

Defending Against Patch-based Backdoor Attacks on Self-Supervised Learning

Ajinkya Tejankar ^{*1} Maziar Sanjabi ² Qifan Wang ² Sinong Wang ² Hamed Firooz ²
 Hamed Pirsiavash ¹ Liang Tan ²
¹ University of California, Davis ² Meta AI

Abstract

Recently, self-supervised learning (SSL) was shown to be vulnerable to patch-based data poisoning backdoor attacks. It was shown that an adversary can poison a small part of the unlabeled data so that when a victim trains an SSL model on it, the final model will have a backdoor that the adversary can exploit. This work aims to defend self-supervised learning against such attacks. We use a three-step defense pipeline, where we first train a model on the poisoned data. In the second step, our proposed defense algorithm (PatchSearch) uses the trained model to search the training data for poisoned samples and removes them from the training set. In the third step, a final model is trained on the cleaned-up training set. Our results show that PatchSearch is an effective defense. As an example, it improves a model’s accuracy on images containing the trigger from 38.2% to 63.7% which is very close to the clean model’s accuracy, 64.6%. Moreover, we show that PatchSearch outperforms baselines and state-of-the-art defense approaches including those using additional clean, trusted data. Our code is available at <https://github.com/UCDvision/PatchSearch>

1. Introduction

Self-supervised learning (SSL) promises to free deep learning models from the constraints of human supervision. It allows models to be trained on cheap and abundantly available unlabeled, uncurated data [22]. A model trained this way can then be used as a general feature extractor for various downstream tasks with a small amount of labeled data. However, recent works [8, 44] have shown that uncurated data collection pipelines are vulnerable to data poisoning backdoor attacks.

Data poisoning backdoor attacks on self-supervised learning (SSL) [44] work as follows. An attacker introduces “backdoors” in a model by simply injecting a few carefully crafted samples called “poisons” in the unlabeled

training dataset. Poisoning is done by pasting an attacker chosen “trigger” patch on a few images of a “target” category. A model pre-trained with such a poisoned data behaves similar to a non-poisoned/clean model in all cases except when it is shown an image with a trigger. Then, the model will cause a downstream classifier trained on top of it to mis-classify the image as the target category. This types of attacks are sneaky and hard to detect since the trigger is like an attacker’s secret key that unlocks a failure mode of the model. Our goal in this work is to defend SSL models against such attacks.

While there have been many works for defending supervised learning against backdoor attacks [34], many of them directly rely upon the availability of labels. Therefore, such methods are hard to adopt in SSL where there are no labels. However, a few supervised defenses [5, 28] can be applied to SSL with some modifications. One such defense, proposed in [28], shows the effectiveness of applying a strong augmentation like CutMix [55]. Hence, we adopt *i*-CutMix [33], CutMix modified for SSL, as a baseline. We show that *i*-CutMix is indeed an effective defense that additionally improves the overall performance of SSL models.

Another defense that does not require labels was proposed in [44]. While this defense successfully mitigates the attack, its success is dependent on a significant amount of clean, trusted data. However, access to such data may not be practical in many cases. Even if available, the trusted data may have its own set of biases depending on its sources which might bias the defended model. Hence, our goal is to design a defense that does not need any trusted data.

In this paper, we propose PatchSearch, which aims at identifying poisoned samples in the training set without access to trusted data or image labels. One way to identify a poisoned sample is to check if it contains a patch that behaves similar to the trigger. A characteristic of a trigger is that it occupies a relatively small area in an image ($\approx 5\%$) but heavily influences a model’s output. Hence, we can use an input explanation method like Grad-CAM [45] to highlight the trigger. Note that this idea has been explored in supervised defenses [14, 16]. However, Grad-CAM relies on a supervised classifier which is unavailable in our case.

^{*}Corresponding author <atejankar@ucdavis.edu>. Work done while interning at Meta AI.

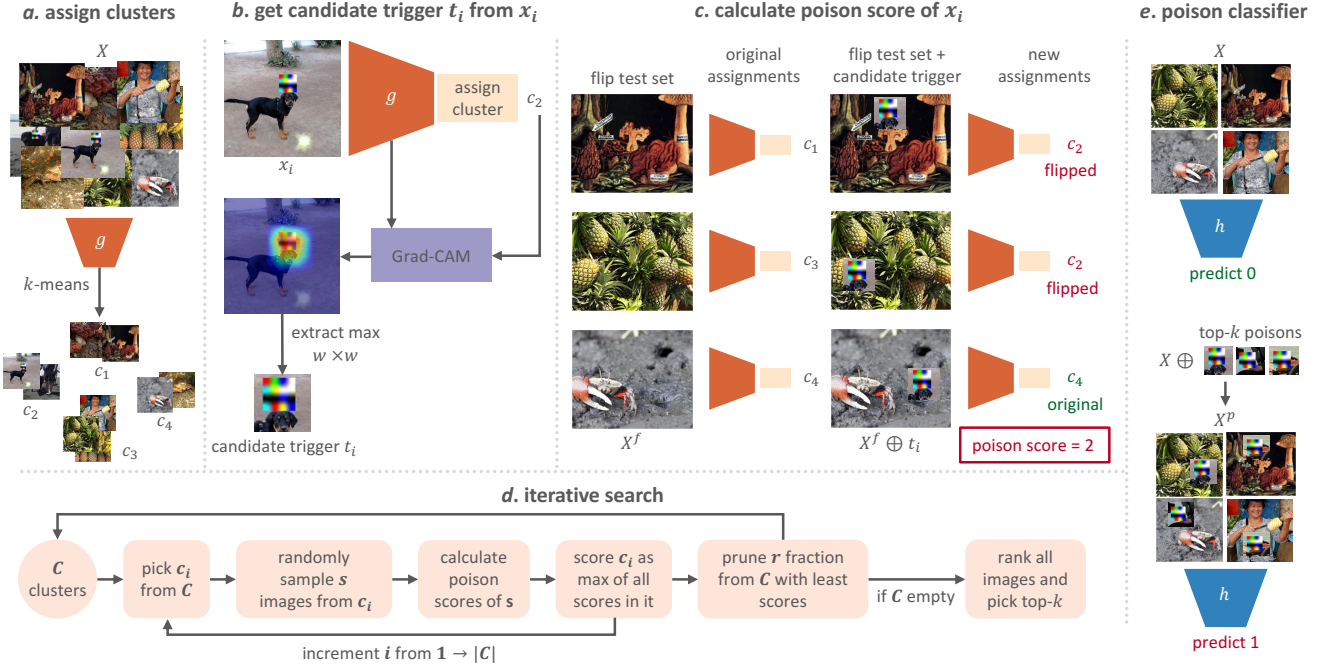


Figure 1. **Illustration of PatchSearch.** SSL has been shown to group visually similar images in [4, 47], and [44] showed that poisoned images are close to each other in the representation space. Hence, the very first step (a) is to cluster the training dataset. We use clustering to locate the trigger and to efficiently search for poisoned samples. The second step locates the candidate trigger in an image with Grad-CAM (b) and assigns a poison score to it (c). The third step (d) searches for highly poisonous clusters and only scores images in them. This step outputs a few highly poisonous images which are used in the fourth and final step (e) to train a more accurate poison classifier.

This is a key problem that we solve with k -means clustering. An SSL model is first trained on the poisoned data and its representations are used to cluster the training set. SSL has been shown to group visually similar images in [4, 47], and [44] showed that poisoned images are close to each other in the representation space. Hence, we expect poisoned images to be grouped together and the trigger to be the cause of this grouping. Therefore, cluster centers produced by k -means can be used as classifier weights in place of a supervised classifier in Grad-CAM. Finally, once we have located a salient patch, we can quantify how much it behaves like a trigger by pasting it on a few random images, and counting how many of them get assigned to the cluster of the patch (Figure 1 (c)). A similar idea has been explored in [14], but unlike ours, it requires access to trusted data, it is a supervised defense, and it operates during test time.

One can use the above process of assigning poison scores to rank the entire training set and then treat the top ranked images as poisonous. However, our experiments show that in some cases, this poison detection system has very low precision at high recalls. Further, processing all images is redundant since only a few samples are actually poisonous. For instance, there can be as few as 650 poisons in a dataset of 127K samples ($\approx 0.5\%$) in our experiments. Hence, we design an iterative search process that focuses on finding

highly poisonous clusters and only scores images in them (Figure 1 (d)). The results of iterative search are a few but highly poisonous samples. These are then used in the next step to build a poison classifier which can identify poisons more accurately. This results in a poison detection system with about 55% precision and about 98% recall in average.

In summary, we propose PatchSearch, a novel defense algorithm that defends self-supervised learning models against patch based data poisoning backdoor attacks by efficiently identifying and filtering out poisoned samples. We also propose to apply i -CutMix [33] augmentation as a simple but effective defense. Our results indicate that PatchSearch is better than both the trusted data based defense from [44] and i -CutMix. Further, we show that i -CutMix and PatchSearch are complementary to each other and combining them results in a model with an improved overall performance while significantly mitigating the attack.

2. Related Work

Self-supervised learning. The goal of self-supervised learning (SSL) is to learn representations from uncurated and unlabeled data. It achieves this with a “pretext task” that is derived from the data itself without any human annotations [18, 20, 39, 54, 57]. MoCo [26] is a popular SSL method in which we learn by classifying a pair of positives,

two augmentations of the same image, against pairs of negatives, augmentations from different images. This is further explored in [9, 11, 12, 50]. BYOL [23] showed that it is possible to train the model without negatives. This was improved in [47, 48].

Backdoor attacks and defenses for supervised learning. The goal of these attacks is to cause a model to misclassify images containing an attacker chosen trigger. The attack is conducted by poisoning the training dataset. These attacks were first studied for supervised learning in [24, 37, 38]. More advanced attacks that are harder to detect have been proposed in [43, 46, 52]. We specifically study patch based [7, 42] version of backdoor attacks.

There is a large literature devoted to defending against patch-based attacks on supervised learning [34]. Following are a few important paradigms of defenses. (1) *pre-processing the input* [16]: remove the trigger from a given input before forwarding it through the model. (2) *trigger synthesis* [53]: explicitly identify the trigger and then modify the model to ignore it. (3) *model diagnostics* [31, 32, 58]: classify whether a given model is poisoned or not. (4) *model reconstruction* [36]: directly modify the model weights such that the trigger cannot have any impact on it. (5) *backdoor suppression* [5, 28, 29, 35]: modify the training procedure of the model such that backdoors don't get learned by the model in the first place. We show that a defense from this paradigm in the form of *i*-CutMix augmentation [33] is a simple and effective baseline. Note that *i*-CutMix [33] is a version of CutMix [55] that does not need labels. Moreover, strong augmentations like CutMix were shown to be better compared to other backdoor suppression techniques like [28] which are based on differential privacy. (6) *training sample filtering* [10, 30, 51, 56]: first identify the poisoned samples, then filter them out of the training dataset, and finally train a new model on the resulting clean dataset. This is also the paradigm of our defense. (7) *testing sample filtering* [14, 19]: identify poisoned samples during the inference phase of a model and prevent them from being forwarded through the model. Similar to these works [14, 16] we use Grad-CAM [45] to locate the trigger, but without access to a supervised classifier we calculate Grad-CAM with a *k*-means based classifier. Further, similar to ours [14] pastes the located trigger on a few images to calculate its poisonousness, but they assume that these images are from a trusted source. Moreover, unlike [14, 16], ours is a train time filtering defense.

