ the network weights³ (we sometimes refer to them as the network coefficients). In the simplest case (white-box watermarking), the watermark is extracted directly by looking at $\mathbf{w}$. In other cases, the model is accessible only in a black-box modality and hence the watermark is extracted by looking at the output of the network in correspondence to a set of specific inputs (black-box watermarking). In yet other cases, like for generative models, the output of the model has a large enough entropic content and hence the watermark can be retrieved from any output y even without knowing the corresponding input (box-free watermarking). Examples of white-box, black-box and free-box DNN watermarking algorithms are described, respectively, in [12], [13] and [14].

In this paper, we focus explicitly on white-box, multi-bit watermarking. In this case, the design of an effective watermarking algorithm boils down to finding a good trade-off between three general requirements (watermarking trade-off triangle), namely: *robustness* against network modifications, *payload*, measured as the number of bits the watermark consists of, and *unobtrusiveness*, that is the capacity of the watermarking technique to embed the desired payload without affecting the performance of the watermarked network. In particular, our work stems from the recognition of the peculiar relationship existing between the above requirements in the case of white-box DNN watermarking. We first argue that due to the particular way DNNs work, unobtrusiveness is linked *very weakly* to the robustness and payload requirements, unless an additional constraint regarding the secrecy of the watermark is considered. Failing to recognise the peculiarity of such a relationship results in watermarking schemes with suboptimal performance in terms of robustness, payload or both. Then, we exploit the above insight to propose a new white-box, multi-bit watermarking algorithm with large payload and improved robustness with respect to existing algorithms. More specifically, in our scheme the values of the DNN weights hosting the watermark are fixed prior to training, making sure that their amplitude is large enough to bear the target payload and survive retraining and other network modifications. The watermarked weights are then frozen and are not updated during the learning phase. In this way, the other weights are determined in a watermark-dependent way, co-operating with the fixed weights to accomplish the network task. To ensure the secrecy of the watermark and avoid that the watermarked weights are easily identifiable, the distribution of the watermarked coefficients is optimised by minimizing the Kullback-Leibler (KL) distance between the watermarked and non-watermarked weights for a given amplitude of the watermark. As done in other systems, robustness is further enhanced by spreading the watermark bits across several host coefficients.

We verified the effectiveness of the proposed watermarking algorithm by considering different architectures and classification tasks addressing different application domains, namely, image classification and image manipulation detection. The results we got show that the proposed

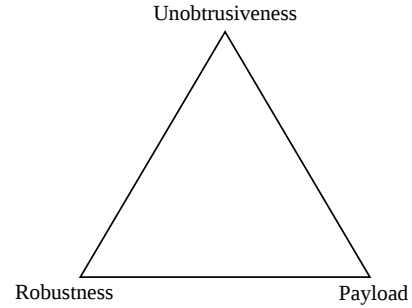


Fig. 1: The watermarking trade-off triangle.

method can achieve good performance even for large payloads, with negligible impact on the accuracy of the underlying classification task. We also verified the strong robustness of the watermark against the most common network modifications, including, pruning, weights quantisation, and, most noticeably, retraining. Moreover, thanks to the optimization of the distribution of the watermarked weights, the marked coefficients are indistinguishable from the non-marked ones.

The rest of this paper is organised as follows: in Section 2, we revisit the watermarking trade-off triangle and lay the basis of the new watermarking algorithm proposed in this paper. We also give a brief review of existing schemes by the light of the previous analysis. In Section 3, we introduce the notation used throughout the paper and state the requirements to be satisfied by the watermark. In Section 4, we describe the proposed watermark embedding and extraction algorithms. Section 5 and 6 are devoted, respectively, to the description of the experimental setting and to the discussion of the experimental results. Finally, in Section 7, we draw our conclusions and highlight directions for future work.

2 THE WATERMARKING TRADE-OFF TRIANGLE RE-VISITED

In classical watermarking theory [15], [16], the main goal of a multi-bit watermarking algorithm is to find a good, possibly optimum, trade-off between three opposite requirements summarised by the watermarking trade-off triangle shown in Figure 1.

Payload and robustness have a similar meaning for media and DNN watermarking, even if the kind of manipulations the watermark should be robust to is completely different in the two cases. In media watermarking, common manipulations include lossy compression, filtering, resampling, cropping etc, while in the DNN case, the manipulations we are interested in are pruning, quantisation and, most noticeably, fine tuning and transfer learning. The unobtrusiveness requirement, however, plays a completely different role in the two cases. In media watermarking, unobtrusiveness requires that the *perceptual* quality of the watermarked media is not degraded due to the presence of the watermark. This, for instance, means that the watermark should be *invisible* in the case of image watermarking, and *inaudible* in the case of audio watermarking. In media watermarking, then, unobtrusiveness has a direct impact on the amplitude or the *strength* of the watermark, usually measured as the

3. In general $\mathbf{w}$ includes also the network biases, however, in our scheme the watermark is embedded only in the weights so, without loss of generality, we will refer to $\mathbf{w}$ as the weights of the network.

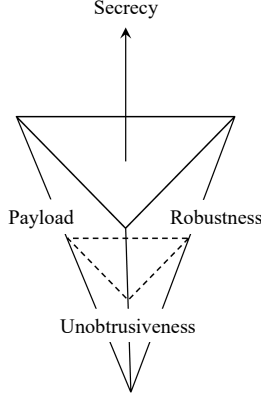


Fig. 2: DNN watermarking trade-off tetrahedron.

distance (possibly the perceptual distance) between the original content and the watermarked one. In turn, this means that increasing the amplitude of the watermark to improve its robustness or increase its payload is possible only up to certain extent. In (white-box) DNN watermarking, the situation is completely different. Here the unobtrusiveness of the watermark requires that its presence has a negligible, if any, impact on the performance achieved by the network with regard to its intended task, hereafter referred to as primary task. This requirement is not linked directly to the amplitude of the watermark, even because in the DNN case an original set of non-watermarked coefficients to measure the distance from does not exist. In addition, the number of coefficients defining a network is usually highly redundant, thus making it possible, for instance, to prune many of them without any noticeable effect on the performance of the network. Eventually, the saturation effect enforced by the most commonly used activation functions limits the impact of very large intermediate values on the final output of the network. The striking conclusion of these observations is that, in principle, one could embed a very large payload by encoding the watermark in the amplitude of very few, in the limit even only one, extremely large coefficients, and freezing them during training. As with any communication channel [17], the absence of constraints on the amplitude, or to better say the power, of the coefficients bearing the watermark, results in unbounded capacity.

The obvious drawback of the approach outlined above is that the watermarked coefficients would be easily detectable, thus compromising irremediably the secrecy of the watermark. A possible way to alleviate this problem is to spread the watermark over many coefficients of limited amplitude, however, the number of coefficients that can be frozen during the training process is bounded by the unobtrusiveness constraint, since freezing a too large fraction of coefficients may degrade the performance of the network on the primary task. It is the presence of the secrecy requirement, then, that creates the trade-off between unobtrusiveness and the other two corners of the triangle, namely robustness and payload. With the above ideas in mind, the trade-off between the various requirements of DNN watermarking is better illustrated by a reverse tetrahedron rather than a triangle (see Figure 2). When the secrecy requirement is relaxed, close to the vertex of the tetrahedron,

the corners of the watermarking trade-off triangle get closer, thus making it easier to satisfy all of them simultaneously. The opposite situation occurs when secrecy is given more importance, in which case the trade-off triangle gets larger, thus creating a stronger tension between its corners.

In view of the previous ideas, it is evident that to design a robust and high payload multi-bit DNN watermarking algorithm, it is advisable to encode the watermark into the amplitude of large coefficients by paying attention to evaluate the impact that the amplitude of the watermarked coefficients has on watermark secrecy.

2.1 Prior art

On the basis of the framework provided by the previous discussion and summarised in Figure 2, it is easy to realize that most white-box watermarking techniques proposed so far fail to address the watermarking trade-off properly. According to the paradigm adopted by most of the methods proposed so far, in fact, watermarking is achieved by adding a watermarking loss term to the loss function used during training. Then, watermarking is carried out simultaneously with training, by letting the weights of the network be the solution of the following optimization problem:

$$\mathbf{w} = \arg \min(\mathcal{L} + \lambda \mathcal{L}_w), \quad (1)$$

where $\mathcal{L}$ is the primary loss, whose minimization ensures that the performance of the network are good, and $\mathcal{L}_w$ guarantees that the watermark can be extracted without (too many) errors. $\mathcal{L}_w$ is usually related to the error probability of the watermark decoder, which, in most cases is designed based on a spread spectrum paradigm. In [18], for instance, and many subsequent works inspired to it, e.g., [2], [13], [19], [20], [21], the watermark bits are encoded in the sign of the projection of the marked coefficients onto a set of n pseudo-random sequences. The watermark loss $\mathcal{L}_w$, then, corresponds to the cross-entropy between the true watermark bits and the soft bit estimates obtained by applying a sigmoid function to the projections. In the above setting, the tradeoff between payload and robustness is achieved by using longer spreading sequences, and, most commonly, by adjusting the weighting parameter λ . It is clear, however, that increasing λ ensures only that the cross-entropy term decreases, at the price of a stronger impact on the network performance, without any guarantee that this increases the robustness of the watermark.

A noticeable exception is the zero-bit watermarking algorithm described in [22], where, similarly to our scheme, the weights hosting the watermark are fixed before training and are not modified during the training process. The strategy adopted to enhance robustness, however, does not rely on the amplitude of the marked coefficients. Rather, a properly designed loss function explicitly making the watermarked weights more sensitive to loss variations with respect to non-watermarked weights, is adopted, thus increasing the robustness of the watermark. In this way [22] can achieve a remarkable robustness against fine-tuning and weights quantization, however it still lacks robustness against transfer learning. In addition, no attention is given to the payload requirement, given that the system proposed in [22] belongs to the class of zero-bit watermarking methods.

