

Explicit Tradeoffs between Adversarial and Natural Distributional Robustness

Mazda Moayeri
mmoayeri@umd.edu

Kiarash Banihashem
kiarash@umd.edu

Soheil Feizi
sfeizi@cs.umd.edu

Department of Computer Science
University of Maryland

Abstract

Several existing works study either adversarial or natural distributional robustness of deep neural networks separately. In practice, however, models need to enjoy both types of robustness to ensure reliability. In this work, we bridge this gap and show that in fact, *explicit tradeoffs* exist between adversarial and natural distributional robustness. We first consider a simple linear regression setting on Gaussian data with disjoint sets of *core* and *spurious* features. In this setting, through theoretical and empirical analysis, we show that (i) adversarial training with ℓ_1 and ℓ_2 norms increases the model reliance on spurious features; (ii) For ℓ_∞ adversarial training, spurious reliance only occurs when the scale of the spurious features is larger than that of the core features; (iii) adversarial training can have *an unintended consequence* in reducing distributional robustness, specifically when spurious correlations are changed in the new test domain. Next, we present extensive empirical evidence, using a test suite of twenty adversarially trained models evaluated on five benchmark datasets (ObjectNet, RIVAL10, Salient ImageNet-1M, ImageNet-9, Waterbirds), that adversarially trained classifiers rely on backgrounds more than their standardly trained counterparts, validating our theoretical results. We also show that spurious correlations in training data (when preserved in the test domain) can *improve* adversarial robustness, revealing that previous claims that adversarial vulnerability is rooted in spurious correlations are incomplete.

1 Introduction

Despite continuously improving upon state of the art accuracy on various benchmarks, deep image classifiers remain brittle to distribution shifts, suffering massive performance drops when evaluated on non-i.i.d. data. For example, the accuracy of object detectors trained on ImageNet [12] reduces by 40-45% on ObjectNet [5], where images are taken within households at various viewpoints and rotations. The reliance of deep models on *spurious features*, which correlate with class labels but are irrelevant to the true labeling function [29], is one cause of poor model robustness, as performance degrades when spurious correlations are broken. Indeed, model reliance on spurious features like texture [17] and background [68, 6] is well documented. Of greater concern, deep models in safety critical applications such as detection of pneumonia [71] and COVID-19 [11] have been observed to rely on hospital specific spurious markers, causing poor generalization to new hospitals.

Adversarial examples [62, 18] pose another troubling distribution shift, where imperceptible input perturbations can cause model accuracy to drop to zero. Many works have been proposed to improve the adversarial robustness of deep models [35, 73, 51, 58, 36, 10, 65, 54], including the widely popular *adversarial training*, where inputs are augmented via adversarial attack during training [39]. While spurious correlation robustness has also attracted lots of attention [2, 29, 38, 49], the

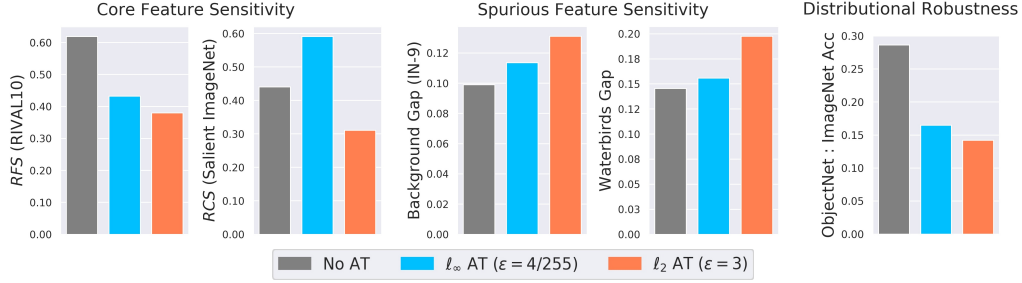


Figure 1: Snapshot of empirical evidence using *RIVAL10*, *Salient ImageNet-1M*, *ImageNet-9*, *Waterbirds*, and *ObjectNet* benchmarks. Results averaged over ResNet18 and ResNet50. **Adversarial training, especially under ℓ_2 norm, reduces (increases) sensitivity to core (spurious) features. The increased reliance on spurious features leads to worse distributional robustness.**

two problems are most often considered independently, despite both being essential to the safe and reliable deployment of deep models in the wild.