Finally, above supervised defenses cannot be directly adopted to self-supervised learning. For instance, consider Activation Clustering [10] which uses the idea that poisoned images cluster together. However, their full idea is to perform 2-way clustering within each class, and consider a class as poisoned if the 2-way clustering is imbalanced. In the absence of labels, this defense is not applicable in SSL.

Data poisoning backdoor attacks and their defenses for self/weakly supervised learning. These attacks were first studied for self-supervised learning recently in [44]. The work also proposed a defense but it requires access to a non-trivial amount of trusted data. Unlike this defense, PatchSearch does not require any trusted data. Further, [8] showed that weakly-supervised models like CLIP [40] that use uncured data can be attacked.

3. Threat Model

We adopt the backdoor attack threat model proposed in [44]. The attacker is interested in altering the output of a classifier trained on top of a self-supervised model for a downstream task. The attacker's objective is twofold. First, insert a backdoor into the self-supervised model such that the downstream classifier mis-classifies an incoming image as the target category if it contains the attacker chosen trigger. Second, hide the attack by making sure that the accuracy of the downstream classifier is similar to a classifier trained on top of a clean model if the image does not contain the trigger. The attacker achieves these objectives by "data poisoning" where a small trigger patch, fixed beforehand, is pasted on a few images of the target category which is also chosen by the attacker. The hypothesis is that the SSL method associates the trigger with the target category resulting in a successful attack. Given this attack, the defender's goal is to detect and neutralize it without affecting the overall performance of the model. The defender must do this without knowing the trigger or the target category and without any access to trusted data or image labels.

4. Defending with PatchSearch

The goal of PatchSearch is to identify the poisonous samples in the training set and filter them out. We do this by first training an SSL model on the poisonous data and using it to identify the poisonous samples. The cleaned up training data is then used to train a final, backdoor-free model. Note that the architecture and the SSL method used for the final model need not be the same as the model used during defense. In fact, we show that using an easy to attack model during defense allows PatchSearch to find the poisons more easily. There are four components of PatchSearch as illustrated in Figure 1. (1) Cluster the dataset. (2) Assign a poison score to a given image. (3) Efficiently search for a few but highly poisonous samples. (4) Train a poison classifier to find poisons with high recall.

We are given an unlabeled dataset X with images $\{x_i\}_{i=1}^N$. First, we cluster the representations of X into l clusters with *k*-means to get cluster centers and cluster assignments (x_i, y_i) where $y_i \in \{c_1 \dots c_l\}$ is the cluster label (Figure 1 (a)). The resulting clusters are used for locating candidate triggers and to efficiently search for poisons.

Scoring an image. The goal is to quantify the poisonousness of a given image by assigning a poison score to it. The steps for this are described next and illustrated in Figures 1 (b) and 1 (c). We use the cluster centers as classification weights by calculating cosine similarity between x_i and cluster centers, and use y_i as the label to calculate Grad-CAM for x_i . The resulting heatmap is used to extract a candidate trigger t_i which is a $w \times w$ patch that contains the highest heatmap values. We then choose a set of images called flip test set X^f , paste t_i on those, forward through the model and get their cluster assignments. Finally, the poison score of t_i or x_i is the number of images in X^f whose cluster assignments flipped to that of x_i . We ensure that the flip test set is difficult to flip by sampling a few images per cluster that are closest to their respective cluster centers. Note that all flip set images do not need to be clean.

Iterative search. In this phase, the goal is to efficiently find a few but highly poisonous images. Hence, instead of assigning poison scores to the entire set X , we propose a greedy search strategy that focuses on scoring images from poisonous clusters. Concretely, we process s random samples per cluster in an iteration (Figure 1 (d)). At the end of an iteration, each cluster is assigned a poison score as the highest poison score of the images in it. Since poisoned images cluster together, just a single iteration should give us a good estimate of clusters that are likely to contain poisons. We use this insight to focus on highly poisonous cluster by removing (pruning) a fraction r of the clusters with least poison scores in each iteration. The algorithm stops when there are no more clusters to be processed. Finally, all scored images are ranked in the decreasing order of their scores and the top- k ranked images are used in the next step.

Poison classification. Here, the goal is to find as many poisons as possible. We do this by building a classifier h to identify the top- k ranked patches found in the previous step (Figure 1 (e)). Specifically, given the training dataset X and a set P of k patches, we construct poisoned training set X^p by randomly sampling a patch p_j from P and pasting it at a random location on x_i . Samples from X are assigned label “0” (non-poisonous) while samples from X^p are assigned the label “1” (poisonous). However, this labeling scheme is noisy since there are poisoned images in X which will be wrongly assigned the label 0. One way to reduce noise is to not include images in X with high poison scores. Another way is to prevent the model from overfitting to the noise by using strong augmentations and early stopping. However, early stopping requires access to validation data which is unavailable. Hence, we construct a proxy validation set by using the top- k samples with label 1 (poisonous) and bottom- k samples with label 0 (non-poisonous). Note that these samples are excluded from the training set. We stop training the classifier if F1 score on the proxy validation set does not change for 10 evaluation intervals. Finally, the re-

sulting binary classifier is applied to X and any images with “1” (poisonous) label are removed from the dataset.

5. Experiments

We follow the experimentation setup proposed in [44].

Datasets. We use ImageNet-100 dataset [49], which contains images of 100 randomly sampled classes from the 1000 classes of ImageNet [41]. The train set has about 127K samples while the validation set has 5K samples.

Architectures and self-supervised training. In addition to ResNet-18 [27], we also conduct experiments on ViT-B [17] with MoCo-v3 [13] and MAE [25]. Unlike [44], we only consider ResNet-18 with MoCo-v2 [12] and BYOL [23] since we find that older methods like Jigsaw [39] and RotNet [20] have significantly worse clean accuracies compared to MoCo-v2 and BYOL. For the code and parameters, we follow MoCo-v3 [13] for ViT-B, and [44] for ResNet-18. Finally, implementation of i -CutMix and its hyperparameter values follow the official implementation from [33].

Evaluation. We train a linear layer on top of a pre-trained models with 1% of randomly sampled, clean and labeled training set. To account for randomness, we report results averaged over 5 runs with different seeds. We use the following metrics over the validation set to report model performance: accuracy (Acc), count of False Positives (FP) for the target category, and attack success rate (ASR) as the percentage of FP ($FP * 100 / 4950$). Note that maximum FP is 4950 (50 samples/category $\times$ 99 incorrect categories). The metrics are reported for the validation set with and without trigger pasted on it, referred to as Patched Data and Clean Data. See Section A.2 of the Appendix for more details.

Attack. We consider 10 different target categories and trigger pairs used in [44]. Unless mentioned otherwise, all results are averaged over the 10 target categories. Further, following [44], we consider 0.5% and 1.0% poison injection rates which translate to 50% and 100% images being poisoned for a target category. Poisons are generated according to [1, 44] (see Figure 2 for an example). The pasted trigger size is 50x50 ($\approx$ 5% of the image area) and it is pasted at a random location in the image while leaving a margin of 25% on all sides of the image.

PatchSearch. Grad-CAM implementation is used from [21]. Unless mentioned, we run the PatchSearch with following parameters: number of clusters $l = 1000$, patch size $w = 60$, flip test set size = 1000, samples per cluster $s = 2$, and clusters to prune per iteration $r = 25\%$. This results in about 8K (6%) training set images being scored. The architecture of the poison classifier is based on ResNet-18 [27] but uses fewer parameters per layer. We use top-20 poisons to train the classifier, and remove top 10% of poisonous samples to reduce noise in the dataset. We use strong augmentations proposed in MoCo-v2 [12] and randomly resize

Model	top-20 Acc (%)	Rec. (%)	Prec. (%)	Total Rem.
BYOL, ResNet-18, 0.5%	99.5	98.0	58.8	1114.6
MoCo-v3, ViT-B, 0.5%	96.5	98.8	48.3	1567.0
MoCo-v3, ViT-B, 1.0%	97.5	99.0	59.5	2368.3

Table 1. **Poison detection results.** We find that iterative search has high *top-20 Acc* which is the accuracy of finding poisons in top-20 ranked images. In the right 3 columns, we show the effectiveness of the poison classification model at filtering out the poisons. *Rec.* is recall, *Prec.* is precision, while *Total Rem.* is the total number of samples removed by the classifier. We can see that the classifier can identify most of the poisons since recall is $\approx 98\%$. Note that removing a small number of clean images does not significantly damage the final SSL model as shown in Table 2.

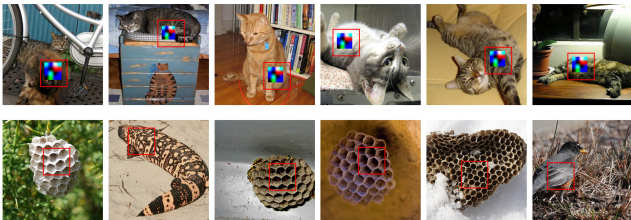


Figure 2. **Top ranked poisonous and non-poisonous images.** We visualize a few top ranked poisonous (1st row) and non-poisonous (2nd row) images found by PatchSearch for BYOL, ResNet-18, poison rate 0.5% and target category Tabby Cat.

the patch between 20x20 and 80x80 (where input is $w \times w$ and default $w = 60$). The classifier is evaluated every 20 iterations with max 2K iterations, and training is stopped if F1 score on the proxy val set does not change for 10 evaluations. Finally, in order to improve the stability of results, we train 5 different classifiers and average their predictions. See Section A.2 in Appendix for more details.

5.1. Main Results

Poison detection results of PatchSearch. We report the results of poison detection for PatchSearch at different stages in Table 1. We can see that at the end of iterative search, most of the top-20 images are indeed poisonous. We can also see that poison classification leads to very high recall with acceptable precision. Finally, we visualize a few top ranked poisonous and non-poisonous images found with PatchSearch in Figure 2.

Effectiveness of PatchSearch and *i*-CutMix. Table 2 lists the results of using *i*-CutMix and PatchSearch for defending against backdoor attacks on 4 different settings. Note that we make it easier for PatchSearch to find poisons by training an easy to attack model during defense, and since *i*-CutMix can suppress the attack, we use models trained without it during PatchSearch. Similarly, we find that the attack is not strong enough to be detectable by PatchSearch for ResNet-18 and MoCo-v2 models (Ta-

ble A5). Therefore, we use the dataset filtered by PatchSearch with ViT-B, MoCo-v3 models. Our results indicate that *i*-CutMix not only improves the clean accuracy of the models but also reduces the attack effectiveness. However, *i*-CutMix cannot fully remove the effect of poisoning and it suppresses the learning of backdoor without explicitly identifying poisons. Here, PatchSearch is a good alternative, which not only significantly suppresses the attack, but also identifies poisoned samples that can be verified by the defender to detect an attack. This may have value in some applications where tracking the source of poisoned samples is important. Finally, the two defenses are complementary to each other, and their combination results in a model that behaves very close to a clean model. PatchSearch removes most poisoned samples while *i*-CutMix suppresses the effect of some remaining ones. Detailed results including per category results, and mean and variance information can be found in Section A.3 in Appendix.

Comparison with trusted data defense. Table 3 compares PatchSearch to the trusted data defense from [44]. Their defense trains a student model from a backdoored teacher with knowledge distillation (KD) [4] on poison-free, trusted data. However, trusted data can be difficult to acquire in SSL and may have its own set of biases. The results in Table 3 show that PatchSearch offers a good trade-off between accuracy and FP despite not using any trusted data. Compared to KD with 25% trusted data, which may not be practical in many cases, PatchSearch is better than KD by ≈ 4 points on clean and patched accuracies at the cost of a slightly higher FP.