A DNN watermarking algorithm that in some way achieves robustness by associating the watermark to large DNN coefficients is the one described in [20]. In this method, the watermark is still embedded by properly including a watermark loss term. However, the message bits are weighted by so called greedy residuals. Because of the way such greedy residuals are constructed, the watermark is indirectly embedded in the largest weights of the network, whose amplitude is increased in order to get large residuals by multiplying the message bits with the correct sign. In addition, to facilitate obtaining large amplitudes, the watermark is embedded in the first layer of the network, where the weights tend to be larger. No particular attention is paid to secrecy, and hence the watermarked weights are easily identifiable, as discussed in Section 6.3.

In this work, we propose a simple but effective way to exploit the large amplitude of the watermarked coefficients to simultaneously achieve robustness and high payload. In particular, the watermark is spread over a number of large-amplitude, fixed coefficients (like in [22]), while at the same time ensuring that the marked coefficients are as indistinguishable as possible from the non-marked ones. This provides a way to simultaneously trade-off between the 4 corners of the tetrahedron illustrated in Figure 2. The experimental results shown in Section 6, prove that in this way our system can achieve a remarkable robustness against the most common network manipulations, including retraining for transfer learning, ensuring a payload which is out of reach of state of the art methods with comparable, or even lower, robustness.

3 NOTATION AND PROBLEM DEFINITION

Let Φ indicate a generic non-protected DNN model, that is, a model trained for a given task without the watermark. The notation Φ^m is used to denote a watermarked model carrying out the same task. We denote with $\mathcal{L}$ the primary loss function of the network, usually corresponding to cross-entropy. The average loss measured across all the samples of the labeled training set $(\mathcal{X}, \mathcal{Y})$ is compactly denoted with $\mathcal{L}(\mathcal{X}, \mathcal{Y}, \Phi)$ ($\mathcal{L}(\mathcal{X}, \mathcal{Y}, \Phi^m)$, for the watermarked model).

We indicate with $\mathbf{b} \in \{0, 1\}^l$ the vector with the watermark bits, (sometimes referred to as the watermark message), where l denotes the number of bits the watermark consists of.

The network weights can be generically represented as tensors. For each convolutional layer j , the dimensionality of the tensor with the weights is determined by the kernel size of the filters, the depth of the input and the number of filters (3-dim tensor). For notational simplicity, we flatten the weight tensor into a vector $\mathbf{w}_j$ (row vector) containing all the weights of layer j . We generically indicate the vector with all the weights of the network (be them watermarked or not) as $\mathbf{w}$. With the above notation, embedding the watermark bits into the weights of the network corresponds to embedding the vector $\mathbf{b}$ into $\mathbf{w}$. The final watermarked vector should be such that the network preserves its functionality.

We denote with Ω , hereafter called *host* index set, the set with the indexes of the weights hosting the watermark, that is, the positions in $\mathbf{w}$ that are selected to carry the

watermark message. The cardinality of Ω is indicated by n (usually $n \geq l$). With this notation, the vector with the watermarked weights should be indicated as $\mathbf{w}^m = (w_{i_1}, w_{i_2}, \dots, w_{i_n}), i_k \in \Omega$, however, for sake of simplicity, we will indicate it as $\mathbf{w}^m = (w_1^m, w_2^m, \dots, w_n^m)$. We also introduce the vector with non-watermarked weights, indicated by $\bar{\mathbf{w}}^m$. With the notation introduced above, we have $\bar{\mathbf{w}}^m = (w_{j_1}, w_{j_2}, \dots, w_{j_n}), j_k \notin \Omega$, more simply written as $\bar{\mathbf{w}}^m = (\bar{w}_1^m, \bar{w}_2^m, \dots, \bar{w}_{N-n}^m)$, where N is the total number of coefficients the model consists of. We observe that the host indexes may be selected from only one layer (single-layer embedding), or multiple layers (multi-layer embedding). We find useful to denote with Ω_k the set of watermark indexes in layer k . By indicating with N_k the number of weights in layer k , then, the percentage of watermarked weights in the network, denoted with p^m , is equal to $p^m = |\Omega|/N \times 100$, while the percentage of watermarked weights in layer k is $p_k^m = |\Omega_k|/N_k \times 100$. We let $f_{\mathbf{w}^m}$ and $f_{\bar{\mathbf{w}}^m}$ denote the marginal empirical distribution of the watermarked and non-watermarked weights, estimated, respectively, on $\mathbf{w}^m$ and $\bar{\mathbf{w}}^m$, that is, the distributions induced by the respective sequences, and let $E_{\mathbf{w}^m}$ (res. $E_{\bar{\mathbf{w}}^m}$) denote the empirical expectation computed over $f_{\mathbf{w}^m}$ (res. $f_{\bar{\mathbf{w}}^m}$). Finally, we denote with $\mathcal{D}$ the Kullback-Leibler (KL) distance between two continuous probability distributions f and g [17], defined as $\mathcal{D}(f||g) = \int_{x \in \mathbb{R}} f(x) \log(f(x)/g(x)) dx$ ⁴.

3.1 Watermarking model and requirements

According to the multi-bit approach adopted in this paper, the watermarker wants to embed a watermark message $\mathbf{b}$ into the weights of the model Φ^m . To do so, he relies on a secret key $K = \{\Omega, \mathbf{s}\}$, consisting of the information on the host layers and weights (Ω), and the spreading sequence $\mathbf{s}$ used to spread the message $\mathbf{b}$ over the host coefficients. In addition to the usual unobtrusiveness, robustness and payload requirements, the watermark must satisfy the *integrity* requirement, asking that in the absence of model modifications, the decoded watermark $\hat{\mathbf{b}}$ should be close (ideally equal) to $\mathbf{b}$, that is, in the absence of modifications the bit error rate should be small (ideally zero).

With regard to *secrecy*, we loosely define it as the impossibility for an attacker to estimate the secret key by observing the weights of the watermarked model. In particular, here we focus on the estimation Ω . Estimating Ω , in fact, would allow the attacker to erase the watermark or overwrite it (while to decode it the knowledge of $\mathbf{s}$ is also necessary).

For the robustness requirement, we consider retraining for fine tuning and in particular transfer learning, and model compression, namely, parameter pruning and weights quantization.

Retraining is a typical modification applied to models. A trained model is further trained for some epochs on the same or a different task. We speak about *fine-tuning* when the network model is retrained to solve the same task in the same domain (often, on the same dataset used for the original training or on a subset of it), for some additional epochs, possibly with a different learning rate. We speak

4. For discrete distributions the integral is replaced by a summation over the alphabet of x .

about *transfer learning* in a more general setting where knowledge is transferred across domains or tasks. Transfer learning is widely adopted in practice since it turns out that transferring the knowledge from different tasks is less computationally expensive than training a new model from scratch. By referring to the categorization adopted in [23], in the *transductive transfer learning* setting, the source and target tasks are the same, while the source and target domains are different. In the *inductive transfer learning* setting, the target task is different from the source task, while the domain can be the same or not. In this case, the model trained on the source task is used as a pre-trained solution to train a network on a new task, possibly related to the original one. Retraining alters the weights of the model, the degree of alteration depending on the difference between the domain and task targeted by the retraining process and those of the original model, and on the number of retraining epochs. Arguably, during transfer learning, the weights are modified by a larger extent than during fine-tuning, thus achieving robustness in the former scenario can be much harder.

Compression. DNN models are often squeezed to deploy them into low power or computationally weak devices, and to reduce the storage demand. Typical methods for model compression are *network pruning* and *weight quantization*. The former cuts off network weights, whose value is smaller than a threshold, the latter reduces the numerical precision of the model coefficients, converting them from a floating point representation to a lower-precision, typically integer, representation. More specifically, model quantization can be described as follows:

$$\mathbf{w}_q = \left\lfloor \frac{\mathbf{w}}{\delta} \right\rfloor \times \delta, \quad \delta = 2w_{max}/2^{n_b}, \quad (2)$$

where $\mathbf{w}_q$ indicates the quantized weights, w_{max} represents the maximum absolute value assumed by the weights, and n_b indicates the number of bits required to represent the quantized values.

We observe that network retraining represents the most critical case of network re-use for existing DNN watermarking schemes. Although some DNN watermarking methods have been proposed that can achieve a certain degree of robustness against fine-tuning [22], to the best of our knowledge, no method can simultaneously achieve high-payload and robustness in a transfer learning scenario.

4 THE PROPOSED DNN WATERMARKING METHOD

By relying on the insights we got by revisiting the watermarking trade-off triangle, our approach to simultaneously achieve a high payload and robustness against network modifications consists in embedding the watermark into several large coefficients and freeze them during the training process. In doing so we must i) ensure that the presence of several, fixed, large coefficients does not prevent training the network (unobtrusiveness), ii) avoid that the large coefficients hosting the watermark are easily identified and attacked (secrecy). In particular, the proposed watermark embedding algorithm works as follows:

- The information message $\mathbf{b}$ is embedded before training and the weights hosting the watermark message are not updated during training.

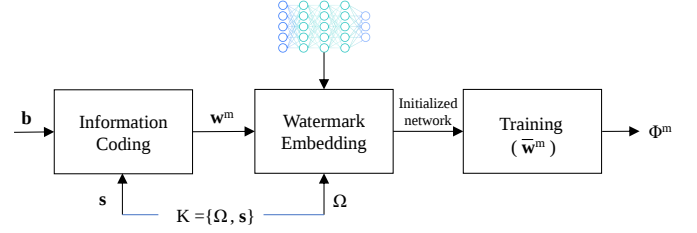


Fig. 3: Watermark embedding procedure.