In the few works that do consider adversarial and spurious correlation robustness in tandem, the prevailing argument is that the origin of adversarial vulnerability is that the model focuses on spurious correlations that can be manipulated [76, 72, 25]. However, recently, noise-based analyses on RIVAL10 [41] and Salient ImageNet-1M [60] datasets suggest adversarial training may actually increase model sensitivity to spurious features; a result that is both counter-intuitive and in direct contrast to existing ideology.

To better understand this observation, we first appeal to a simple linear regression problem on Gaussian data with disentangled *core* and *spurious* features. In this setting, we theoretically show

- Adversarial training under ℓ_2 and ℓ_1 norms increase model reliance on spurious features, as using spurious features forces an attacker to spread its budget over additional features.
- For ℓ_∞ adversarial training, increased spurious feature reliance only occurs when *the scale* of the spurious feature is larger than that of the core features. That is, spurious features are used when perturbations that corrupt core features are too small to disrupt spurious correlations.
- Due to increased spurious feature reliance, there is an *explicit tradeoff* between adversarial and distributional robustness. Specifically, we show that adversarial training decreases model robustness to distribution shifts in the test domain where spurious correlations are broken.

To validate our theory, we evaluate twenty models adversarially trained using ℓ_2 and ℓ_∞ projected gradient descent [50]. Specifically, we inspect performance on multiple spurious robustness benchmarks over synthetic (ImageNet-9 [68], Waterbirds [49]) and real (RIVAL10 [41], Salient ImageNet-1M [60], ObjectNet [5]) datasets. Figure 1 summarizes our experiments, where we find that adversarially trained models consistently show greater sensitivity to spurious features compared to standardly trained baselines, with the effect more dramatic for ℓ_2 adversarial training than ℓ_∞ ¹. Finally, we show that the presence of spurious correlations in training data (when preserved in test domain) can improve adversarial robustness, with stronger spurious correlations leading to greater accuracy under attack.

Our work combines two prevalent but often separately considered notions of robustness, yielding surprising theoretically-derived and empirically-supported results. We hope our contributions grant insight to both adversarial and distributional robustness communities, and emphasize the need for holistic evaluations of model robustness.

2 Review of Literature

Adversarial Robustness. Since adversarial examples were first observed in deep models [62, 18], the phenomenon has been extensively studied. New attacks [8, 34, 42, 15] and defenses [35, 43, 40, 37]

¹High *RCS* for ℓ_∞ AT models is due to reduced *scale* of contextual bias in Salient ImageNet since the data diversity weakens background correlations. In ℓ_1 and ℓ_2 adversarial training, models rely on spurious features regardless of their scale. See details in Section 3.

are introduced frequently, in a game of cat and mouse where the attacker generally has the upper hand [3]. Certified defenses seek to break this cycle by offering provable robustness guarantees [58, 36, 51, 10]. Arguably the most popular defense is *adversarial training* [39] where images are augmented with adversarial perturbations during training, amounting to a min max optimization.

Natural Distributional Robustness. In contrast to synthetic adversarial perturbations, many works seek to characterize the robustness of deep models to naturally occurring distribution shifts, for instance due to common corruptions (noise, blur, etc) [22] or changes in rendition [21, 63]. [30] compiles ten benchmarks of realistic distribution shifts over diverse applications (medical, economic, etc). Many algorithms have been proposed to improve out-of-distribution robustness [38, 49, 2, 33], though in comprehensive evaluations, their gains over empirical risk minimization are marginal, as they often only hold for *certain* distribution shifts [70, 21]. [70] identifies *diversity* and *correlation* shifts as two key dimensions to OOD robustness; our work focuses on the latter.