Finetuning as a defense. Since the downstream task data is clean and labeled, a very simple defense strategy would be to finetune SSL models on it instead of only training a linear layer. Hence, we compare finetuning with our defense in Table 4. We find that the attack with linear probing increases the FP from 17 to 1708 while the attack with finetuning the whole model increases the FP from 18 to 685. We see that 685 is still much higher than 18, so the attack is still strong with finetuning the whole model. Obviously, one can do finetuning with a very large learning rate to forget the effect of poisoning, but that reduces the accuracy of the model due to overfitting to the downstream task and forgetting the pretraining. Further, with finetuning, different model weights need to be stored for different downstream tasks instead of a single general model, which may not be desirable. Moreover, in settings like few shot learning, the available data can be too little for effective finetuning. Finally, the results also show that all three defenses, PatchSearch, *i*-CutMix, and finetuning, are complementary to each other and can be combined to get a favorable trade-off between accuracy (Acc) and FP (last row of Table 4).

Backdoored MAE vs. defended MoCo-v3. The results from [44] indicate that MAE is robust against back-

Model Type	Clean Data			Patched Data			Clean Data			Patched Data		
	Acc	FP	ASR	Acc	FP	ASR	Acc	FP	ASR	Acc	FP	ASR
<i>ViT-B</i>	<i>MoCo-v3, poison rate 0.5%</i>						<i>MoCo-v3, poison rate 1.0%</i>					
Clean	70.5	18.5	0.4	64.6	27.2	0.5	70.5	18.5	0.4	64.7	27.2	0.5
Clean + <i>i</i> -CutMix	75.6	15.6	0.3	74.4	14.6	0.3	75.6	15.6	0.3	74.4	14.6	0.3
Backdoored	70.6	17.4	0.4	46.9	1708.9	34.5	70.7	14.8	0.3	37.7	2535.9	51.2
Backdoored + <i>i</i> -CutMix	75.6	14.9	0.3	72.2	242.2	4.9	75.5	12.5	0.3	70.2	434.2	8.8
PatchSearch	70.2	23.1	0.5	64.5	39.8	0.8	70.1	43.4	0.9	63.7	76.0	1.5
PatchSearch + <i>i</i> -CutMix	75.2	19.7	0.4	74.2	19.0	0.4	75.0	39.0	0.8	74.0	41.4	0.8
<i>ResNet-18</i>	<i>MoCo-v2, poison rate 0.5%</i>						<i>BYOL, poison rate 0.5%</i>					
Clean	49.7	34.0	0.7	46.5	35.4	0.7	65.7	20.6	0.4	60.6	24.8	0.5
Clean + <i>i</i> -CutMix	55.9	26.3	0.5	54.3	27.1	0.5	65.8	20.5	0.4	63.7	19.6	0.4
Backdoored	50.0	27.2	0.5	40.9	703.3	14.2	66.5	18.6	0.4	35.3	2079.6	42.0
Backdoored + <i>i</i> -CutMix	55.4	24.4	0.5	53.0	124.0	2.5	67.0	18.5	0.4	61.0	365.3	7.4
PatchSearch	49.7	33.1	0.7	46.3	37.2	0.8	66.4	22.9	0.5	61.5	33.4	0.7
PatchSearch + <i>i</i> -CutMix	55.3	30.6	0.6	53.8	29.2	0.6	66.8	22.4	0.5	64.5	22.5	0.5

Table 2. **Defense results.** We compare clean, backdoored, and defended models in different configurations. *Acc* is the accuracy on validation set, *FP* is the number of false positives for the target category, and *ASR* is FP as percentage. All results are averaged over 5 evaluation runs and 10 categories. We observe the following: (1) Presence of *i*-CutMix significantly boosts clean accuracy for models trained with MoCo, but this boost is small for BYOL models. (2) While *i*-CutMix reduces the effectiveness of attacks in all cases, there is still room for improvement. (3) PatchSearch alone can significantly mitigate backdoor attacks. (4) *patched accuracies* of models defended with PatchSearch + *i*-CutMix are very close to clean models in all cases. See Section A.3 for detailed, category-wise results.

Model	Trusted Data	Clean Data		Patched Data	
		Acc	FP	Acc	FP
Clean	100%	49.9	23.0	47.0	22.8
Backdoored	0%	50.1	26.2	31.8	1683.2
KD Defense	25%	44.6	34.5	42.0	37.9
KD Defense	10%	38.3	40.5	35.7	44.8
KD Defense	5%	32.1	41.0	29.4	53.7
PatchSearch	0%	49.4	40.1	45.9	50.3

Table 3. **Comparison with trusted-data defense.** We compare PatchSearch with the trusted data and knowledge distillation (KD) defense [44]. For a fair comparison, we use the setting from [44] without *i*-CutMix: ResNet-18, MoCo-v2, and poison rate 1.0%. All results except the last row are from [44]. We find that despite not relying on any trusted data, our defense can suppress the attack to the level of KD + 5% trusted data. Moreover, our model has significantly higher accuracies: 49.4% vs 32.1% on clean data and 45.9% vs 29.4% on patched data. Compared to KD + 25% trusted data, our model has higher accuracies with only slightly higher FP.

door attacks. However, there are a few confounding factors such as finetuning and differences in model architecture which need more investigation. We compare finetuned ViT-B models trained with MoCo-v3 and MAE in Table 5. We find that MoCo-v3 defended with PatchSearch and *i*-CutMix is better than MAE both in terms of Acc and FP for 1% labeled finetuning data, but MAE quickly catches up in

Evaluation	Clean Data		Patched Data	
	Acc	FP	Acc	FP
<i>Backdoored</i>				
Linear	70.6	17.4	46.9	1708.9
Finetuned	72.2	18.2	63.3	685.8
<i>PatchSearch + i-CutMix</i>				
Linear	75.2	19.7	74.2	19.0
Finetuned	78.1	16.9	76.6	16.8

Table 4. **Finetuning as defense.** We explore finetuning an entire model on the downstream dataset as a defense. While not as good as PatchSearch, we find that finetuning can mitigate the attack. However, the model loses its appeal as a general-purpose backbone and becomes specific to a single downstream task. Finally, note that finetuning is complementary to PatchSearch and *i*-CutMix. Setting: ViT-B, MoCo-v3, and poison rate 0.5%.

the 10% regime. It is possible that these observations are specific to our dataset or model training setup. Hence, we show a similar pattern even for officially pre-trained [2, 3] (on ImageNet-1000) models in Table 6.

Attack on a downstream task. Following [44], so far we have only considered the case where evaluation dataset is a subset of the training dataset. However, SSL models are intended to be used for downstream datasets that are different from the pre-training dataset. Hence, we show in Table 7 that it is possible to attack one of the categories used

Method	Clean Data		Patched Data	
	Acc	FP	Acc	FP
<i>Finetuned with 1% labeled data</i>				
MAE	65.7	18.7	53.8	97.6
MoCo-v3	78.2	20.2	76.8	17.1
<i>Finetuned with 10% labeled data</i>				
MAE	84.0	9.5	75.6	40.1
MoCo-v3	84.9	11.5	83.0	10.8

Table 5. **Backdoored MAE vs. defended MoCo-v3.** While the inherent robustness of MAE to backdoor attacks [44] is appealing, we show that a properly defended MoCo-v3 model has better clean accuracy and lower FP compared with a backdoored MAE. Setting: ViT-B, poison rate 0.5%, and average of 4 target categories (Rottweiler, Tabby Cat, Ambulance, and Laptop).

Method	1%	Finetuning Data		
		10%	100%	100%
MAE	51.5	73.0	83.5	83.6 [25]
MoCo-v3	64.5	74.2	83.0	83.2 [25]

Table 6. **MoCo-v3 is better than MAE in finetuning on less data.** It is possible that observations in Table 5 are specific to our dataset or model training setup. Hence, we show that finetuning official, unlabeled ImageNet-1K pre-trained, ViT-B models for MAE and MoCo-v3 still results in MoCo-v3 outperforming MAE in the 1% regime. However, MAE catches up in 10% and 100% regimes. Note that only this table uses the ImageNet-1K dataset.

Model	Clean Data		Patched Data	
	Acc	FP	Acc	FP
Clean	47.6	33.8	39.4	28.0
Backdoored	47.9	27.5	12.4	3127.3
PatchSearch	47.4	35.0	40.7	32.3

Table 7. **Attack on downstream Food101 classification.** We choose a target category from Food101 [6] dataset, poison 700 of its images, and add the poisoned images to the pre-training dataset (ImageNet-100). The models are evaluated on Food101 that has 5050 val images in total (50 val images per class $\times$ 101 classes). We find that the attack is successful and also that PatchSearch can defend against the attack. Setting: BYOL & ResNet-18. See Table A7 of Appendix for detailed results.

in the downstream task of Food101 [6] classification. See Section A.1 and Table A7 of Appendix for more details.

5.2. Ablations

***i*-CutMix.** It is interesting that a method as simple as *i*-CutMix can be a good defense. Therefore, we attempt to understand which components of *i*-CutMix make it a good defense. *i*-CutMix has two components: creating a composite image by pasting a random crop from a random image, and replacing the original one-hot loss with a weighted one. The weights are in proportion to the percentages of the

Experiment	Clean Data		Patched Data	
	Acc	FP	Acc	FP
No <i>i</i> -CutMix	70.5	24.8	42.9	1412.4
<i>i</i> -CutMix	75.6	28.0	70.5	268.4
+ remove weighted loss	71.4	26.4	68.7	104.2
+ paste black box	71.5	24.6	42.3	1333.0
+ paste random noise	70.5	24.2	46.7	1480.0

Table 8. **Analysis of *i*-CutMix.** Removing weighted loss means that we use one-hot loss instead of the weighted one. We find that simply pasting crops from other images helps greatly, while the weighted loss term can additionally improve the accuracy. We believe the reason is that random crop pasting forces the model to make decisions based on the entire image instead of small patches. We attempt to visualize this in Figure 3. Note that pasting black box is effectively CutOut augmentation [15] and pasting gaussian noise box is effectively RandomErase augmentation [59]. Setting: ViT-B, MoCo-v3, poison rate 0.5%, target category Rottweiler.

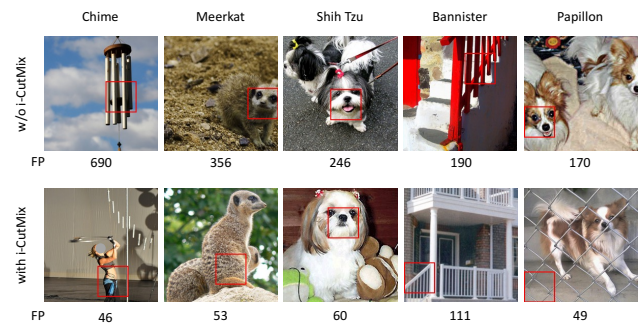


Figure 3. **Effect of *i*-CutMix on potency of patches.** We hypothesize that *i*-CutMix encourages the model to use global features rather than focusing on small regions. To explore this, we use PatchSearch as an interpretation algorithm and analyze the impact of *i*-CutMix on clean patches. We use a *clean* model to calculate FP of all clean training images with PatchSearch (no iterative search) and use the resulting FP as potency score. FP for each category is calculated as the max FP of all patches in it (shown below images). We find that *i*-CutMix suppresses the impact of individual patches (resulting in smaller potency score), which is aligned with our hypothesis. The total potency (sum of FP) of patches from all categories is reduced by 71.3% for BYOL, ResNet-18 and 45.4% for MoCo-v3, ViT-B. We show images for five random categories with and without *i*-CutMix (BYOL, ResNet-18).

mixed images. We observe the effect of these components in Table 8. We can also get an insight into the working of *i*-CutMix by using PatchSearch as an interpretation algorithm. See Figure 3.