- The message bits are encoded in $\mathbf{w}^m$ via direct sequence spread spectrum watermarking [15]. As a consequence, the message $\mathbf{b}$ is spread over several host weights. Specifically, the l bits of the message are used to modulate a pseudo-random sequence $\mathbf{s}$ of length $n \geq l$, where $S = n/l$ is the spreading factor. The strength of the watermark is controlled by means of a parameter γ .
- To ensure that the watermarked weights are indistinguishable from the other weights, the distribution of the spreading sequence is chosen in such a way to be as close as possible to that of the non-watermarked weights. In particular, for a given γ , the optimum distribution of the marked coefficients is derived theoretically by minimizing the KL distance between the distribution of the watermarked and non-watermarked coefficients (see Section 4.2).

For a given payload of l bits, the trade-off between unobtrusiveness and robustness is determined by the spreading factor S , while γ controls the robustness and the secrecy of the watermark, with a too large γ resulting in an easy-to-detect watermark.

The main steps of the watermark embedding algorithm are illustrated in Fig. 3 and described in the following.

4.1 Watermark embedding

Let $\mathbf{u} \in \{-1, +1\}^l$ be the antipodal sequence associated to $\mathbf{b}$, that is $u_i = 1$ (res. -1) if $b_i = 1$ (res. 0). Before embedding, the vector $\mathbf{u}$ is used to modulate a spread-spectrum sequence $\mathbf{s}$. In particular, each u_i is used to modulate a block of coefficients of $\mathbf{s}$ (direct sequence spread spectrum [15]).

Formally, given a spreading factor S , we generate a secret pseudo-random sequence $\mathbf{s} \in \mathbb{R}^n$ of length $n = S \cdot l$ (the generation of $\mathbf{s}$ is discussed in Section 4.2). Then, for each $i \in [1, l]$, u_i is used to modulate the sign of the elements of $\mathbf{s}$ corresponding to the indexes from $(i-1)S + 1$ to iS . Specifically, the vector of watermarked coefficients $\mathbf{w}^m$, at the output of the information coding block in Fig. 3, is obtained as follows:

$$w_j^m = u_i \cdot s_j, \quad j \in \{(i-1)S + 1, \dots, iS\}, \quad i \in [1, l]. \quad (3)$$

Then, $(w_{(i-1)S+1}^m, \dots, w_{iS}^m)$ is the vector of the weights associated to b_i . In the embedding phase, the network weights in the positions indicated by Ω are initialized to the values given by $\mathbf{w}^m$ (watermark embedding block in Fig. 3). With regard to the coefficients hosting the watermark, they are chosen at random according to the secret key K . In addition to security, a further advantage of choosing the positions of the watermarked weights randomly is that the weights associated to the same bits are more likely to

undergo independent changes during subsequent network modifications.

The network, then, is trained as follows: the network weights are initialized randomly, thus getting the initial vector of the weights $\mathbf{w}^{(0)}$. The weights in the positions indicated by Ω are set to the values given by $\mathbf{w}^m$. At this point, the network is trained as usual, by minimizing the loss across all the samples of the training set. The watermarked weights $\mathbf{w}^m$ are frozen during the training process, and only the non-watermarked weights $\bar{\mathbf{w}}^m$ are updated through backpropagation.

As we said, we must ensure that the watermarked weights are as indistinguishable as possible from the non-watermarked ones, in order to avoid that they can be easily identified and attacked. To this purpose, the distribution of $\mathbf{s}$, and, as a consequence, the distribution of $\mathbf{w}^m$, is chosen in such a way to minimize the distinguishability between watermarked and non-watermarked weights, as described in the following section.

4.2 Optimization of watermarked weights distribution

In this section, we address the problem of finding the optimum distribution to generate the pseudo-random spreading sequence $\mathbf{s}$. Although our final goal is to optimize the distribution of $\mathbf{s}$ (and as a consequence the distribution of $\mathbf{w}^m$), we find convenient to initially formulate the optimization problem as a minimization carried out over the weights $\mathbf{w}^m$. Then, we will see that, under some mild assumptions, we can rephrase the optimization over the weights as an optimization over the weights distribution $f_{\mathbf{w}^m}$.

The problem of training a DNN watermarked model can be formalized as follows⁵:

$$\min_{\mathbf{w}^m} \mathcal{L}(\mathcal{X}, \mathcal{Y}, \bar{\mathbf{w}}^m \cup \mathbf{w}^{m,*}) \quad (4)$$

where, in order to minimize the distinguishability between watermarked and non-watermarked weights, $\mathbf{w}^{m,*}$ is obtained by solving the following optimization:

$$\mathbf{w}^{m,*} = \arg \min_{\mathbf{w}^m} \mathcal{D}(f_{\mathbf{w}^m} || f_{\bar{\mathbf{w}}^m}), \quad (5)$$

which also depends on $\bar{\mathbf{w}}^m$, with the additional constrain that

$$E_{\mathbf{w}^m}[w|w > 0] = \gamma, \quad E_{\mathbf{w}^m}[w|w < 0] = -\gamma, \quad (6)$$

to control the average amplitude of the weights bearing the watermark.

In the following, we require that the watermarked weights distribution is symmetric (given that the distribution of the non-watermarked weights is typically symmetric, this goes w.l.o.g.). With this assumption, in fact, the watermarked weights w_j^m obtained via Eq. (3) follows the same distribution of the pseudo-random sequence $\mathbf{s}$, regardless of $\mathbf{u}$. Because of the symmetry of the distribution, we can rewrite equation (6) as $E_{\mathbf{w}^m}[|w|] = \gamma$. A large γ results in a strong - hence expectedly more robust - watermark. In addition to robustness, γ affects the secrecy of the watermark. A very large value of γ , i.e., such that $E_{\bar{\mathbf{w}}^m}[|w|] \ll \gamma$, would,

5. Rigorously speaking, the set of trainable parameters also includes the biases, however, for sake of simplicity, we refer only to the weights since the watermark is not embedded in the biases.

in fact, make the presence of the watermark easily detectable upon inspection of the values of the weights.

To go on, we observe that (4) and (5) are entangled equations due to the presence of $\mathbf{w}^m$ and $\bar{\mathbf{w}}^m$ in both of them, so they can not be solved easily. However, under the assumption that the presence of the watermark does not affect the statistical distribution of the non-watermarked weights⁶, the two minimizations can be separated as follows:

$$\begin{aligned} \mathbf{w}^{m,*} &= \arg \min_{\mathbf{w}^m} \mathcal{D}(f_{\mathbf{w}^m} || \tilde{f}_{\bar{\mathbf{w}}^m}) \\ E_{\mathbf{w}^m}[|w|] &= \gamma, \end{aligned} \quad (7)$$

$$\bar{\mathbf{w}}^{m,*} = \arg \min_{\bar{\mathbf{w}}^m} \mathcal{L}(\mathcal{X}, \mathcal{Y}, \bar{\mathbf{w}}^m \cup \mathbf{w}^{m,*}), \quad (8)$$

where $\tilde{f}_{\bar{\mathbf{w}}^m}$ is the distribution of the non-watermarked weights of a non-watermarked model Φ solving the same task. To a closer look, we observe that the minimization in (7) depends only on the probability density function of $\bar{\mathbf{w}}^m$ rather than on the specific sequence, so it can be conveniently reformulated as a minimization over the distribution $f_{\bar{\mathbf{w}}^m}$:

$$\begin{aligned} f_{\mathbf{w}^m}^* &= \arg \min_{f_{\mathbf{w}^m}} \mathcal{D}(f_{\mathbf{w}^m} || \tilde{f}_{\bar{\mathbf{w}}^m}) \\ E[|w|] &= \gamma. \end{aligned} \quad (9)$$

Then, the sequence $\mathbf{w}^{m,*}$ appearing in (7) can be *any* sequence generated according to $f_{\mathbf{w}^m}^*$, thus allowing to keep the specific sequence $\mathbf{w}^{m,*}$ secret.

Solving the problem in (9) for a general distribution of the non-watermarked weights $\bar{\mathbf{w}}^m$ is not easy. In the following, we solve it by assuming that the distribution of the non-watermarked weights can be approximated by a Laplacian distribution (our experiments on several network architectures and tasks confirmed the goodness of this assumption, see Section 6.1.1). Let, then, $f_{\bar{\mathbf{w}}^m}(w) = \frac{1}{2\lambda} e^{-|w|/\lambda}$, that is $\bar{w}_i^m \sim \text{Laplace}(0, \lambda)$. Under such an assumption, problem (9) is solvable in closed form, as stated in the following theorem.

Theorem 1. *Let $\mathcal{F}$ denote the set of symmetric distributions. The minimization problem*

$$\begin{aligned} \min_{f_w \in \mathcal{F}} \mathcal{D}\left(f_w || \frac{1}{2\lambda} e^{-|w|/\lambda}\right), \\ \text{s.t. } E[|w|] = \gamma \end{aligned} \quad (10)$$

is equivalent to the problem of finding the maximum entropy distribution over all the symmetric probability density functions f_w satisfying $E[|w|] = \gamma$, whose solution is $f_w^ = \frac{1}{2\gamma} e^{-|w|/\gamma}$. Hence the optimum distribution for the watermarked weights is a Laplacian distribution with 0 mean and scale parameter γ , that is $w_i \sim \text{Laplace}(0, \gamma)$, for $i \in \Omega$.*

6. We checked experimentally that the distribution of the weights corresponding to the non-host indexes of the non-watermarked model and that of the watermarked model are approximately the same, for all the watermark settings considered in this paper. More details are reported in Section 6.1.1.