Spurious Correlations. Solely optimizing for accuracy leads deep models to rely on *any* patterns predictive of class in the training domain. This includes *spurious* features, which are irrelevant to the true labeling function. Natural image datasets are riddled with spurious features [31, 59]. Spurious feature reliance becomes problematic under distribution shifts that break their correlation to class labels: sidewalk segmentation struggles in the absence of cars [56], familiar objects cannot be recognized in unfamiliar poses [1] or uncommon settings [27, 5], etc. A natural and ubiquitous spurious correlation in vision is image backgrounds, observed in numerous prior works to be leveraged by models for classification [68, 41] and object detection [47]. Spurious correlations also relate to algorithmic biases [13, 17], with implications for fairness [19, 7, 9, 28], reflecting the importance of this issue.

Accordingly, many works seek to improve spurious correlation robustness. Families of approaches include optimizing for worst group accuracy [49, 23, 48, 74, 38], learning invariant latent spaces [46, 2], appealing to meta-learning [45] or causality [44, 4]. Our work does not focus on mitigation methods, but instead sheds insight on how optimizing for a different notion of robustness (i.e. adversarial) affects spurious feature reliance, and consequently, natural distributional robustness. Generally, models trained under ERM are believed to have a propensity to use spurious features, especially when they are easy to learn, due to bias of learners (algorithmic and human) to absorb simple features first [55] and take shortcuts [16]. However, recent work suggests that core features are still learned under ERM even when spurious features are favored, and simple finetuning on data without the spurious feature can efficiently reduce spurious feature reliance without full model retraining [29].

Unintended Outcomes of Adversarial Training. Adversarial training achieves improved accuracy under attack, but comes at the cost of standard accuracy, with multiple works provably demonstrating this tradeoff [73, 14]. Notably, [26] inspires our theoretical analysis, though we focus on the effect of adversarial training on out-of-distribution (OOD) robustness to spurious correlation shifts, rather than its effect on standard accuracy. A more positive outcome is that adversarial training leads to perceptually aligned gradients [53], with applications to model debugging [61, 66], and further, transfer learning on adversarially robust features yields better accuracy on downstream tasks compared to features learned from standard training [50], despite having lower accuracy on the original task.

To our knowledge, robustness to adversarial and natural distribution shifts have not been studied in tandem. However, spurious correlations are at times mentioned with adversarial robustness, usually in claims that the origin of adversarial vulnerability is in model’s focus on (imperceptible) spurious features [76, 72, 25]. Our results create tension with the contrapositive of their argument, as we show that mitigating adversarial vulnerability (via adversarial training) results in *increased* spurious feature reliance. We do this analytically in a simple linear regression setting (Section 3), and empirically on multiple benchmarks, with an emphasis on natural spurious features (i.e., backgrounds) in our experiments (Section 3). Further, we even demonstrate a case where the presence of a spurious feature leads to *improved* adversarial robustness (Section 4.3)). We note that the spurious features we observe to be positively associated with adversarial robustness may be distinct from those that prior works claim contribute to adversarial vulnerability. However, our result of adversarial training leading to increased spurious feature reliance (of any kind) is novel and contrary to common understanding. Given the critical nature of adversarial and spurious correlation robustness for model security, reliability, and fairness, the significance of our result in revealing potential misconceptions on the interplay of these two crucial modes of robustness should not be understated.

[41] and [60] recently observed decreased core sensitivity on a handful of ℓ_2 adversarially trained models, in spirit with our findings, but with no explanation. We offer the first rigorous analysis

of this counterintuitive phenomenon, evaluating 5 to 10 times as many models in 5 times as many settings. More importantly, we theoretically prove that adversarial training increases spurious feature reliance, contributing novel fundamental insight as to how optimizing for adversarial robustness can lead to reduced robustness to natural distribution shifts, uncovering important effects like the norm of adversarial training and the scale of spurious features at play.