Effect of *w*. An important hyperparameter of PatchSearch is *w* which controls the size of the extracted candidate trigger. Since the defender does not know the size of the trigger, wrong choice of *w* can lead to poor poison detection performance as shown in Figure 4. However, it is possible to automate the process of finding a good value for *w*. We know that if the choice of *w* is wrong, most poison

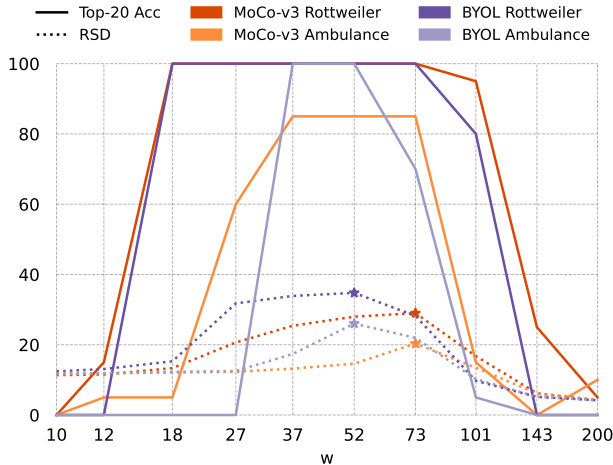


Figure 4. **Effect of w .** We plot the top-20 Acc against w values sampled uniformly in log-space from 10 to 200. We find highest accuracies for w near 50 which is the actual trigger size. However, the defender does not know the actual trigger size. Hence, we propose relative standard deviation (RSD) of the population of poison scores as a heuristic to pick w . RSD is standard deviation divided by the mean. In all cases, w with highest RSD (marked with $\star$ in the figure) also corresponds to the highest top-20 accuracy.

scores will be similar. Hence, the w resulting in high variance in poisons scores should be closer to actual trigger size used by the attacker. However, we need to adjust the scale of the variance with mean since different w values result in populations with different scales of scores. Hence, we use relative standard deviation (RSD¹), which is standard deviation divided by the mean. Figure 4 shows that picking the w with highest RSD also results in highest top-20 accuracy. Therefore, a defender can try different values of w and pick the one with highest RSD without making any assumptions about the trigger size.

Other ablations. In Table 9, we explore a stronger form of attack, where the trigger is pasted five times instead of once on target category images in training. This results in the trigger occupying roughly 25% area of the image. In Table 10, we consider the effect of applying PatchSearch to clean data since a defender does not know whether a dataset is poisoned. We see only a small degradation in accuracy. In Figure 5, we consider the effect of removing the poison classifier, and find that resulting image ranking shows poor precision vs. recall trade-off in some cases. See Section A.1 of Appendix for more ablations and details.

6. Conclusion

We introduced a novel defense algorithm called PatchSearch for defending self-supervised learning against patch based data poisoning backdoor attacks. PatchSearch iden-

Model	Clean Data		Patched Data	
	Acc	FP	Acc	FP
Clean	70.5	18.5	64.6	27.2
Backdoored	70.7	21.1	42.0	1996.7
PatchSearch	70.3	25.2	64.7	42.0

Table 9. **Repeated patches during attack.** We consider a stronger attack, where five triggers instead of one are pasted on target category images in training. We find that PatchSearch successfully suppresses the attack. Setting: ViT-B, MoCo-v3, poison rate 0.5%.

Model	Clean Data		Patched Data	
	Acc	FP	Acc	FP
Clean without PatchSearch	70.5	18.5	64.6	27.2
Clean with PatchSearch	69.2	20.0	63.2	33.3

Table 10. **Effect of defending a clean model.** Applying PatchSearch on a clean dataset will remove a small set of clean images (6K or about 5% of the training set), which results in only a small reduction in accuracy. Setting: ViT-B and MoCo-v3.

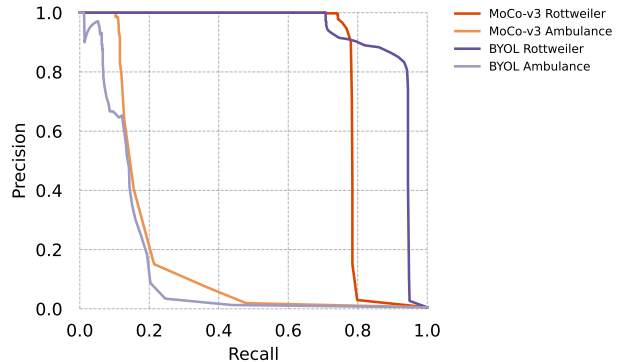


Figure 5. **Effect of removing poison classifier.** We demonstrate the need for the poison classifier by removing it from PatchSearch, and analyzing the results of images ranked by poison score. Despite scoring every image in the dataset with PatchSearch (no iterative search), poisons cannot be precisely detected in some cases.

tifies poisons and removes them from the training set. We also show that the augmentation strategy of *i*-CutMix is a good baseline defense. We find that PatchSearch is better than *i*-CutMix and the SOTA defense that uses trusted data. Further, PatchSearch and *i*-CutMix are complementary to each and their combination improves the model performance while effectively mitigating the attack. However, our defense assumes that the trigger is patch-based and it smaller than the objects of interest. These assumptions may limit its effectiveness against future attacks. We hope that our paper will encourage the development of better backdoor attacks and defense methods for SSL models.

Acknowledgement. This work is partially supported by DARPA Contract No. HR00112190135, HR00112290115, and FA8750-19-C-0098, NSF grants 1845216 and 1920079,

¹https://en.wikipedia.org/wiki/Coefficient_of_variation

NIST award 60NANB18D279, and also funding from Shell Inc., and Oracle Corp. We would also like to thank K L Navaneet and Aniruddha Saha for many helpful discussions.

References

- [1] Code for Backdoor Attacks on Self-Supervised Learning paper. <https://github.com/UMBCvision/SSL-Backdoor>. 4
- [2] MAE models and code. <https://github.com/facebookresearch/mae>. 6
- [3] MoCo-v3 models and code. <https://github.com/facebookresearch/moco-v3>. 6
- [4] Soroush Abbasi Koohpayegani, Ajinkya Tejankar, and Hamed Pirsiavash. Compress: Self-supervised learning by compressing representations. *Advances in Neural Information Processing Systems*, 33:12980–12992, 2020. 2, 5
- [5] Eitan Borgnia, Valeriia Cherepanova, Liam Fowl, Amin Ghiasi, Jonas Geiping, Micah Goldblum, Tom Goldstein, and Arjun Gupta. Strong data augmentation sanitizes poisoning and backdoor attacks without an accuracy tradeoff. In *ICASSP 2021-2021 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, pages 3855–3859. IEEE, 2021. 1, 3
- [6] Lukas Bossard, Matthieu Guillaumin, and Luc Van Gool. Food-101 – mining discriminative components with random forests. In *European Conference on Computer Vision (ECCV)*, 2014. 7, 12, 15
- [7] Tom B Brown, Dandelion Mané, Aurko Roy, Martín Abadi, and Justin Gilmer. Adversarial patch. *arXiv preprint arXiv:1712.09665*, 2017. 3
- [8] Nicholas Carlini and Andreas Terzis. Poisoning and backdoorizing contrastive learning. In *International Conference on Learning Representations (ICLR)*, 2022. 1, 3
- [9] Mathilde Caron, Ishan Misra, Julien Mairal, Priya Goyal, Piotr Bojanowski, and Armand Joulin. Unsupervised learning of visual features by contrasting cluster assignments. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2020. 3
- [10] Bryant Chen, Wilka Carvalho, Nathalie Baracaldo, Heiko Ludwig, Benjamin Edwards, Taesung Lee, Ian Molloy, and Biplav Srivastava. Detecting backdoor attacks on deep neural networks by activation clustering. *arXiv preprint arXiv:1811.03728*, 2018. 3
- [11] Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey Hinton. A simple framework for contrastive learning of visual representations. *arXiv preprint arXiv:2002.05709*, 2020. 3
- [12] Xinlei Chen, Haoqi Fan, Ross Girshick, and Kaiming He. Improved baselines with momentum contrastive learning. *arXiv preprint arXiv:2003.04297*, 2020. 3, 4
- [13] Xinlei Chen, Saining Xie, and Kaiming He. An empirical study of training self-supervised vision transformers. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 9640–9649, 2021. 4
- [14] Edward Chou, Florian Tramer, and Giancarlo Pellegrino. Sentinet: Detecting localized universal attacks against deep learning systems. In *2020 IEEE Security and Privacy Workshops (SPW)*, pages 48–54. IEEE, 2020. 1, 2, 3
- [15] Terrance DeVries and Graham W Taylor. Improved regularization of convolutional neural networks with cutout. *arXiv preprint arXiv:1708.04552*, 2017. 7
- [16] Bao Gia Doan, Ehsan Abbasnejad, and Damith C Ranasinghe. Februus: Input purification defense against trojan attacks on deep neural network systems. In *Annual Computer Security Applications Conference*, pages 897–912, 2020. 1, 3
- [17] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *International Conference on Learning Representations*, 2021. 4
- [18] Alexey Dosovitskiy, Jost Tobias Springenberg, Martin Riedmiller, and Thomas Brox. Discriminative unsupervised feature learning with convolutional neural networks. In *Advances in neural information processing systems (NeurIPS)*, 2014. 2
- [19] Yansong Gao, Change Xu, Derui Wang, Shiping Chen, Damith C Ranasinghe, and Surya Nepal. Strip: A defence against trojan attacks on deep neural networks. In *Proceedings of the 35th Annual Computer Security Applications Conference*, 2019. 3
- [20] Spyros Gidaris, Praveer Singh, and Nikos Komodakis. Unsupervised representation learning by predicting image rotations. In *International Conference on Learning Representations (ICLR)*, 2018. 2, 4
- [21] Jacob Gildenblat and contributors. Pytorch library for cam methods. <https://github.com/jacobgil/pytorch-grad-cam>, 2021. 4
- [22] Priya Goyal, Mathilde Caron, Benjamin Lefaudeaux, Min Xu, Pengchao Wang, Vivek Pai, Mannat Singh, Vitaliy Liptchinsky, Ishan Misra, Armand Joulin, et al. Self-supervised pretraining of visual features in the wild. *arXiv preprint arXiv:2103.01988*, 2021. 1
- [23] Jean-Bastien Grill, Florian Strub, Florent Altché, Corentin Tallec, Pierre H Richemond, Elena Buchatskaya, Carl Doherty, Bernardo Avila Pires, Zhaohan Daniel Guo, Mohammad Gheshlaghi Azar, et al. Bootstrap your own latent: A new approach to self-supervised learning. *arXiv preprint arXiv:2006.07733*, 2020. 3, 4
- [24] Tianyu Gu, Brendan Dolan-Gavitt, and Siddharth Garg. Badnets: Identifying vulnerabilities in the machine learning model supply chain. *arXiv preprint arXiv:1708.06733*, 2017. 3
- [25] Kaiming He, Xinlei Chen, Saining Xie, Yanghao Li, Piotr Dollár, and Ross Girshick. Masked autoencoders are scalable vision learners. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 16000–16009, 2022. 4, 7, 12
- [26] Kaiming He, Haoqi Fan, Yuxin Wu, Saining Xie, and Ross Girshick. Momentum contrast for unsupervised visual representation learning. In *Proceedings of the IEEE/CVF Confer-*