Proof. By observing that

$$\begin{aligned} \mathcal{D}(f_w || \frac{1}{2\lambda} e^{-|w|/\lambda}) &= \\ \int_w f_w(w) \left[\log f_w(w) - \log \frac{1}{2\lambda} e^{-|w|/\lambda} \right] dw &= \\ \int_w f_w(w) \log f_w(w) dw - \log \frac{1}{2\lambda} + \frac{1}{\lambda} \log e \cdot E[|w|], \end{aligned} \quad (11)$$

and given that $E[|w|] = \gamma$, eq. (10) can be rephrased as

$$\max_{f_w: E[|w|] = \gamma} h(f_w), \quad (12)$$

where $h(f) = -\int_x f(x) \log f(x) dx$ denotes the differential entropy. Equation (12) is equivalent to finding the *maximum entropy distribution* over all the probability density functions for which $E[|w|] = \gamma$ [17].

By rewriting the differential entropy integrating over the positive and negative supports separately, and by exploiting the symmetry of f_w , we get

$$\begin{aligned} h(f_w) &= - \int_{\mathbb{R}} f_w(w) \log f_w(w) dw = \\ &= - 2 \int_0^\infty f_w(w) \log f_w(w) dw \\ &= - \int_0^\infty f'_w(w) \log f'_w(w) dw + 1 \\ &= h(f'_w) + 1, \end{aligned} \quad (13)$$

where in the second-to-last equality we let:

$$f'_w(w) = \begin{cases} 2f_w(w) & w \geq 0 \\ 0 & w < 0 \end{cases}. \quad (14)$$

Therefore, in order to find the entropy-maximizing distribution in (12), we can equivalently solve the problem

$$\begin{aligned} \max_{f'_w} h(f'_w) \\ E[w] = \gamma \\ f'_w(w) = 0, \text{ for } w < 0. \end{aligned} \quad (15)$$

The distribution we search for in (15) is the maximum entropy distribution over all the probability density functions with support $[0, \infty)$ satisfying $E[w] = \gamma$. The solution of this problem is known ([17], Ex 12.2.5) and corresponds to $f'_w = \frac{1}{\gamma} e^{-|w|/\gamma}$, $w \geq 0$, hence yielding:

$$f_w^*(w) = \frac{1}{2\gamma} e^{-|w|/\gamma}, -\infty \leq w \leq \infty. \quad (16)$$

□

Based on the result of Theorem 1, the watermarked weights must be generated following the distribution $f_{w^m}^* = \frac{1}{2\gamma} e^{-|w|/\gamma}$, that is, $\mathbf{w}^{m,*} \sim \text{Laplace}(0, \gamma)$. To do so, we proceed as follows: we first generate the pseudo-random sequence $\mathbf{s}$ according to a $\text{Laplace}(0, \gamma)$ distribution. Then, we apply Eq. (3) to modulate the spreading sequence according to the watermark message, thus getting the vector of watermarked weights $\mathbf{w}^{m,*}$ still following a $\text{Laplace}(0, \gamma)$ distribution.

4.3 Watermark extraction

Watermark retrieval requires the knowledge of the indexes of the watermarked weights and the pseudo-random sequence $\mathbf{s}$. The vector $\mathbf{w}^m$ is first obtained by reading the weights of Φ^m in the positions indicated by Ω , then the i -th bit of the watermark is extracted as follows:

$$\hat{b}_i = \begin{cases} 1 & \text{if } \sum_{j=(i-1)S+1}^{iS} s_j \cdot w_j^m \geq 0 \\ 0 & \text{otherwise} \end{cases}. \quad (17)$$

The Bit Error Rate (BER) is calculated as $\text{BER} = (\sum_i (\mathbf{b} \oplus \hat{\mathbf{b}})_i / l) \times 100$, where $\oplus$ denotes the bitwise XOR operation.

We notice that, in the absence of modifications of the watermarked model, $\hat{\mathbf{b}}$ is equal to $\mathbf{b}$ by construction. In classical watermarking theory, such a property is achieved by informed watermarking algorithms [15], whereby embedding is performed by applying a signal-dependent perturbation to the host signal. Our DNN watermarking algorithm adopts a somewhat dual approach. The watermarked weights are first fixed, then the other weights are defined in a watermark-dependent way, in such a way that they *co-operate* with the fixed weights to accomplish the network task.

5 EXPERIMENTAL METHODOLOGY

We validated the proposed DNN watermarking technique by considering different architectures and tasks.

5.1 Host networks and tasks

We focused on tasks taken from two different application areas, that is, image forensics and pattern recognition. Specifically, we considered the distinction of images generated by Generative Adversarial Networks (GANs) from natural images [24] (GAN detection), object recognition (CIFAR-10 and 100 classification [25]) and traffic sign classification (GTSRB classification [26]). These tasks permit to test the effectiveness of our method in different scenarios, including binary classification and multi-class classification. For each task, we chose network architectures among those achieving state-of-the-art performance.

For the GAN detection task, we considered the discrimination of natural face images and images generated by the StyleGAN2 model [27], by means of XceptionNet, which can achieve state-of-the-art performance [28]. We trained the network for 10 epochs with SGD optimizer, learning rate 0.01 and batch size 32. More details on network training and the dataset can be found in the authors' repository [29].

For CIFAR-10 classification, we trained a ResNet18 and a DenseNet169 architecture [30], following the parameter setting reported in [31], achieving the benchmark accuracy for this task. Specifically, with both DenseNet and ResNet architecture, the network was trained for 200 epochs with SGD optimizer, learning rate 0.01 with multi-step decay every 50 epochs and batch size 32. In all our experiments, we consider augmentation. CIFAR-100 and GTSRB are considered for the transfer-learning experiments.

Notably, the architectures we used to assess the effectiveness of the algorithm are quite diverse. They all have a similar number of parameters (in the order of 10^7), however,

they differ in terms of depth and internal structure connections. While the block connections in the XceptionNet architecture are pretty standard, in DenseNet, there are dense blocks where each convolutional layer is connected to every other layer in a feed-forward fashion. ResNet instead relies on residual blocks for feature extraction.

The DNN watermarking algorithm and network training have been implemented by using the PyTorch 1.8 library and the code is made publicly available for reproducibility (<https://github.com/andreacos/Deep-Neural-Networks-Robust-Watermark>).

5.2 Watermarking algorithm setting

We tested the performance of our watermarking algorithm by considering both single-layer and multi-layer embedding. The multi-layer solution is necessary for large payloads and/or spreading factors, that is, for large n , when the percentage of watermarked weights in the watermarked layer for the single-layer embedding is too large, thus deteriorating the network performance (as confirmed by our experiments).

In all the settings, we embedded the watermark in the convolutional layers from intermediate to deep. All the information on the host layers k , the total number of weights in the layer (N_k), and the variance of the distribution of the non-watermarked weights for each layer (σ_k^2), can be found at the link <https://github.com/andreacos/Deep-Neural-Networks-Robust-Watermark>.

We observe the following noticeable differences among the architectures we have considered: regardless of the primary task, the variance of the distribution of the weights of DenseNet and ResNet is much lower than the variance of the weights of XceptionNet. More precisely, the variance of non-watermarked weights for XceptionNet is in the range $[0.4-2.5]$, while it is in the range $[4 \cdot 10^{-6}-6 \cdot 10^{-6}]$ for DenseNet and $[2 \cdot 10^{-5}-1 \cdot 10^{-4}]$ for ResNet.

In the multi-layer case, for simplicity, we watermarked the same number of weights in each layer, that is $|\Omega_k|$ is the same for all k 's. Given that the host layers have different number of parameters, the percentages of watermarked weights p_k^m in different layers are different. Such percentages obviously depend on the specific watermarking setting, i.e., on the payload l and the spreading factor S .

For a given layer k , the watermark is embedded by considering different strengths γ proportional to the standard variation of the distribution of the non-watermarked weights of the layer. More precisely, in the experiments, we let $\gamma = C\sigma_k/\sqrt{2}$ and adjusted the watermark strength by varying the parameter C . We observe that, with this definition of γ , the variance of the distribution of the watermarked weights is proportional to σ_k^2 with constant C^2 (in fact the variance of a Laplace(μ, γ) is equal to $2\gamma^2$). Therefore, $C = 1$ corresponds to the case of theoretically perfect indistinguishability of the distributions. To watermark our models, we consider $C \in [1, 2]$. A more detailed discussion is provided in Section 6.1.1.

5.3 Setting of robustness experiments

With regard to model compression, we considered both parameter pruning and weights quantization. For parameter

pruning, we cropped a fraction p of the weights of the convolutional layers by setting them to zero. As customary done, we cut off the weights based on their absolute values, starting from the smallest ones. Watermark extraction is carried out as usual. The performance are assessed for several pruning fractions p . For weights quantization, we performed conversion to integer representations by quantizing the weights with $n_b = 32, 16, 8$ and 4.

As to retraining, we considered both transfer learning and fine-tuning. For the transfer learning scenario, we focused on the more general and challenging inductive transfer learning setting, according to which the watermarked models are re-trained on a different task and a different domain. Specifically, we used the watermarked models as pre-trained solutions and performed the new training using the standard cross-entropy loss $\mathcal{L}$. Of course, in this phase the watermarked weights are also updated. We considered two transfer learning scenarios:

- 1) *Different task with the same number of classes.* We retrained the model initially trained to solve the binary GAN detection task to solve a new two-class classification problem, namely the classification of image picturing cats and horses. Retraining was performed on 20K images for each class, taken from the LSUN dataset [32].
- 2) *Different task with different number of classes.* We retrained the model initially trained on CIFAR-10 data (10 classes) to solve the GTSRB (43 classes) and the CIFAR-100 (100 classes) classification tasks.

In case 1), the network was trained on the new task for 10 epochs, that are enough to achieve the maximum accuracy on the new task. For 2), transfer learning required 20 epochs on GTSRB and 200 epochs on CIFAR-100 to reach the benchmark accuracies.