3 Theoretical Analysis on Linear models

We begin by analysing the effects of adversarial training on a simple linear regression model. Consider the model $Y = \langle X, \theta^{\text{opt}} \rangle + W$ where $X \in \mathbb{R}^m$ is the input variable, $\theta^{\text{opt}} \in \mathbb{R}^m$ is the optimal parameter, $\langle \cdot, \cdot \rangle$ represents the inner product and $W \in \mathbb{R}$ is a noise variable. We assume that the input variables follow a multivariate Gaussian distribution $N(0, \Sigma)$ where $\Sigma \in \mathbb{R}^{m \times m}$ is the covariance matrix and further assume that W is sampled from the Gaussian distribution $N(0, \sigma_w^2)$. We assume that the set of features $[m]$ consists of two groups, the *core features* C and the *spurious features* S . Without loss of generality, we assume that $C = \{1, \dots, p\}$ and $S = \{p+1, \dots, m\}$. We assume that the optimal parameter θ^{opt} has non-zero entries on the set C only. This implies that the output depends on the input only through the core features and conditioned on the core features, it is independent of the spurious ones. More formally, we assume that $Y \perp X_S | X_C$ where X_S and X_C represent the core and spurious subsets of the input, respectively.

For loss functions, we define the *standard loss* function as $L(\theta) = \mathbb{E}[(y - \langle X, \theta \rangle)^2]$, and the *adversarial loss* function as

$$L_{p,\epsilon}(\theta) = \mathbb{E} \left[\max_{\|\delta\|_p \leq \epsilon} (Y - \langle X + \delta, \theta \rangle)^2 \right], \quad (1)$$

where p is the attack norm and ϵ represents the norm budget.

We first show an equivalent form of (1) that is more amenable to analysis.

Theorem 1. Assume that $Y = \langle X, \theta^{\text{opt}} \rangle + W$ where $W \sim N(0, \sigma_w^2)$ is independent of X and $\theta^{\text{opt}} \in \mathbb{R}^m$ is a fixed parameter. Assume further that X follows the distribution $N(0, \Sigma)$ and define σ_θ^2 as $(\theta - \theta^{\text{opt}})^T \Sigma (\theta - \theta^{\text{opt}}) + \sigma_w^2$. The loss function (1) is equivalent to

$$L_{p,\epsilon}(\theta) = c_2 \cdot \sigma_\theta^2 + (c_1 \sigma_\theta + \epsilon \cdot \|\theta\|_q)^2 \quad (2)$$

where $c_1 = \sqrt{\frac{2}{\pi}} < 1$, $c_2 = 1 - c_1^2$ and $\|\cdot\|_q$ is the dual norm of $\|\cdot\|_p$, i.e., $\frac{1}{p} + \frac{1}{q} = 1$. Furthermore, the above formulation is convex in θ .

The above result is similar to Proposition 3.2 in [26] which provides characterization results for the ℓ_2 norm. The key distinction of our results is providing a simple convex formulation of the robust minimization problem, allowing the results to be easily generalized for an arbitrary ℓ_p norm. The theorem shows that the optimal value $\hat{\theta}$ minimizing the adversarial loss $L_{p,\epsilon}(\theta)$ is not θ^{opt} and in general, may be non-zero on the set of spurious features S . This means that adversarial training directs the model towards using the spurious correlations in order to increase robustness.

The proof of the theorem is provided in the Appendix. The main structure of the proof is similar to that of [26]. We first show that the inner maximization problem of (1) can be rewritten as

$$\max_{\|\delta\| \leq \epsilon} (Y - \langle X + \delta, \theta \rangle)^2 = \left(|Y - \langle X, \theta \rangle| + \epsilon \cdot \|\theta\|_q \right)^2. \quad (3)$$

This allows us to rewrite (1) as

$$\begin{aligned} L_{p,\epsilon} &= \mathbb{E} \left[\left(|Y - \langle X, \theta \rangle| + \epsilon \cdot \|\theta\|_q \right)^2 \right] \\ &= \mathbb{E} \left[(Y - \langle X, \theta \rangle)^2 \right] + \epsilon^2 \cdot \|\theta\|_q^2 + 2 \cdot \epsilon \cdot \|\theta\|_q \cdot \mathbb{E} [|Y - \langle X, \theta \rangle|] \\ &\stackrel{(a)}{=} \mathbb{E} \left[(\langle X, \theta - \theta^{\text{opt}} \rangle + W)^2 \right] + \epsilon^2 \cdot \|\theta\|_q^2 + 2 \cdot \epsilon \cdot \|\theta\|_q \cdot \mathbb{E} [\langle X, \theta - \theta^{\text{opt}} \rangle + W], \end{aligned}$$

where for (a) we have used the fact that $Y = \langle X, \theta \rangle + W$. Using the fact that X and W are Gaussian, we show that this is equal to $\sigma_\theta^2 + \epsilon^2 \cdot \|\theta\|_q^2 + 2 \cdot c_1 \cdot \epsilon \cdot \|\theta\|_q \cdot \sigma_\theta$ which we can further simplify to