- ence on Computer Vision and Pattern Recognition (CVPR), 2020. 2
- [27] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016. 4
- [28] Sanghyun Hong, Varun Chandrasekaran, Yiğitcan Kaya, Tudor Dumitras, and Nicolas Papernot. On the effectiveness of mitigating data poisoning attacks with gradient shaping. *arXiv preprint arXiv:2002.11497*, 2020. 1, 3
- [29] Kunzhe Huang, Yiming Li, Baoyuan Wu, Zhan Qin, and Kui Ren. Backdoor defense via decoupling the training process. In *International Conference on Learning Representations*, 2022. 3
- [30] Shanjiaoyang Huang, Weiqi Peng, Zhiwei Jia, and Zhuowen Tu. One-pixel signature: Characterizing cnn models for backdoor detection. In *European Conference on Computer Vision*, pages 326–341. Springer, 2020. 3
- [31] Xijie Huang, Moustafa Alzantot, and Mani Srivastava. Neuroninspect: Detecting backdoors in neural networks via output explanations. *arXiv preprint arXiv:1911.07399*, 2019. 3
- [32] Soheil Kolouri, Aniruddha Saha, Hamed Pirsiavash, and Heiko Hoffmann. Universal litmus patterns: Revealing backdoor attacks in cnns. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 301–310, 2020. 3
- [33] Kibok Lee, Yian Zhu, Kihyuk Sohn, Chun-Liang Li, Jinwoo Shin, and Honglak Lee. *i*-mix: A domain-agnostic strategy for contrastive representation learning. In *International Conference on Learning Representations*, 2021. 1, 2, 3, 4
- [34] Yiming Li, Yong Jiang, Zhifeng Li, and Shu-Tao Xia. Backdoor learning: A survey. *IEEE Transactions on Neural Networks and Learning Systems*, 2022. 1, 3
- [35] Yige Li, Xixiang Lyu, Nodens Koren, Lingjuan Lyu, Bo Li, and Xingjun Ma. Anti-backdoor learning: Training clean models on poisoned data. *Advances in Neural Information Processing Systems*, 34:14900–14912, 2021. 3
- [36] Kang Liu, Brendan Dolan-Gavitt, and Siddharth Garg. Fine-pruning: Defending against backdooring attacks on deep neural networks. In *International Symposium on Research in Attacks, Intrusions, and Defenses*, pages 273–294. Springer, 2018. 3
- [37] Yingqi Liu, Shiqing Ma, Yousra Aafer, Wen-Chuan Lee, Juan Zhai, Weihang Wang, and Xiangyu Zhang. Trojaning attack on neural networks. 2017. 3
- [38] Yuntao Liu, Yang Xie, and Ankur Srivastava. Neural trojans. In *2017 IEEE International Conference on Computer Design (ICCD)*, 2017. 3
- [39] Mehdi Noroozi and Paolo Favaro. Unsupervised learning of visual representations by solving jigsaw puzzles. In *European Conference on Computer Vision (ECCV)*, 2016. 2, 4
- [40] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. In *International Conference on Machine Learning*, pages 8748–8763. PMLR, 2021. 3
- [41] Olga Russakovsky, Jia Deng, Hao Su, Jonathan Krause, Sanjeev Satheesh, Sean Ma, Zhiheng Huang, Andrej Karpathy, Aditya Khosla, Michael Bernstein, Alexander C. Berg, and Li Fei-Fei. ImageNet Large Scale Visual Recognition Challenge. *International Journal of Computer Vision (IJCV)*, 115(3):211–252, 2015. 4
- [42] Aniruddha Saha, Akshayvarun Subramanya, Koninika Patil, and Hamed Pirsiavash. Role of spatial context in adversarial robustness for object detection. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops*, pages 784–785, 2020. 3
- [43] Aniruddha Saha, Akshayvarun Subramanya, and Hamed Pirsiavash. Hidden trigger backdoor attacks. In *Proceedings of the AAAI Conference on Artificial Intelligence*, 2020. 3
- [44] Aniruddha Saha, Ajinkya Tejankar, Soroush Abbasi Koohpayegani, and Hamed Pirsiavash. Backdoor attacks on self-supervised learning. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2022. 1, 2, 3, 4, 5, 6, 7
- [45] Ramprasaath R Selvaraju, Michael Cogswell, Abhishek Das, Ramakrishna Vedantam, Devi Parikh, and Dhruv Batra. Grad-cam: Visual explanations from deep networks via gradient-based localization. In *Proceedings of the IEEE international conference on computer vision*, pages 618–626, 2017. 1, 3
- [46] Ali Shafahi, W Ronny Huang, Mahyar Najibi, Octavian Suciu, Christoph Studer, Tudor Dumitras, and Tom Goldstein. Poison frogs! targeted clean-label poisoning attacks on neural networks. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2018. 3
- [47] Koohpayegani Soroush Abbasi, Ajinkya Tejankar, and Hamed Pirsiavash. Mean shift for self-supervised learning. In *International Conference on Computer Vision (ICCV)*, 2021. 2, 3
- [48] Ajinkya Tejankar, Soroush Abbasi Koohpayegani, Vipin Pillai, Paolo Favaro, and Hamed Pirsiavash. Isd: Self-supervised learning by iterative similarity distillation. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, 2021. 3
- [49] Yonglong Tian, Dilip Krishnan, and Phillip Isola. Contrastive multiview coding. In *European conference on computer vision*, pages 776–794. Springer, 2020. 4
- [50] Yonglong Tian, Chen Sun, Ben Poole, Dilip Krishnan, Cordelia Schmid, and Phillip Isola. What makes for good views for contrastive learning? In *Advances in Neural Information Processing Systems (NeurIPS)*, 2020. 3
- [51] Brandon Tran, Jerry Li, and Aleksander Madry. Spectral signatures in backdoor attacks. *Advances in neural information processing systems (NeurIPS)*, 2018. 3
- [52] Alexander Turner, Dimitris Tsipras, and Aleksander Madry. Clean-label backdoor attacks. 2018. 3
- [53] Bolun Wang, Yuanshun Yao, Shawn Shan, Huiying Li, Bimal Viswanath, Haitao Zheng, and Ben Y Zhao. Neural cleanse: Identifying and mitigating backdoor attacks in neural networks. In *2019 IEEE Symposium on Security and Privacy (SP)*, pages 707–723. IEEE, 2019. 3
- [54] Zhirong Wu, Yuanjun Xiong, Stella X Yu, and Dahua Lin. Unsupervised feature learning via non-parametric instance

- discrimination. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, 2018. 2
- [55] Sangdoo Yun, Dongyoon Han, Seong Joon Oh, Sanghyuk Chun, Junsuk Choe, and Youngjoon Yoo. Cutmix: Regularization strategy to train strong classifiers with localizable features. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 6023–6032, 2019. 1, 3
- [56] Yi Zeng, Won Park, Z Morley Mao, and Ruoxi Jia. Rethinking the backdoor attacks’ triggers: A frequency perspective. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 16473–16481, 2021. 3
- [57] Richard Zhang, Phillip Isola, and Alexei A Efros. Colorful image colorization. In *Proceedings of the European Conference on Computer Vision (ECCV)*, 2016. 2
- [58] Songzhu Zheng, Yikai Zhang, Hubert Wagner, Mayank Goswami, and Chao Chen. Topological detection of trojaned neural networks. *Advances in Neural Information Processing Systems*, 34:17258–17272, 2021. 3
- [59] Zhun Zhong, Liang Zheng, Guoliang Kang, Shaozi Li, and Yi Yang. Random erasing data augmentation. In *Proceedings of the AAAI conference on artificial intelligence*, volume 34, pages 13001–13008, 2020. 7

A. Appendix

A.1. Additional Ablations

We consider the effect of number of images scored by iterative search in Table A1 and Figure A1. Number of scored images can be changed by varying two hyperparameters s (number of samples scored per cluster) and r (percentage of least poisonous clusters removed in each iteration). The results show consistently good performance when processing more than 6% of samples. Next, we consider the effect of varying the size of flip test set X^f in Table A2. We find that we can get reasonable results with $|X^f|$ as small as 316. Next, we evaluate intermediate checkpoints of one of our models to understand the relationship between overall performance of the model (measured with clean data Acc) and the attack effectiveness (measured with patched data FP) in Table A3. We can see that attack effectiveness improves as the overall model performance improves in the early epochs. In later epochs, the attack effectiveness fluctuates but is still high relative to earlier epochs. This indicates that it is possible to reduce the compute requirements for the model used during defense. We only need to train this model until the attack effectiveness becomes broadly comparable to a fully trained model. Next, we show that our defense is effective even when changing datasets in Table A4. Next, we consider the effect of varying top- k in Table A5, and find that any top- $k \leq 20$ is a good choice. Next, we consider changing the trigger size w during the attack in Table A6. We find that our defense works despite changing w from 50 to 75 and 100. Next, we consider attacking images from a downstream dataset instead of pre-training dataset. We use Food101 [6] classification dataset for this purpose. It consists of 101 fine-grained food categories and 750 images per category. For each category, we randomly select 700 images for training and 50 images for validation. We backdoor BYOL, ResNet-18 models by picking random category from Food101, poisoning all of its 700 training images, and adding them to the pre-training dataset (ImageNet-100). The resulting pre-trained model is evaluated on the task of Food101 classification. We use the same linear evaluation setup as other ResNet-18 models in our experiments. We use the training set of Food101 (700×101) for training the linear layer while the evaluation set (50×101) is used for evaluation. The results are presented in Table A7. We find that the attack is successful, but the models can also be successfully defended with PatchSearch.

A.2. Implementation Details

Below, we describe implementation details for various settings. Note that running PatchSearch is relatively inexpensive compared to the pre-training stage. For instance, with ImageNet-100 (126K images) and ViT-B on 4x3090

GPUs, PatchSearch takes 1.5 hrs, which is small as compared to the training time, about 16 hrs, for MoCo-v3.

Poison classifier training in PatchSearch. The training has following parameters: SGD (lr=0.01, batch size=32, max iterations=2000, weight decay= $1e^{-4}$, and cosine lr scheduler). The architecture of the classifier is ResNet-18 but each layer has only a single BasicBlock instead of the default two².

ResNet-18 model training. As mentioned in the main paper, MoCo-v2 and BYOL training for ResNet-18 models is exactly the same as code³ from [41]. The models are trained on 4 NVIDIA A100 GPUs.

MoCo-v3 model training. We use the code⁴ from MoCo-v3 [11] for training ViT-B models. We use the default hyperparameters from the code except SGD (batch size=1024, epochs=200). The models are trained on 8 NVIDIA A100 GPUs.

MAE model training. We use the code⁵ from MAE [23] to train ViT-B models. All hyperparameters are unchanged except following: SGD (batch size=32 and accum iter=4). Here, *accum iter* refers to the number iterations used for averaging the gradients before updating parameters. The models are trained on 8 NVIDIA A100 GPUs. Hence, the effective batch size is $32 \times 8 \times 4 = 1024$.