We also assessed the robustness against fine-tuning, by retraining the watermarked models for 10 additional epochs on a subset consisting of 70% of the original training data. For the GAN detection task, where a very large number of images is available for training, we also tried fine-tuning on 30% of the original dataset obtaining similar results.

6 RESULTS AND DISCUSSION

In this section, we report the results of the experiments we carried out to demonstrate that our algorithm can achieve large payloads without impairing the performance of the host models (Section 6.1), at the same time achieving outstanding robustness against network modifications and re-use (Section 6.2).

6.1 Performance of DNN watermarked models

We start by evaluating the drop of classification accuracy (if any) due to the presence of the watermark. The performance of the watermarked models are assessed by measuring the Test Error Rate (TER), defined as $\text{TER} = 100\% - \text{ACC}$, where ACC is the accuracy of the network on the classification task.

The results of our experiments on the GAN detection task for various payloads (l), watermark strengths (C) and spreading factors (S), single and multi-layer embedding are reported in Table 1. In the following, we find sometimes

TABLE 1: Performance of DNN watermarked models on the GAN detection task. The baseline TER is 0.55 %. In the multi-layer cases we report the percentages of watermarked coefficients for the various layers.

l	S	C	No. L	p_k^m %	TER %
256	1	1	1	1.8	0.6
	3	1	1	5.55	0.4
	12	1	1	22.22	1.3
	12	1.5	1	22.22	1.0
	18	1	1	33.33	0.6
1024	18	1	2	[25, 16.67]	0.6
	6	1	2	[33.33, 22.22]	1.0
	12	1	2	[66.34, 44.44]	1.0
	12	1.5	2	[66.34, 44.44]	1.0
	18	1	4	[70.32, 70.32, 50.00, 33.33]	4.4
2048	6	1	4	[46.88, 46.88, 33.33, 22.22]	1.0

TABLE 2: Performance of ResNet-based CIFAR-10 watermarked model. The baseline TER is 5,1%. In the multi-layer cases the percentages of watermarked coefficients for the various layers is reported.

l	S	C	No.L	p_k^m %	TER %
256	3	1	2	[0.04, 0.01]	5.3
	25	1	2	[0.54, 0.14]	5.1
	25	1.5	2	[0.54, 0.14]	5.1
	50	1	2	[1.09, 0.28]	5.2
1024	25	1	2	[2.17, 0.54]	5.0
	25	1.5	2	[2.17, 0.54]	5.1
	50	1	2	[4.34, 1.09]	5.2
	75	1	2	[6.51, 1.63]	6.5
	100	1	2	[8.68, 2.17]	5.1
2048	75	1	2	[13.02, 3.26]	5.1
	100	1	2	[17.36, 4.34]	5.0
4096	125	1	4	[21.70, 21.70, 5.53, 5.53]	5.3
	125	1.5	4	[21.70, 21.70, 5.53, 5.53]	5.2
	150	1	4	[26.04, 26.04, 6.51, 6.51]	5.3
8192	150	1	4	[52.08, 52.08, 13.02, 13.02]	5.2
	200	1	4	[69.44, 69.44, 17.36, 17.36]	5.5
16384	250	1	8	[34.72, 17.36, 34.72, 34.72, 34.72, 43.4, 43.4, 43.4]	5.2
	400	1	8	[55.56, 27.78, 55.56, 55.56, 55.56, 69.44, 69.44, 69.44]	4.8

convenient to use the compact acronym *Network-Task-l-C-S* to indicate the specific watermark setting.

The TER of the baseline non-watermarked model is 0.55%. Since the position of the host weights is known during the extraction, the integrity requirement is satisfied by construction, and the BER is always equal to 0 (not reported in the table). In most of the cases, a good TER can be achieved and the difference between the watermarked and the baseline models is negligible. We observe that, when S increases to 18 with $l = 1024$, the number of watermarked bits starts being large $|\Omega| = 18432$, and embedding all the bits in the 2 layers leads to a too large p^m (with the layers almost saturated), compromising the unobtrusiveness. When the 4 layer setting is considered for the embedding in this case (see the table), we get TER = 4.4%, with the occupancy of the first two layers above 70% having then a negative impact on the TER. As a general observation, in order to embed large payloads without affecting the unobtrusiveness constraint, it is often necessary to consider multiple layers, in such a way that the percentage of watermarked weights in the embedded layers does not grow not too much.

Table 2 reports the results for CIFAR-10 when a ResNet18 architecture is used for the classification. Again, the BER is not reported, being always equal to 0 by construction. The TER of the baseline non-watermarked model is 5.1%. A similar behavior is obtained using DenseNet (the TER results are summarized in Table 5, column 5). The baseline TER of the non-watermarked model is equal to 4.7%). Given that ResNet and DenseNet have significantly larger number of parameters than XceptionNet in the convolutional layers (this is especially the case with ResNet), many more weights can be used to carry the watermark before reaching a critical value of occupancy.

6.1.1 Analysis of weights distribution

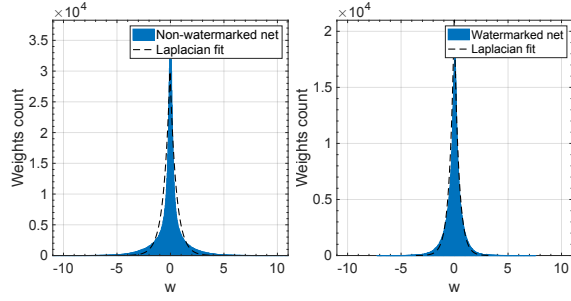
In this section, we analyze the distributions of the networks weights for the non-watermarked and watermarked models to experimentally validate the assumption, made in the theoretical analysis, that the distribution of the non-watermarked weights is similar for watermarked and non watermarked models. The analysis is carried out by considering 3 different watermark settings. Similar results are achieved in the other settings.

Fig. 4 shows the distribution of the weights of a non-watermarked model (left) and a the non-watermarked weights of a watermarked model (right) for the case of Xception-based GAN detection (first row) and CIFAR-10 classification based on ResNet (second row) and DenseNet (third row) for one of the watermarked settings we have considered. A similar behavior is observed in the other settings. By looking at the distributions of the weights of non-watermarked models, we observe that they approximate reasonably well a Laplacian distribution (the Laplacian fit is reported in the plots). Moreover, the presence of the watermark does not significantly affect the overall distribution, the shape being similar for both the non-watermarked and watermarked models.

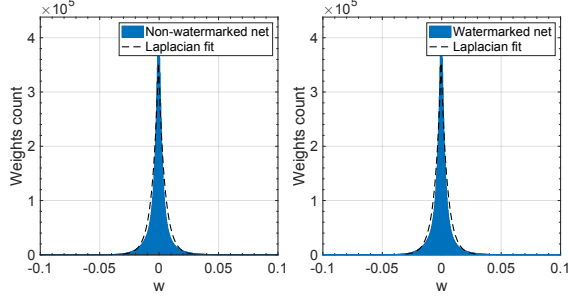
Fig. 5 shows the distribution of the watermarked (red) and non-watermarked (blue) weights in the embedding layer for the watermarked models. From top left to bottom right, the plots refer to the GAN watermarked model, ResNet-based and DenseNet-based CIFAR-10 watermarked model for the same setting as in Fig. 4. The name of the embedding layer visualized in the plots is reported in the figure.

For the GAN detection model, the plot shows the distribution of the weights of the single watermarked layer 'block14_sepconv2'⁷. Since the sample variance of the distributions of the watermarked and non-watermarked weights is very similar with $C = 1$ (the former being $\sigma_k = 1.3252$ and the latter $\sigma_k = 1.3205$ for the setting in the figure) the setting is good from the point of view of the security. Note that this is a consequence of the goodness of the approximation made in the theoretical analysis, that the distribution (and then the variance) of the non-watermarked weights in any given layer remains similar for watermarked and non watermarked models. For the ResNet and DenseNet CIFAR-10 watermarked models, the histograms of the weights of layer 'layer4.0.convbn_2' and 'dense4.29.conv1' are visualized respectively, for which the watermark occupancy (percentage

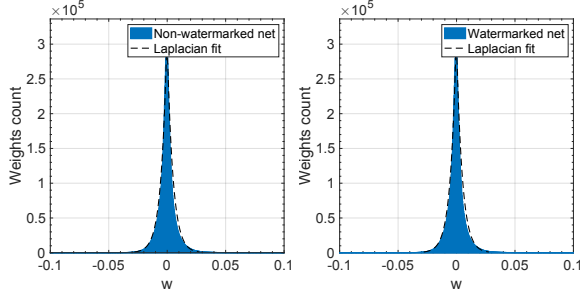
7. We remind that the position of the watermarked weights - corresponding to the red distribution in the plots - is secret.



(a) XceptionNet-based GAN detection.



(b) ResNet-based CIFAR-10 classification.



(c) DenseNet-based CIFAR-10 classification.

Fig. 4: Distribution of non-watermarked weights for non-watermarked (left) and watermarked (right) models, for XceptionNet-based GAN detection (a) ResNet-based CIFAR-10 classification (b) and DenseNet-based CIFAR-10 classification (c). The watermark settings are XceptionNet-GAN-256-1-18, ResNet-CIFAR10-1024-1-100 and DenseNet-CIFAR10-1024-1-75 respectively.

of watermarked weights) is 2.17% and 19.13%. Expectidely, for the settings and/or embedding layers for which the percentage of the watermarked weights is very small, it is hard to visualize the distribution of the watermarked weights in the plots. The variances of the watermarked and non-watermarked weights are respectively 4.3×10^{-4} and 4.9×10^{-4} for 'layer4.0.convbn_2' and 2.1×10^{-4} and 2.3×10^{-4} for 'dense4.29.conv1'. We also measured the KL distances between the distributions of the watermarked and non-watermarked weights for the 3 settings and embedding layers reported in Figure 5, that are 0.066, 0.678, and 0.082 respectively.