(2) with some algebraic manipulation. Next, using standard techniques for analysing composition of convex functions, we show that σ_θ^2 and $(c_1\sigma_\theta + \epsilon \cdot \|\theta\|_q)^2$ are both convex in θ , implying that $L_{p,\epsilon}(\theta)$ is convex in θ as well.

Using the convex formulation in Theorem 1, we evaluate the linear model for different values of parameters to better understand the reliance of the model on spurious features. In our experiments, we consider a simple model with 5 features, the first two of which are core features.

We let η be a parameter controlling the correlation degree between the core and spurious features, with larger values corresponding to higher correlation. For the distribution of the data, define the matrix $\tilde{Q}$ and Q matrices as

$$\tilde{Q} = \begin{bmatrix} 1 & \frac{1}{2} & 0 & 0 & 0 \\ \frac{1}{2} & 1 & 0 & 0 & 0 \\ \eta & \eta & 1 & 0 & 0 \\ \eta & \eta & 0 & 1 & 0 \\ \eta & \eta & 0 & 0 & 1 \end{bmatrix}, \quad Q_{i,j} = \frac{\tilde{Q}_{i,j}}{\sqrt{\sum_{i,j'} \tilde{Q}_{i,j'}^2}} \quad (4)$$

Note that Q is obtained by normalizing the rows of $\tilde{Q}$. Each row of the Q matrix corresponds to an input feature. We let Σ take the value QQ^T . This is equivalent to sampling X from the distribution $Q\mathcal{N}(0, I)$. Throughout our experiments, we set $\sigma_w = 0.1$.

For an arbitrary vector θ , we define its Norm Fraction over Spurious features NFS(θ) as

$$\text{NFS}(\theta) = \frac{\sum_{i \in S} \theta_i^2}{\sum_j \theta_j^2} \quad (5)$$

Intuitively, NFS measures the degree to which a model relies on spurious features. For our first experiment, we perform adversarial training with varying parameter ϵ to obtain a predictor $\hat{\theta}$ and evaluate its NFS. We consider different choices of the ℓ_p norm as well as the correlation parameter η . The results can be seen in Figure 2.

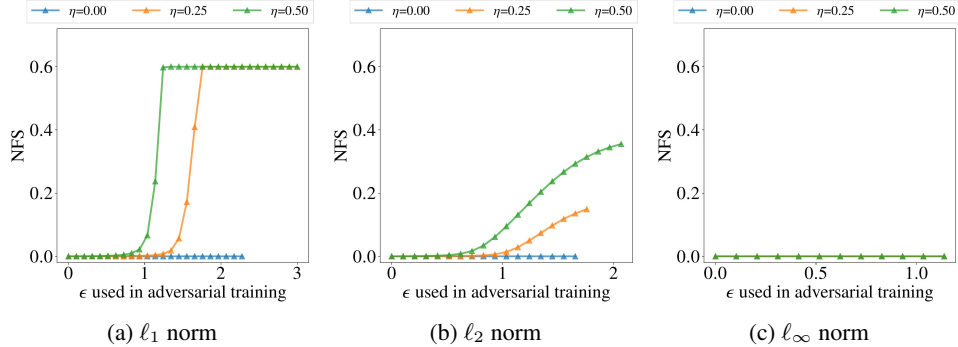


Figure 2: Reliance of the adversarially trained model on spurious features as measured by NFS (see Equation (5)) for different choices of ℓ_p norm and different values of adversarial budget ϵ and spurious correlation parameter η .