ResNet-18 linear evaluation. We use the linear layer training procedure proposed in code⁶ from ComPress [4] for evaluating ResNet-18 models. A single linear layer is trained on top of a frozen backbone. The output of the backbone is processed according to following steps before passing it to the linear layer. (1) A mini-batch of features is normalized to have unit l_2 norm. (2) The mini-batch is also normalized to have zero mean and unit variance. Note that the mean and variance used for normalization in the second step comes from the l_2 normalized features of the entire training dataset. All hyperparameters values are set to default values from the original code.

ViT-B linear evaluation. We use the linear layer training procedure from MoCo-v3 [11] code⁷. We set all hyperparameters to their default values from the code except SGD (epochs=30, batch size=256).

ViT-B fine-tuning evaluation. We use the code⁸ from MAE [25] for fine-tuning ViT-B models. Strong augmentations like mixup, random erase, and cutmix are turned off during fine-tuning since we find that their presence hurts the

²<https://github.com/pytorch/vision/blob/main/torchvision/models/resnet.py>

³<https://github.com/UMBCvision/SSL-Backdoor>

⁴<https://github.com/facebookresearch/moco-v3>

⁵<https://github.com/facebookresearch/mae>

⁶https://github.com/UMBCvision/CompPress/blob/master/eval_linear.py

⁷https://github.com/facebookresearch/moco-v3/blob/main/main_lincls.py

⁸https://github.com/facebookresearch/mae/blob/main/engine_finetune.py

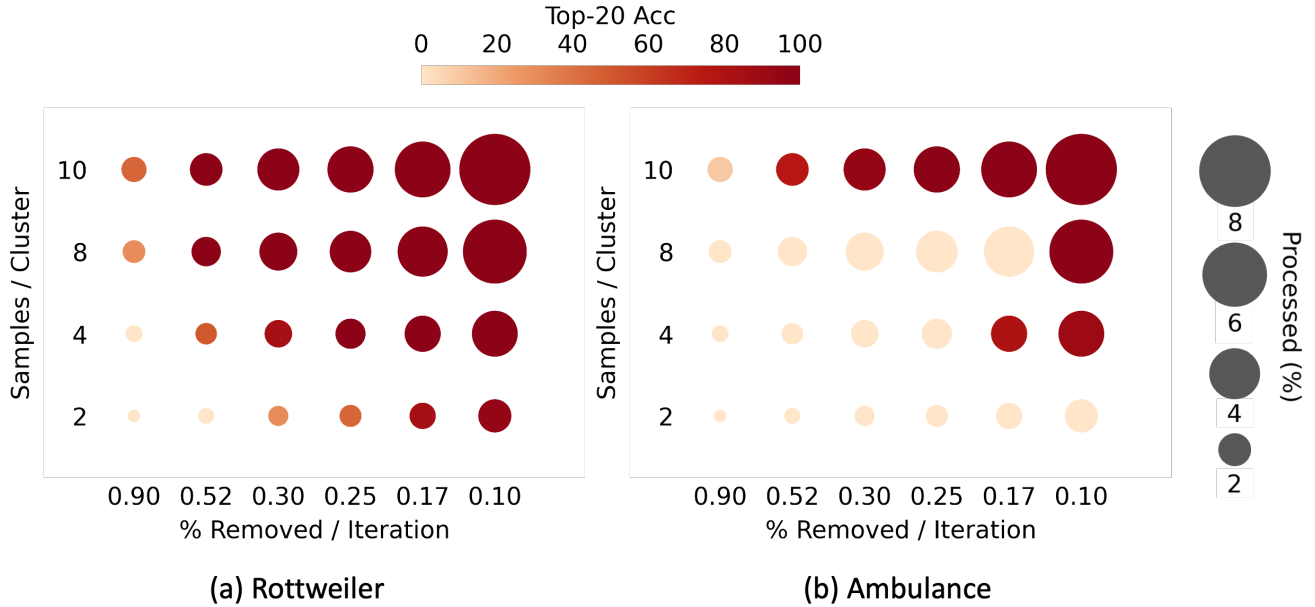


Figure A1. **Effect of number of images scored.** We vary different hyperparameters that control the number of images scored by Patch-Search and observe the effect on the accuracy of finding poisons in top-20 ranked images. The color of the circles denotes the accuracy while their size denotes the number of scored images. We find that a large value for samples per cluster (s) has better performance compared to small s values with comparable number of processed images. Finally, processing more than 6% of images results in a good model performance. See Table A1 for detailed results. Setting: MoCo-v3, ViT-B, and poison rate 0.5%.

Num. Clusters	Samples Per Cluster	Pruned per iteration (%)					
		0.90	0.52	0.30	0.25	0.17	0.10
100	2	0.0 / 0.18	0.0 / 0.31	0.0 / 0.54	0.0 / 0.70	0.0 / 0.97	0.0 / 1.60
100	4	0.0 / 0.35	0.0 / 0.62	0.0 / 1.07	0.0 / 1.30	80.0 / 1.94	90.0 / 3.21
100	8	0.0 / 0.70	0.0 / 1.24	0.0 / 2.14	0.0 / 2.60	0.0 / 3.87	100.0 / 6.42
100	10	10.0 / 0.88	75.0 / 1.55	95.0 / 2.68	100.0 / 3.30	100.0 / 4.84	100.0 / 8.02
1000	2	20.0 / 1.80	45.0 / 3.00	75.0 / 5.30	80.0 / 6.40	100.0 / 9.40	100.0 / 16.00
1000	4	65.0 / 3.50	100.0 / 6.10	100.0 / 10.60	95.0 / 12.70	100.0 / 18.60	100.0 / 30.90
1000	8	100.0 / 7.00	100.0 / 12.20	100.0 / 21.00	100.0 / 25.00	100.0 / 35.40	100.0 / 52.30
1000	10	100.0 / 8.80	100.0 / 15.20	100.0 / 26.10	100.0 / 30.80	100.0 / 42.40	100.0 / 59.30

Table A1. **Effect of varying scored count.** We explore the effect of processing different amount of images by varying number of clusters, s (samples per cluster) and r (percentage of clusters pruned per iteration). Each table entry has the format *accuracy of finding poisons (%) in top-20 / percentage of training set scored*. We find that even with large r and a small number of clusters, increasing s can improve the model performance. Finally, we observe that processing more than 6% of the training set results in good model performance most of the times. Setting: MoCo-v3, ViT-B, poison rate 0.5%, target category Ambulance.

overall performance of the model (Table A8). For MoCo-v3 models, fine-tuning runs for 30 epochs while for MAE models it runs for 90 epochs. Finally, MAE uses global pooling of tokens following the default option in the original code while MoCo-v3 models use $[\text{CLS}]$ token. Rest of the hyperparameters are unchanged from their default values. For fine-tuning results on ImageNetx-1k reported in Table 6, all settings except epochs=50 are the same as their default values. Note that strong augmentations is kept on as in the default settings in order to be comparable to the numbers

reported in MAE [23].

A.3. Per Category Results

We list per category results for Table 2. The results are in Tables A9, A10, A11, and A12.

		Target Category	Flip test set size				
			10	32	100	316	1000
top-20 Acc (%)	Ambulance		5.0	20.0	60.0	80.0	80.0
	Rottweiler		55.0	100.0	100.0	100.0	100.0

Table A2. **Effect of flip test set size.** We explore the effect of changing the flip test set size $|X^f|$. The larger this set the more diverse samples will be used to obtain the poison score for an image. We find that PatchSearch is not greatly sensitive to this hyperparameter and even a small value like 316 could be effective. Setting: ViT-B, MoCo-v3, and poison rate 0.5%.

Checkpoint Epoch	Clean Data		Patched Data		Clean Data		Patched Data	
	Acc	FP	Acc	FP	Acc	FP	Acc	FP
<i>Rottweiler</i>					<i>Ambulance</i>			
20	27.2	58.8	24.9	72.8	26.9	47.8	24.7	67.6
40	43.6	47.4	36.4	391.8	43.3	20.6	39.6	66.2
60	53.1	36.6	17.1	3145.0	53.0	14.0	46.5	284.6
80	58.3	32.8	35.5	1787.4	58.0	13.4	51.2	262.4
100	60.9	28.8	28.0	2689.6	61.4	12.6	48.0	1041.6
120	62.9	27.6	35.5	2118.4	63.1	12.4	51.6	804.2
140	64.2	30.0	40.7	1675.0	64.6	9.2	51.4	908.6
160	65.8	26.4	43.5	1420.4	66.1	9.8	54.8	704.6
180	66.3	25.8	40.9	1895.8	66.9	7.8	54.4	800.6
200	66.4	26.0	39.9	1926.8	66.8	7.6	53.6	895.6

Table A3. **Relationship between overall performance and attack effectiveness.** We evaluate intermediate checkpoints to understand the relationship between overall model performance (measured with *Clean Data Acc*) and the attack effectiveness (measured with *Patched Data FP*). We find that attack effectiveness has a strong correlation with overall model performance in early epochs (≤ 60). While this correlation is not strict for later epochs, the attack effectiveness maintains its overall magnitude. This observation suggests that a model used during defense need not be trained to convergence. We only need to train it until the attack effectiveness is comparable to the fully trained model in overall magnitude. Setting: BYOL, ResNet-18, and poison rate 0.5%.

Model	Clean Data			Patched Data		
	Acc	FP	ASR	Acc	FP	ASR
Clean	79.2	205	2.1	54.3	741	7.4
Backdoored	79.0	206	2.1	26.5	6187	61.9
PatchSearch	78.7	224	2.2	57.0	816	8.2

Table A4. **CIFAR-10.** We show results for the attack and defense on CIFAR-10. We find that the defense is successfully able to mitigate the attack. In fact, we find that the poison classifier is not needed since the iterative search itself is sufficient. Simply removing the top 10% of the clusters removes 100% of poisons. To account for smaller, 32x32, images, we also reduce the size of the trigger to 8x8. Setting: ResNet-18, BYOL, poison rate 0.5% and target category Airplane.

Model	Acc (%) in top- k				
	5	10	20	50	100
ResNet-18, BYOL, 0.5%	100.0	100.0	99.5	93.8	84.0
ResNet-18, MoCo-v2, 0.5%	52.0	55.0	52.5	44.6	30.9
ViT-B, MoCo-v3, 0.5%	100.0	100.0	96.5	87.6	73.5
ViT-B, MoCo-v3, 1.0%	98.0	98.0	97.5	84.0	77.3

Table A5. **Effect of top- k in PatchSearch.** Note that PatchSearch is not able to find the patches for ResNet-18, MoCo-v2 for any k . Hence, we use an easy to backdoor model like ViT-B to defend it.

Model	Clean Data			Patched Data		
	Acc	FP	ASR	Acc	FP	ASR
$w = 75$						
Clean	65.7	28.8	0.6	55.0	19.0	0.4
Backdoored	66.6	32.4	0.7	36.4	1588.0	32.0
PatchSearch	66.5	29.2	0.6	55.8	25.2	0.5
$w = 100$						
Clean	65.7	28.8	0.6	45.6	15.2	0.3
Backdoored	65.8	30.2	0.6	20.4	2481.6	50.6
PatchSearch	65.6	32.8	0.7	46.7	34.6	0.7

Table A6. **Changing trigger size w in the attack.** We explore the effect of changing the trigger size w during the attack. We find that our defense can successfully defend against different w by searching for most effective w during the defense (outlined in Figure 4). Note that since we are dealing with large w we paste the trigger anywhere in the image without any boundary. This is different from the default setting of 25% margin on all sides. Setting: ResNet-18, BYOL, poison rate 0.5% and target category Rottweiler.