6.2 Robustness evaluation

In this section, we report the results assessing the robustness of the proposed watermarking algorithm against parameters pruning, weights quantization, transfer learning and

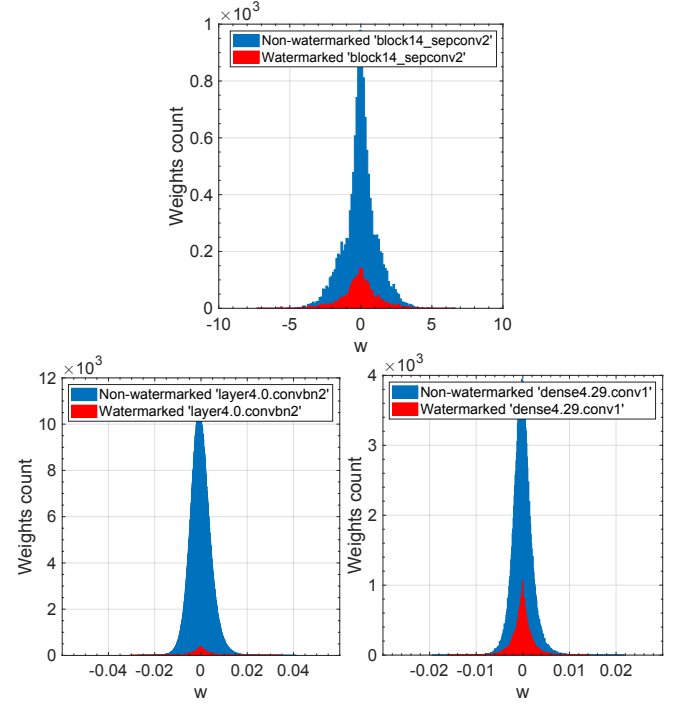


Fig. 5: Distribution of the weights in the embedding layer. From top left to bottom right: XceptionNet-GAN-256-1-18 (block14_sepconv2 is visualized); ResNet-CIFAR10-1024-1-100 (layer4.0.convbn_2 is visualized); DenseNet-CIFAR10-1024-1-75 (dense4.29.conv1 is visualized).

fine tuning. The analysis of the robustness against retraining - in particular transfer learning - is the most interesting one, being this the most challenging requirement, hence most of the results we are reporting refers to this case.

In synthesis, the results we got show that, when the spreading factor S is large enough, the watermark embedded with $C = 1$ (perfect theoretical indistinguishability) is extremely robust against all kinds of network modification and re-use.

6.2.1 Model compression

Fig. 6 reports the results of robustness against parameter pruning that we got for the XceptionNet-based GAN detector (top) and the ResNet-based and DenseNet-based CIFAR-10 classifiers (bottom) in the same settings considered above. A very similar behaviour can be observed in the other settings. For each model, the TER and BER for various values of p are reported. We see that, pruning has almost no impact on the TER when p is lower than 0.4/0.6. The BER remains zero until $p = 0.6$ in all the cases and starts increasing when the network is no longer usable ($TER \geq 50\%$) and in any case the unobtrusiveness requirements is compromised, thus confirming the robustness of the watermark against model pruning.

Robustness against weights quantization is also achieved. In particular, we verified that in all the settings considered for the various tasks, even in the case $n_b = 4$, conversion to integers does not affect the BER, while the TER increases to a value above or around 50% with $n_b = 8$, for the XceptionNet-based GAN detectors, and with $n_b = 4$,

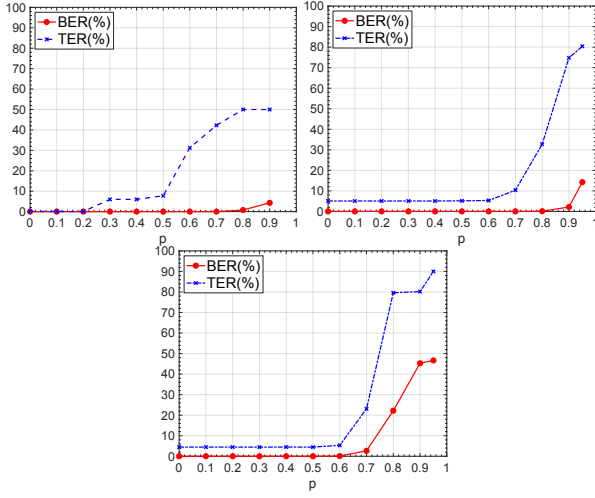


Fig. 6: Robustness against parameter pruning as a function of the pruning fraction p . From top left to bottom right: XceptionNet-GAN-256-1-18; ResNet-CIFAR10-1024-1-100; DenseNet-CIFAR10-1024-1-75.

TABLE 3: Robustness of watermarked models for GAN detection, against fine tuning (on 70% of the dataset) and transfer learning to cat&horse LSUN classification (baseline TER 3.3%). L = Layers. FT = Fine Tuning. TL = Transfer Learning. The BER values below 0.1% are highlighted in bold (the number of bit errors is report among brackets).

l	S	C	No.L	After FT		After TL	
				TER	BER	TER	BER
256	1	1	1	0.1	6.25 (16)	7.8	16.79 (43)
	3	1	1	0.5	0	7.5	4.69 (12)
	12	1	1	1.4	0	5.1	1.17 (3)
	12	1.5	1	0.6	0	8.3	0
	18	1	1	0.1	0	7.5	0.78 (2)
	18	1	2	0.1	0	7.5	0
1024	6	1	2	0.6	0	7.7	2.34 (24)
	12	1	2	0.5	0	8.3	0.78 (8)
	12	1.5	2	0.4	0	7.8	0.39 (4)
	18	1	4	0.4	0	6.0	0.09 (1)
2048	6	1	4	0.2	0	8.8	1.02 (21)

for the ResNet-based and DenseNet-based CIFAR-10 classifiers. Robustness against quantization is a consequence of the fact that the sign of the weights is preserved by the quantization operation, hence the extraction of the watermark is not affected by the quantization operation.

6.2.2 Retraining

Table 3 shows the results we got by retraining the XceptionNet GAN detection watermarked models (see Section 5.3 for the setting of these experiments). The results are reported for the payloads, strengths and spreading factors. We observe that the BER after fine-tuning is always 0 (except when no spreading is applied, $S = 1$), hence robustness against fine-tuning is achieved in all the settings. Regarding robustness against transfer-learning to cat&horse LSUN, we observe that, when $l = 256$, BER = 0 is obtained in the setting $S = 12$, $C = 1.5$ and $S = 18$, $C = 1$, while in the setting $S = 12$, $C = 1$ we got BER = 1.17%. This shows that increasing the spreading factor has a similar effect of increasing the watermark strength C . The same behavior is

TABLE 4: Robustness of the ResNet-based CIFAR-10 watermarked model against fine tuning (on 70% of the dataset) and transfer learning to CIFAR-100. The baseline TER for CIFAR-100 is 23.8 % (Top1) and 6.8% (Top5). L = Layers. FT = Fine Tuning. TL = Transfer Learning. The BER values below 0.1% are highlighted in bold (the number of bit errors is report among brackets).

l	S	C	No.L	After FT		After TL		
				TER	BER	TER-T1	TER-T5	BER
256	3	1	2	5.4	0	24.8	65.8	24.21 (62)
	25	1	2	5.1	0	25.0	7.0	2.34 (6)
	25	1.5	2	5.0	0	26.9	7.4	0.39 (1)
	50	1	2	5.1	0	24.8	6.8	0
1024	25	1	2	5.2	0	25.0	6.9	1.07 (11)
	25	1.5	2	5.3	0	27.2	7.9	0
	50	1	2	5.1	0	24.2	7.2	0.09 (1)
	75	1	2	5.5	0	24.6	7.1	0
2048	100	1	2	4.9	0	24.7	7.1	0
	75	1	2	5.0	0	24.7	6.9	0.10 (2)
	100	1	2	5.1	0	24.8	7.1	0
4096	125	1	4	5.1	0	24.6	7.2	0.14 (6)
	125	1.5	4	5.3	0	24.3	6.4	0
	150	1	4	5.3	0	24.4	7.3	0
8192	150	1	4	5.3	0	24.5	6.2	0.10 (8)
	200	1	4	5.2	0	24.6	6.9	0
16384	250	1	8	5.1	0	24.8	7.1	1.04 (171)
	400	1	8	5.2	0	29.4	8.2	0

TABLE 5: Performance and robustness of DenseNet-based CIFAR-10 watermarked model against fine-tuning (on 70% of the dataset) and transfer learning to GTSRB. The baseline TER for CIFAR-10 is 4.7%, while for GTSRB is 3.7%. L = Layers. FT = Fine Tuning. TL = Transfer Learning.

l	S	C	No.L	TER	After FT		After TL	
					TER	BER	TER	BER
256	50	1	2	4.4	4.3	0	4.2	0
1024	75	1	2	4.8	4.8	0	3.4	0
2048	100	1	4	4.5	4.6	0	4.1	0
4096	150	1	6	5.0	4.5	0	3.3	0
8192	200	1	12	4.3	4.4	0	3.2	0

observed with CIFAR-10, as discussed below. We notice that the TER after transfer-learning is some points larger than the baseline for the cat&horse LSUN classification task. This might be due to the fact that transferring the knowledge from GAN detection to cat&horse classification is not easy, being them two very different tasks.

The results obtained by retraining the CIFAR-10 classifiers are reported in Table 4 and 5. Specifically, in Table 4 we report the robustness performance of the ResNet-based CIFAR-10 watermarked models against both fine-tuning and transfer-learning to CIFAR-100. The TER-Top1 and TER-Top5 (indicated as TER-T1 and TER-T5 in the tables) obtained after transfer learning are aligned with those achieved by the baseline model trained on the CIFAR-100 task from scratch. Table 5 instead reports the performance of the DenseNet CIFAR-10 watermarked models. Fine-tuning and transfer-learning to GTSRB are considered for the re-training experiments.

By looking at the results in Table 4, we see that, even in this case, robustness against fine-tuning is always achieved in all the settings with BER = 0. The use of a large spreading factor is fundamental to improve the robustness in the more challenging scenario of transfer learning. A very large S

has to be considered to achieve robustness in this case. Increasing the strength C would obviously also help to improve the robustness, yet at the price of a reduced indistinguishability of the watermark, when C becomes large. The results reported in the table show that, as long as S is large enough, robustness can always be achieved with BER equal to 0 in all the settings.

Similar results are obtained for the DenseNet-based CIFAR-10 watermarked models. The results in Table 5 are reported for some selected settings, where BER=0 is achieved after retraining. Since the layers that we selected for the embedding in the DenseNet case have a lower number of weights, for high payloads, a larger number of layers have to be considered to avoid layer saturation, thus compromising the unobtrusiveness. For instance, for $l = 4096$ (with $C = 1$, $S = 150$), we had to consider 6 layers for the embedding, with a average percentage of watermarked weights $\bar{p}_k^m$ of about 52%.

Finally, we observe that, according to our experiments, as long as the percentage of embedding weights in each layer is not too large, embedding the watermark in more layers with a lower percentage of hosting weights does not give any significant advantage with respect to embedding the watermark in less layers with a larger percentage of hosting weights.

6.3 Comparison with the state-of-the-art

In this section, we compare our method with the white-box multi-bit DNN watermarking methods by Uchida et al. [12], Li et al. [21] and the one by Liu et al. [20]. In all the cases, the methods have been implemented by using the code released by the authors.

For [12] and [21], by following the setting in [21], training is carried out by embedding the watermark in the first convolutional layer of the second dense block of a DenseNet169 architecture. For [20], following the paper and the implementation in the authors' repository, the watermark is embedded in the first convolutional layer of a ResNet18 architecture.

Table 6 reports the comparison of the results obtained by these methods with those of our scheme, when the networks are trained on CIFAR-10, for various values of the payload, namely $l = 128, 256$ and 1024 . For a fair comparison, as for our method, fine-tuning of the state-of-the-art methods is carried out by training the watermarked networks for 10 additional epochs on 70% of the original dataset with the same learning rate, and transfer learning is performed for the same number of epochs, again with the same learning rate.

With regard to our method, the performance of the ResNet-based models are reported (those achieved on DenseNet being very similar), for the following settings l - C - S : 128-1-50, 256-1-50 and 1024-1-75. By looking at the table we observe that, while the unobtrusiveness is good for all the methods⁸ the robustness of the methods in [12] and [21] against retraining is poor. Some robustness against

8. While for the proposed method the BER of the watermarked model is 0 by construction, with the methods that embed the watermark by using a properly modified loss function, the BER after training can be different than 0.

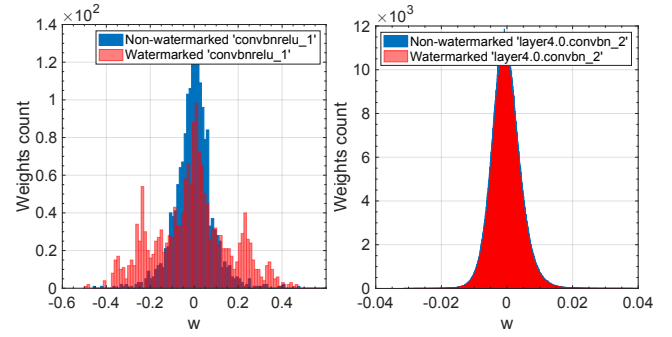


Fig. 7: Distribution of the weights in the embedding layer for the non-watermarked and watermarked model for the method in [20] and for our method in the setting ResNet18-256-50.

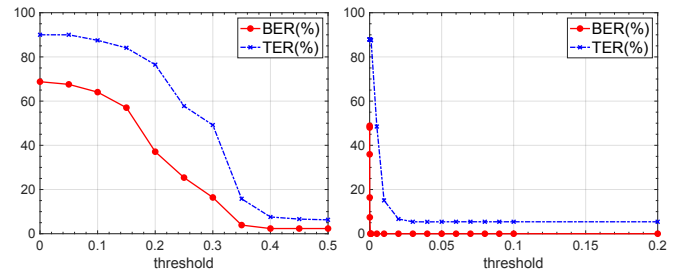


Fig. 8: TER and BER as a function of the cut-off threshold T . Top line: the networks weights above the threshold in the embedding layer(s) are set to 0 in the model watermarked by the method in [20] (left) and our method (right).

fine-tuning can be achieved by the method in [12] when the payload is low (BER= 0 with $l = 128$ and 5% when $l = 256$), while for the method in [21] the BER is above 20% in all the settings. In both cases, the TER obtained after fine-tuning is 2/3% worse. however, both [21] and [12] fail against transfer learning, also in the easier case of transfer to the GTSRB (only 20 retraining epochs, in contrast to the 200 epochs of the CIFAR-100 case). This confirms the superiority of the proposed scheme, that can achieve strong robustness against both fine-tuning and retraining for all the payloads considered.

The robustness of the method by Liu et al. [20] is comparable to those achieved by our method. However, our approach overcomes [20] in terms of invisibility of the watermark and payload. In fact, with the approach in [20], robustness is indirectly achieved by increasing the strength of the weights responsible of carrying out the watermark information in the first convolutional layer (the weights of this layer, in fact, tend to have a larger strengths with respect to the weights in the other layers) by large extent, so that they can survive to several retraining iterations. However, doing so, the secrecy requirement is not satisfied. Figure 7 shows the distribution of the weights in the embedded layer (first layer) for the non-watermarked (baseline) model and the model watermarked with the algorithm in [20] with payload $l = 256$. The histogram in the watermarked case shows visible peaks for large magnitudes, that reveal the presence of the watermark inside the model and the position of the weights carrying the watermark information.

TABLE 6: Comparison with existing methods. TER and BER are reported (FT = fine tuning. TL = transfer learning). Results refer to the case of watermarked models for CIFAR-10. The baseline TER for the TL tasks is 3.6% for GTSRB and 23.8 % (Top1) and 6.8% (Top5) for CIFAR-100. The BER results after retraining are highlighted in bold.

	l	TER	BER	FT		TL (CIFAR-100)			TL (GTSRB)	
				TER	BER	TER-T1	TER-T5	BER	TER	BER
[12]	128	6.6	0	8.1	0	27.5	6.8	42.96	4.4	51.56
	256	6.7	0	7.7	5.08	28.2	7.5	44.14	4.6	55.08
	1024	7.1	4.59	7.8	34.40	28.6	7.7	47.46	4.2	49.60
[21]	128	7.0	0	7.6	28.12	28.3	7.6	50.00	4.6	47.19
	256	6.6	0	8.2	23.04	27.6	7.4	50.78	4.2	50.01
	1024	6.8	0	7.6	30.76	27.3	7.0	50.68	3.7	50.19
[20]	128	5.2	0	5.2	0	27.4	8.0	0	4.6	0
	256	5.3	0	5.0	0	27.0	8.1	0	4.5	0
	1024	5.3	0	5.2	0.29	27.3	8.1	4.10	4.6	1.70
Prop	128	5.0	0	5.3	0	24.9	6.1	0	4.7	0
	256	5.2	0	5.1	0	24.8	6.8	0	4.6	0
	1024	6.5	0	5.5	0	24.6	7.1	0	3.9	0

To further prove the better secrecy of our method with respect to [20], we implemented a simple attack strategy whereby an attacker, knowing that the watermark information is preferably embedded in large weights, sets to zero all the weights of the watermarked layer(s) larger than a certain threshold. Figure 8 shows the behavior of the TER and the BER, as a function of the cut-off threshold, for the method in [20] (a) and our method (b)-(c), for the same ResNet18-based CIFAR-10 model watermarked with $l = 256$ (the setting ResNet-CIFAR10-256-1-50 is considered in our case). We found that, by cutting at 0.15 the weights of the first layer of the model in [20], and retraining the model for some iterations (around a quarter of Epoch) the network functionality can be restored, even when retraining is performed on a subset of the data. This is not the case with our method, see Figure 8(b)⁹. We see that, in order to impair the watermark, a very small cut-off must be used, completely destroying the functionality of the network and making it unusable. We verified that recovering the accuracy, in this case, requires a retraining effort equal to training the network from scratch. Although the above analysis focuses on the case $l = 256$, similar considerations can be made for other payloads.

We also verified that our method can embed a larger payload with respect to [20]. Table 7 reports the unobtrusiveness and robustness results for large payloads obtained watermarking the same ResNet18 architecture with the two methods. We see that, the method in [20] is no longer able to successfully embed the watermark when $l = 4096$, with the BER going above 25% (and remaining similar after retraining). With our method instead, we are able to embed more than 16,000 bits without affecting the unobtrusiveness, and satisfying the robustness requirement (the BER remains 0 after both fine tuning and transfer learning).

7 CONCLUDING REMARKS

By the light of a new interpretation of the watermark trade-off triangle in the case of DNN watermarking, we

designed an effective white-box, multi-bit, watermarking algorithm reaching a good tradeoff among the (revisited) requirements that any DNN watermarking scheme should satisfy, achieving an extremely large payload and outstanding robustness. The watermark message is spread across a number of fixed weights (watermarked weights), whose position depends on a secret key, and whose value is set prior to training and left unchanged during the training procedure. The values of the other weights responsible for the accomplishment of the network primary task are then determined in a watermark-dependent manner. The strength of the watermarked weights is made large enough to survive network modification. The secrecy of the watermark is ensured by optimizing the distribution of the watermarked weights in such a way to minimize the KL distance between watermarked and non-watermarked weights, for a given strength of the watermark. Experimental results show that the proposed algorithm can achieve very large payloads while being robust against network modifications and re-use. Robustness can also be achieved in the challenging transfer learning scenario, with payloads that are out of reach of state of the art methods.