Target Category	Model	Clean Data		Patched Data	
		Acc	FP	Acc	FP
Chicken Curry (10)	Clean	47.6	27	39.1	24
	Backdoored	48.2	19	9.2	3438
	PatchSearch	46.6	29	40.3	21
Steak (11)	Clean	47.6	17	39.5	24
	Backdoored	48.0	26	13.5	2596
	PatchSearch	47.4	22	40.7	57
Panna Cotta (12)	Clean	47.6	41	39.8	23
	Backdoored	47.5	42	19.1	2158
	PatchSearch	47.4	37	40.9	16
Deviled Eggs (13)	Clean	47.6	50	39.3	41
	Backdoored	47.9	23	7.9	4317
	PatchSearch	48.0	52	41.0	35
<i>Mean</i>	Clean	47.6	33.8	39.4	28.0
	Backdoored	47.9	27.5	12.4	3127.3
	PatchSearch	47.4	35.0	40.7	32.3

Table A7. **Attack on downstream task: Food101 classification.** Instead of attacking one of the categories in the pre-training dataset, we consider an attack on the downstream task of Food101 classification [6]. We poison 700 images of a category from Food101 and add them to the pre-training dataset (ImageNet-100). We evaluate the resulting models on Food101 classification and find that the targeted attack is successful. Next, we defend the backdoored models with PatchSearch and find that it is able to defend against the attack successfully. Averaged across the four categories, PatchSearch has 99.7% recall and 54.8% precision for filtering out poisons. Setting: BYOL & ResNet-18.

Model	Clean Data		Patched Data	
	Acc	FP	Acc	FP
with strong aug	63.1	21.0	53.5	63.1
without strong aug	65.7	18.7	53.8	97.6

Table A8. **Effect of strong augmentations during fine-tuning.** We find that strong augmentations like cutmix, mixup, and random erase can degrade the overall model performance. Setting: ViT-B, MAE and average of 4 target categories (Rottweiler, Tabby Cat, Ambulance, and Laptop).

Target				Clean Data		Patched Data	
Category	Model	PatchSearch	<i>i</i> -CutMix	Acc	FP	Acc	FP
Rottweiler (10)	clean	✗	✗	65.7	29.4	60.6	24.4
	clean	✓	✓	65.8	33.0	64.0	28.6
	backdoored	✗	✗	66.4	26.0	39.9	1926.8
	defended	✗	✓	66.8	26.4	60.4	272.0
	defended	✓	✗	66.9	31.8	61.1	43.8
	defended	✓	✓	65.6	32.6	63.4	32.2
Tabby Cat (11)	clean	✗	✗	65.7	5.0	60.8	9.4
	clean	✗	✓	65.8	4.4	63.8	3.4
	backdoored	✗	✗	65.9	1.6	31.9	2338.0
	defended	✗	✓	67.1	6.4	62.5	299.0
	defended	✓	✗	66.1	4.2	61.9	7.0
	defended	✓	✓	67.0	3.2	64.6	1.6
Ambulance (12)	clean	✗	✗	65.7	8.6	60.5	10.2
	clean	✗	✓	65.8	9.2	63.9	10.4
	backdoored	✗	✗	66.8	7.6	53.6	895.6
	defended	✗	✓	66.5	7.8	62.4	116.0
	defended	✓	✗	66.4	9.2	61.2	9.2
	defended	✓	✓	67.3	7.2	65.1	7.0
Pickup Truck (13)	clean	✗	✗	65.7	14.2	60.6	13.2
	clean	✗	✓	65.8	14.2	63.7	15.6
	backdoored	✗	✗	66.4	13.6	48.0	1430.6
	defended	✗	✓	68.3	16.8	64.1	131.0
	defended	✓	✗	65.7	14.8	60.8	13.0
	defended	✓	✓	67.1	19.0	64.6	18.8
Laptop (14)	clean	✗	✗	65.7	30.8	60.4	33.0
	clean	✗	✓	65.8	27.4	64.1	15.8
	backdoored	✗	✗	65.6	29.6	16.1	355.8
	defended	✗	✓	66.2	25.2	60.5	356.2
	defended	✓	✗	65.8	30.4	60.9	37.6
	defended	✓	✓	66.3	29.0	63.9	17.0
Goose (15)	clean	✗	✗	65.7	10.4	60.3	18.0
	clean	✗	✓	65.8	8.6	63.6	9.0
	backdoored	✗	✗	66.4	10.0	26.5	3391.0
	defended	✗	✓	67.4	11.2	62.4	288.2
	defended	✓	✗	66.4	12.2	61.6	23.6
	defended	✓	✓	67.0	25.6	65.1	24.8
Pirate Ship (16)	clean	✗	✗	65.7	4.2	60.4	4.6
	clean	✗	✓	65.8	5.4	63.5	6.0
	backdoored	✗	✗	66.9	4.6	38.3	2438.6
	defended	✗	✓	67.2	4.4	61.4	328.2
	defended	✓	✗	66.4	4.0	61.5	9.8
	defended	✓	✓	66.9	3.6	64.4	2.4
Gas Mask (17)	clean	✗	✗	65.7	19.8	60.7	28.6
	clean	✗	✓	65.8	22.0	63.8	36.4
	backdoored	✗	✗	66.6	21.2	34.8	2722.4
	defended	✗	✓	66.8	18.6	57.2	931.6
	defended	✓	✗	66.8	30.8	62.1	52.8
	defended	✓	✓	67.2	29.0	64.9	55.8
Vacuum Cleaner (18)	clean	✗	✗	65.7	60.0	60.6	74.8
	clean	✗	✓	65.8	58.6	63.4	38.2
	backdoored	✗	✗	66.6	49.2	20.4	3210.0
	defended	✗	✓	67.0	51.0	59.6	328.4
	defended	✓	✗	66.8	58.4	61.8	85.8
	defended	✓	✓	67.0	59.0	64.8	39.2
American Lobster (19)	clean	✗	✗	65.7	24.0	61.0	31.4
	clean	✗	✓	65.8	21.8	63.5	32.6
	backdoored	✗	✗	67.0	22.2	43.5	2086.8
	defended	✗	✓	66.5	17.6	59.5	602.2
	defended	✓	✗	66.3	32.8	61.9	51.0
	defended	✓	✓	66.6	15.8	64.2	26.6
Mean and STD	clean	✗	✗	65.7 ± 0.0	20.6 ± 16.8	60.6 ± 0.2	24.8 ± 20.2
	clean	✗	✓	65.8 ± 0.0	20.5 ± 16.5	63.7 ± 0.2	19.6 ± 13.1
	backdoored	✗	✗	66.5 ± 0.4	18.6 ± 14.3	35.3 ± 11.9	2079.6 ± 967.5
	defended	✗	✓	67.0 ± 0.6	18.5 ± 13.7	61.0 ± 2.0	365.3 ± 239.4
	defended	✓	✗	66.4 ± 0.4	22.9 ± 17.1	61.5 ± 0.5	33.4 ± 25.6
	defended	✓	✓	66.8 ± 0.5	22.4 ± 16.8	64.5 ± 0.5	22.5 ± 17.1

Table A9. Per category results for BYOL, ResNet-18, and poison rate 0.5%. Detailed results for Table 2.

Target				Clean Data		Patched Data	
Category	Model	PatchSearch	<i>i</i> -CutMix	Acc	FP	Acc	FP
Rottweiler (10)	clean	✗	✗	49.7	36.8	46.3	28.4
	clean	✗	✓	55.9	32.0	54.0	29.4
	backdoored	✗	✗	50.2	39.4	33.4	1094.4
	defended	✗	✓	55.1	38.6	52.4	117.2
	defended	✓	✗	50.0	40.2	46.3	26.6
	defended	✓	✓	55.6	39.0	54.4	35.2
Tabby Cat (11)	clean	✗	✗	49.7	7.2	46.4	8.2
	clean	✗	✓	55.9	8.0	54.4	6.6
	backdoored	✗	✗	50.0	5.8	33.5	1901.6
	defended	✗	✓	55.3	5.6	52.4	181.8
	defended	✓	✗	49.9	11.6	46.7	13.8
	defended	✓	✓	55.7	8.8	54.2	9.0
Ambulance (12)	clean	✗	✗	49.7	18.4	46.4	15.0
	clean	✗	✓	55.9	13.6	54.4	19.2
	backdoored	✗	✗	50.3	14.0	46.2	103.2
	defended	✗	✓	55.4	10.6	53.8	27.6
	defended	✓	✗	49.0	14.4	45.4	15.0
	defended	✓	✓	55.7	16.4	54.1	15.8
Pickup Truck (13)	clean	✗	✗	49.7	16.6	46.6	15.4
	clean	✗	✓	55.9	16.6	54.2	18.6
	backdoored	✗	✗	50.6	13.0	46.4	115.0
	defended	✗	✓	55.8	15.2	53.8	75.2
	defended	✓	✗	49.1	18.2	45.8	17.8
	defended	✓	✓	55.5	16.6	53.9	18.6
Laptop (14)	clean	✗	✗	49.7	37.0	46.0	35.8
	clean	✗	✓	55.9	33.0	54.3	24.2
	backdoored	✗	✗	49.8	33.8	41.8	466.2
	defended	✗	✓	55.4	24.0	53.7	92.6
	defended	✓	✗	49.7	41.6	45.8	47.6
	defended	✓	✓	54.8	46.4	53.5	29.4
Goose (15)	clean	✗	✗	49.7	37.2	46.6	39.4
	clean	✗	✓	55.9	38.2	54.2	39.2
	backdoored	✗	✗	49.6	33.8	45.2	194.0
	defended	✗	✓	55.5	31.0	53.8	46.8
	defended	✓	✗	49.5	39.6	46.3	46.0
	defended	✓	✓	55.1	41.0	53.4	37.6
Pirate Ship (16)	clean	✗	✗	49.7	8.2	46.1	9.0
	clean	✗	✓	55.9	5.0	54.4	6.4
	backdoored	✗	✗	49.7	6.2	42.1	573.4
	defended	✗	✓	55.7	4.6	53.6	68.8
	defended	✓	✗	49.7	12.8	46.8	17.4
	defended	✓	✓	55.0	5.6	53.3	8.2
Gas Mask (17)	clean	✗	✗	49.7	43.8	46.5	57.0
	clean	✗	✓	55.9	39.6	54.2	56.8
	backdoored	✗	✗	49.7	43.6	42.3	561.2
	defended	✗	✓	55.2	37.2	53.2	92.6
	defended	✓	✗	50.0	47.4	46.6	62.4
	defended	✓	✓	55.6	48.6	53.8	61.8
Vacuum Cleaner (18)	clean	✗	✗	49.7	51.4	46.5	63.2
	clean	✗	✓	55.9	53.6	54.3	40.0
	backdoored	✗	✗	49.8	48.8	41.7	424.4
	defended	✗	✓	55.4	54.2	53.7	55.6
	defended	✓	✗	49.9	58.0	46.4	64.6
	defended	✓	✓	54.8	51.2	53.4	37.0
American Lobster (19)	clean	✗	✗	49.7	34.0	46.5	35.4
	clean	✗	✓	55.9	23.4	54.1	30.4
	backdoored	✗	✗	49.8	34.0	36.1	1599.2
	defended	✗	✓	55.2	23.2	49.7	482.2
	defended	✓	✗	50.0	47.2	46.7	61.2
	defended	✓	✓	55.6	32.2	54.3	39.2
Mean and STD	clean	✗	✗	49.7 ± 0.0	29.1 ± 15.3	46.4 ± 0.2	30.7 ± 19.2
	clean	✗	✓	55.9 ± 0.0	26.3 ± 15.6	54.3 ± 0.1	27.1 ± 15.6
	backdoored	✗	✗	50.0 ± 0.3	27.2 ± 16.0	40.9 ± 4.9	703.3 ± 626.1
	defended	✗	✓	55.4 ± 0.2	24.4 ± 16.1	53.0 ± 1.3	124.0 ± 132.9
	defended	✓	✗	49.7 ± 0.4	33.1 ± 17.1	46.3 ± 0.5	37.2 ± 21.3
	defended	✓	✓	55.3 ± 0.4	30.6 ± 17.2	53.8 ± 0.4	29.2 ± 16.6

Table A10. Per category results for MoCo-v2, ResNet-18, and poison rate 0.5%. Detailed results for Table 2.