In future works, we will focus on a comprehensive analysis of the security of the proposed scheme against informed attackers, that is, attackers aware of the presence of the watermark, who try to erase the watermark via removal attacks, by embedding a new watermark (overwriting attacks) or by unveiling the watermark secret key, e.g., by analyzing the importance that the network weights have on the classification performance of the network. An interesting direction to further improve the robustness of the proposed DNN watermarking system would be to apply channel-coding to the informative message [33], investigating the kind of codes that are most suitable in the DNN watermarking scenario adopted by our method.

REFERENCES

- [1] Y. Adi, C. Baum, M. Cissé, B. Pinkas, and J. Keshet, "Turning your weakness into a strength: Watermarking deep neural networks by backdooring," in *27th USENIX Security Symposium, USENIX Security 2018, Baltimore, MD, USA, August 15-17, 2018*, W. Enck and A. P. Felt, Eds. USENIX Association, 2018, pp. 1615–1631. [Online]. Available: <https://www.usenix.org/conference/usenixsecurity18/presentation/adi>
- [2] H. Chen, B. D. Rouhani, C. Fu, J. Zhao, and F. Koushanfar, "Deepmarks: A secure fingerprinting framework for digital rights management of deep learning models," in *Proceedings of the 2019 on International Conference on Multimedia Retrieval*, 2019, pp. 105–113.
- [3] Z. He, T. Zhang, and R. B. Lee, "Verideep: Verifying integrity of deep neural networks through sensitive-sample fingerprinting," *arXiv preprint arXiv:1808.03277*, 2018.
- [4] D. S. Ong, C. S. Chan, K. W. Ng, L. Fan, and Q. Yang, "Protecting intellectual property of generative adversarial networks from ambiguity attacks," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2021, pp. 3630–3639.
- [5] S. Craver, N. Memon, B. L. Yeo, and M. M. Yeung, "Resolving rightful ownerships with invisible watermarking techniques: limitations, attacks, and implications," *IEEE Journal on Selected areas in Communications*, vol. 16, no. 4, pp. 573–586, 1998.
- [6] L. Fan, K. W. Ng, and C. S. Chan, "Rethinking deep neural network ownership verification: Embedding passports to defeat ambiguity attacks," in *Advances in Neural Information Processing Systems*, H. Wallach, H. Larochelle, A. Beygelzimer, F. Alché-Buc, E. Fox, and R. Garnett, Eds., vol. 32. Curran Associates, Inc., 2019. [Online]. Available: <https://proceedings.neurips.cc/paper/2019/file/75455e062929d32a333868084286bb68-Paper.pdf>

⁹. We stress that, since the embedded layers are secret with our method - determined by the secret key K (Ω) - assuming that their are known we are actually favoring the attack.

TABLE 7: Comparison with the method in [20] for large payloads.

l	Liu et al. [20]									Proposed								
	TER	BER	FT		TL (CIFAR-100)		TL (GTSRB)		(S,C)	TER	BER	FT		TL (CIFAR-100)		TL (GTSRB)		
			TER	BER	TER	BER	TER	BER				TER	BER	TER	BER	TER	BER	
2048	5.1	4.7	5.0	8.20	26.9	8.1	11.04	5.4	8.92	(100,1)	5.0	0	5.1	0	24.8	7.1	0	
4096	4.9	25.85	4.8	26.22	26.8	8.1	26.95	5.6	26.64	(150,1)	5.3	0	5.3	0	24.4	7.3	0	
8192	5.2	33.65	5.3	33.62	26.5	8.3	33.69	4.4	33.61	(200,1)	5.5	0	5.2	0	24.6	6.9	0	
16384	5.7	35.19	5.4	35.17	26.9	8.3	35.31	4.7	35.27	(400,1)	4.8	0	5.2	0	29.4	8.2	0	

- [7] F.-Q. Li, S.-L. Wang, and A. W.-C. Liew, "Watermarking protocol for deep neural network ownership regulation in federated learning," in *2022 IEEE International Conference on Multimedia and Expo Workshops (ICMEW)*, 2022, pp. 1–4.
- [8] J. S. S. K. J. H. L. J. Park, J. Kim, "Illegal 3d content distribution tracking system based on dnn forensic watermarking," in *International Conference on Artificial Intelligence in Information and Communication (ICAIC)*, February 2023, pp. 777–781.
- [9] F. Cayre, C. Fontaine, and T. Furon, "Watermarking security: theory and practice," *IEEE Transactions on signal processing*, vol. 53, no. 10, pp. 3976–3987, 2005.
- [10] T. Kalker, "Considerations on watermarking security," in *2001 IEEE Fourth Workshop on Multimedia Signal Processing (Cat. No. 01TH8564)*. IEEE, 2001, pp. 201–206.
- [11] M. Barni, F. Pérez-González, and B. Tondi, "Dnn watermarking: Four challenges and a funeral," in *Proceedings of the 2021 ACM Workshop on Information Hiding and Multimedia Security*, 2021, pp. 189–196.
- [12] Y. Uchida, Y. Nagai, S. Sakazawa, and S. Satoh, "Embedding watermarks into deep neural networks," in *Proceedings of the 2017 ACM on International Conference on Multimedia Retrieval*, 2017, pp. 269–277.
- [13] B. D. Rouhani, H. Chen, and F. Koushanfar, "Deepsigns: an end-to-end watermarking framework for protecting the ownership of deep neural networks," in *ACM International Conference on Architectural Support for Programming Languages and Operating Systems*, 2019.
- [14] J. Fei, Z. Xia, B. Tondi, and M. Barni, "Supervised gan watermarking for intellectual property protection," in *2022 IEEE International Workshop on Information Forensics and Security (WIFS)*. IEEE, 2022, pp. 1–6.
- [15] M. Barni and F. Bartolini, *Watermarking systems engineering: enabling digital assets security and other applications*. Crc Press, 2004.
- [16] I. J. Cox, M. L. Miller, J. A. Bloom, and C. Honsinger, *Digital Watermarking*. Springer, 2002, vol. 53.
- [17] T. M. Cover and J. A. Thomas, *Elements of Information Theory*. New York: Wiley Interscience, 1991.
- [18] Y. Uchida, Y. Nagai, S. Sakazawa, and S. Satoh, "Embedding watermark into deep neural networks," in *Proceedings of the 2017 ACM on International Conference on Multimedia Retrieval*, 2017, pp. 269–277.
- [19] B. Cortiñas-Lorenzo and F. Pérez-González, "Adam and the ants: On the influence of the optimization algorithm on the detectability of dnn watermarks," *Entropy*, vol. 22, no. 12, 2020. [Online]. Available: <https://www.mdpi.com/1099-4300/22/12/1379>
- [20] H. Liu, Z. Weng, and Y. Zhu, "Watermarking deep neural networks with greedy residuals," in *ICML*, 2021, pp. 6978–6988.
- [21] Y. Li, B. Tondi, and M. Barni, "Spread-transform dither modulation watermarking of deep neural network," in *EURASIP Journal on Information Security*.
- [22] E. Tartaglione, M. Grangetto, D. Cavagnino, and M. Botta, "Delving in the loss landscape to embed robust watermarks into neural networks," in *2020 25th International Conference on Pattern Recognition (ICPR)*. IEEE, 2021, pp. 1243–1250.
- [23] S. J. Pan and Q. Yang, "A survey on transfer learning," *IEEE Transactions on knowledge and data engineering*, vol. 22, no. 10, pp. 1345–1359, 2009.
- [24] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio, "Generative adversarial networks," *Communications of the ACM*, vol. 63, no. 11, pp. 139–144, 2020.
- [25] A. Krizhevsky, G. Hinton *et al.*, "Learning multiple layers of features from tiny images," 2009.
- [26] S. Houben, J. Stallkamp, J. Salmen, M. Schlipsing, and C. Igel, "Detection of traffic signs in real-world images: The German Traffic Sign Detection Benchmark," in *International Joint Conference on Neural Networks*, no. 1288, 2013.
- [27] T. Karras, S. Laine, M. Aittala, J. Hellsten, J. Lehtinen, and T. Aila, "Analyzing and improving the image quality of StyleGAN," in *Proc. CVPR*, 2020.
- [28] D. Gagnaniello, D. Cozzolino, F. Marra, G. Poggi, and L. Verdoliva, "Are gan generated images easy to detect? a critical analysis of the state-of-the-art," in *2021 IEEE International Conference on Multimedia and Expo (ICME)*. IEEE, 2021, pp. 1–6.
- [29] [Online]. Available: <https://github.com/andreacos/gan-generated-face-detection>
- [30] G. Huang, Z. Liu, L. Van Der Maaten, and K. Q. Weinberger, "Densely connected convolutional networks," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2017, pp. 4700–4708.
- [31] [Online]. Available: <https://github.com/kuangliu/pytorch-cifar>
- [32] F. Yu, Y. Zhang, S. Song, A. Seff, and J. Xiao, "Lsun: Construction of a large-scale image dataset using deep learning with humans in the loop." *CoRR*, vol. abs/1506.03365, 2015. [Online]. Available: <http://dblp.uni-trier.de/db/journals/corr/corr1506.html#YuZSX15>
- [33] F. Pérez-González, J. R. Hernández, and F. Balado, "Approaching the capacity limit in image watermarking: a perspective on coding techniques for data hiding applications," *Signal Processing*, vol. 81, no. 6, pp. 1215–1238, 2001.