Target		Clean Data				Patched Data	
Category	Model	PatchSearch	<i>i</i> -CutMix	Acc	FP	Acc	FP
Rottweiler (10)	clean	✗	✗	70.5	25.4	65.1	16.6
	clean	✗	✓	75.6	31.0	74.5	26.8
	backdoored	✗	✗	70.5	24.8	42.9	1412.4
	defended	✗	✓	75.6	28.0	70.5	268.4
	defended	✓	✗	70.4	28.2	64.8	23.8
	defended	✓	✓	75.7	32.6	74.9	28.6
Tabby Cat (11)	clean	✗	✗	70.5	3.2	64.4	4.0
	clean	✗	✓	75.6	5.0	74.5	3.4
	backdoored	✗	✗	70.7	4.6	42.0	2354.6
	defended	✗	✓	75.4	5.2	73.3	133.4
	defended	✓	✗	70.6	6.0	65.2	7.6
	defended	✓	✓	74.7	4.8	73.8	4.4
Ambulance (12)	clean	✗	✗	70.5	9.8	64.4	13.4
	clean	✗	✓	75.6	8.2	74.4	8.0
	backdoored	✗	✗	70.5	10.2	56.9	748.8
	defended	✗	✓	75.3	8.6	74.4	49.6
	defended	✓	✗	70.1	9.6	63.8	16.6
	defended	✓	✓	75.2	6.8	74.1	8.6
Pickup Truck (13)	clean	✗	✗	70.5	13.8	65.3	10.4
	clean	✗	✓	75.6	11.0	74.3	13.0
	backdoored	✗	✗	70.7	11.8	57.9	805.0
	defended	✗	✓	75.7	12.4	74.3	54.2
	defended	✓	✗	69.9	14.0	65.9	12.8
	defended	✓	✓	75.1	13.0	74.3	13.6
Laptop (14)	clean	✗	✗	70.5	29.6	64.9	39.4
	clean	✗	✓	75.6	34.4	74.5	26.6
	backdoored	✗	✗	70.5	34.6	42.3	2172.4
	defended	✗	✓	75.8	31.4	71.7	365.8
	defended	✓	✗	69.8	43.0	62.9	93.2
	defended	✓	✓	75.3	40.8	74.3	31.0
Goose (15)	clean	✗	✗	70.5	6.4	64.8	9.6
	clean	✗	✓	75.6	6.6	74.4	6.8
	backdoored	✗	✗	70.7	6.8	48.6	1693.2
	defended	✗	✓	76.1	5.2	74.0	66.6
	defended	✓	✗	69.6	10.0	64.8	17.4
	defended	✓	✓	75.1	7.2	74.0	5.0
Pirate Ship (16)	clean	✗	✗	70.5	2.0	64.2	1.4
	clean	✗	✓	75.6	2.6	74.1	1.6
	backdoored	✗	✗	70.7	2.2	56.2	915.6
	defended	✗	✓	75.8	2.0	73.5	93.4
	defended	✓	✗	70.4	2.0	62.1	1.6
	defended	✓	✓	75.7	1.6	74.6	1.2
Gas Mask (17)	clean	✗	✗	70.5	29.0	64.5	55.6
	clean	✗	✓	75.6	12.4	74.3	23.0
	backdoored	✗	✗	70.3	20.0	39.2	2558.0
	defended	✗	✓	75.4	14.2	71.0	381.6
	defended	✓	✗	70.4	42.6	64.6	71.2
	defended	✓	✓	74.9	28.0	73.8	44.2
Vacuum Cleaner (18)	clean	✗	✗	70.5	52.0	64.4	113.8
	clean	✗	✓	75.6	36.2	74.5	25.4
	backdoored	✗	✗	70.6	49.8	43.8	1847.4
	defended	✗	✓	75.2	35.4	66.3	723.4
	defended	✓	✗	70.4	60.8	65.2	128.0
	defended	✓	✓	75.0	49.4	73.8	36.2
American Lobster (19)	clean	✗	✗	70.5	13.6	64.5	8.0
	clean	✗	✓	75.6	8.2	74.3	11.4
	backdoored	✗	✗	70.7	9.4	39.0	2581.8
	defended	✗	✓	75.9	6.4	73.3	185.4
	defended	✓	✗	70.8	14.4	65.7	25.8
	defended	✓	✓	75.1	12.4	74.2	16.8
Mean and STD	clean	✗	✗	70.5 ± 0.0	18.5 ± 15.6	64.6 ± 0.4	27.2 ± 34.8
	clean	✗	✓	75.6 ± 0.0	15.6 ± 13.0	74.4 ± 0.1	14.6 ± 10.0
	backdoored	✗	✗	70.6 ± 0.1	17.4 ± 15.1	46.9 ± 7.5	1708.9 ± 714.2
	defended	✗	✓	75.6 ± 0.3	14.9 ± 12.2	72.2 ± 2.5	232.2 ± 212.5
	defended	✓	✗	70.2 ± 0.4	23.1 ± 19.6	64.5 ± 1.2	39.8 ± 42.6
	defended	✓	✓	75.2 ± 0.3	19.7 ± 16.8	74.2 ± 0.4	19.0 ± 15.0

Table A11. **Per category results for MoCo-v3, ViT-B, and poison rate 0.5%.** Detailed results for Table 2.

Target	Clean Data				Patched Data		
Category	Model	PatchSearch	<i>i</i> -CutMix	Acc	FP	Acc	FP
Rottweiler (10)	clean	✗	✗	70.5	25.4	65.1	16.6
	clean	✗	✓	75.6	31.0	74.5	26.8
	backdoored	✗	✗	70.8	21.8	30.6	3127.4
	defended	✗	✓	75.6	21.4	69.0	510.8
	defended	✓	✗	70.0	32.6	64.8	26.8
	defended	✓	✓	75.3	37.4	74.3	32.0
Tabby Cat (11)	clean	✗	✗	70.5	3.2	64.4	4.0
	clean	✗	✓	75.6	5.0	74.5	3.4
	backdoored	✗	✗	70.9	3.4	23.8	3669.8
	defended	✗	✓	75.7	5.0	71.0	429.6
	defended	✓	✗	70.4	24.4	62.9	86.8
	defended	✓	✓	74.8	29.6	74.0	18.0
Ambulance (12)	clean	✗	✗	70.5	9.8	64.4	13.4
	clean	✗	✓	75.6	8.2	74.4	8.0
	backdoored	✗	✗	70.7	7.4	49.0	1511.2
	defended	✗	✓	75.4	7.4	73.0	180.8
	defended	✓	✗	69.8	25.8	64.2	36.0
	defended	✓	✓	75.0	21.0	73.9	21.0
Pickup Truck (13)	clean	✗	✗	70.5	13.8	65.3	10.4
	clean	✗	✓	75.6	11.0	74.3	13.0
	backdoored	✗	✗	70.7	12.4	54.6	1007.0
	defended	✗	✓	75.2	9.6	73.0	117.8
	defended	✓	✗	69.7	18.4	64.9	19.2
	defended	✓	✓	75.2	18.6	73.9	21.4
Laptop (14)	clean	✗	✗	70.5	29.6	64.9	39.4
	clean	✗	✓	75.6	34.4	74.5	26.6
	backdoored	✗	✗	70.8	27.0	39.4	2566.4
	defended	✗	✓	75.4	25.8	66.9	724.8
	defended	✓	✗	69.9	61.2	60.3	112.8
	defended	✓	✓	75.6	54.2	74.7	55.0
Goose (15)	clean	✗	✗	70.5	6.4	64.8	9.6
	clean	✗	✓	75.6	6.6	74.4	6.8
	backdoored	✗	✗	70.5	5.8	38.6	2437.8
	defended	✗	✓	75.7	5.2	71.2	314.2
	defended	✓	✗	70.1	32.8	63.4	59.0
	defended	✓	✓	74.7	33.4	73.7	25.0
Pirate Ship (16)	clean	✗	✗	70.5	2.0	64.2	1.4
	clean	✗	✓	75.6	2.6	74.1	1.6
	backdoored	✗	✗	70.3	2.6	44.9	2093.8
	defended	✗	✓	75.8	2.4	73.5	90.8
	defended	✓	✗	70.3	20.2	62.4	36.0
	defended	✓	✓	75.0	18.2	74.0	24.2
Gas Mask (17)	clean	✗	✗	70.5	29.0	64.5	55.6
	clean	✗	✓	75.6	12.4	74.3	23.0
	backdoored	✗	✗	70.4	21.2	31.6	3112.0
	defended	✗	✓	76.1	8.8	67.7	728.6
	defended	✓	✗	70.2	79.8	64.7	138.8
	defended	✓	✓	74.6	66.2	73.5	102.8
Vacuum Cleaner (18)	clean	✗	✗	70.5	52.0	64.4	113.8
	clean	✗	✓	75.6	36.2	74.5	25.4
	backdoored	✗	✗	71.2	38.4	40.6	2203.4
	defended	✗	✓	75.3	30.6	64.9	903.0
	defended	✓	✗	70.4	85.8	64.8	182.8
	defended	✓	✓	74.6	69.2	73.8	63.2
American Lobster (19)	clean	✗	✗	70.5	13.6	64.5	8.0
	clean	✗	✓	75.6	8.2	74.3	11.4
	backdoored	✗	✗	70.3	8.0	24.1	3630.2
	defended	✗	✓	75.1	9.0	71.6	341.6
	defended	✓	✗	70.0	53.4	64.3	61.6
	defended	✓	✓	75.1	42.2	74.3	51.0
Mean and STD	clean	✗	✗	70.5 ± 0.0	18.5 ± 15.6	64.6 ± 0.4	27.2 ± 34.8
	clean	✗	✓	75.6 ± 0.0	15.6 ± 13.0	74.4 ± 0.1	14.6 ± 10.0
	backdoored	✗	✗	70.7 ± 0.3	14.8 ± 11.8	37.7 ± 10.2	2535.9 ± 873.6
	defended	✗	✓	75.5 ± 0.3	12.5 ± 9.8	70.2 ± 2.9	434.2 ± 279.3
	defended	✓	✗	70.1 ± 0.2	43.4 ± 25.0	63.7 ± 1.5	76.0 ± 53.9
	defended	✓	✓	75.0 ± 0.3	39.0 ± 18.8	74.0 ± 0.3	41.4 ± 27.0

Table A12. **Per category results for MoCo-v3, ViT-B, and poison rate 1.0%.** Detailed results for Table 2.