As seen in Figure 2, for the ℓ_1 and ℓ_2 norms, increasing the adversary's budget causes the model to rely more on the spurious features. To understand why this happens, it is helpful to consider a game theoretic perspective: If the model only looks at the core features, then the adversary will only need to perturb these features. Thus, even though the spurious features are normally less suitable for prediction, they now have the advantage of being less perturbed. Assuming ϵ is large enough, the model will be better off looking at these features in forming its prediction. Of course, if the model only looks at the spurious features, core features will become even more informative as they would be unperturbed as well. In the game's equilibrium, the model would use both the core and spurious features, relying more on the spurious features with increased values of ϵ .

Interestingly, this reasoning does not always apply for the ℓ_∞ norm. Indeed, if the model were to only look at the core features, the adversary may still perturb the spurious features with no extra cost. This is because the ℓ_∞ norm only measures the *maximum* perturbation in each feature and as long as the perturbations on the spurious features are not larger than the perturbations on the core features,

the norm would not change. The results of Figure 2c support this as the adversary does not rely on the spurious features for its predictions.

Importantly, however, in some cases the adversary may require a larger budget to perturb the spurious features compared to core features. Indeed, if we were to scale a feature by multiplication with a large number, then the adversary would require a larger budget to perturb that feature as the budget for that feature has effectively decreased. We can therefore expect that for the ℓ_∞ norm, scaling up the spurious features would cause the model to rely on them for prediction. Figure 3 shows that this is indeed the case. The figure shows the norm fraction over spurious (NFS) after scaling the spurious features for different values of ϵ . The scaling is done by multiplying the rows of Q corresponding to the spurious features by a scaling parameter. As seen in Figure 3, larger values of the scaling parameter, as well as larger values of ϵ , cause the model to become more reliant on the spurious features.

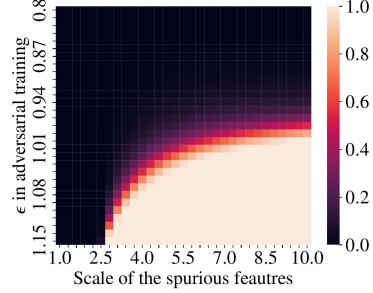


Figure 3: Norm fraction over spurious features (NFS) with an ℓ_∞ constrained adversary as a function of the scale of the spurious features and perturbation budget ϵ .

Next, we measure the effect of using spurious features on the model’s distributional robustness. To do this, we train two adversarial models. The first model, which we denote by “core”, uses only the core features while the second model, denoted by “total”, uses all of the features. We then simulate a distribution shift that breaks the spurious correlations by adding random Gaussian noise to the spurious rows of matrix Q defined in (4). Specifically, for each entry of Q in a spurious row, we add a noise sampled from $N(0, \sigma_Q^2)$ where σ_Q is a noise parameter. We use a fixed value of $\eta = 0.25$ and vary the parameter ϵ , the norm p used in adversarial training as well as the variance of the Gaussian noise. The results are shown in Figure 4.

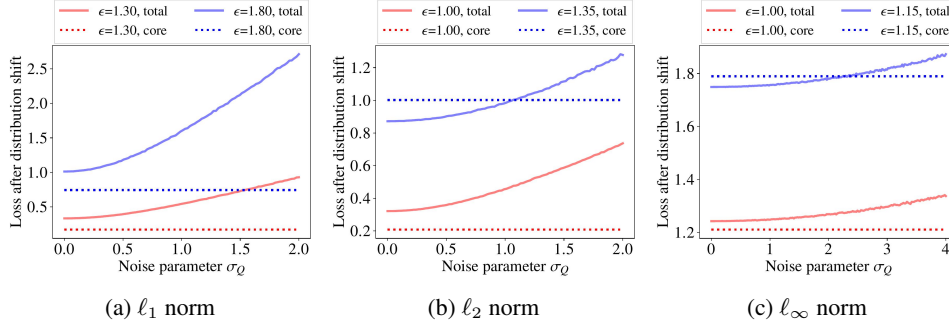


Figure 4: Effect of reliance on spurious features on distributional robustness. Each figures compares two models, one using only the core features, and another using all of the features (denoted by “core” and “total” respectively). For the ℓ_1 and ℓ_2 norms, the scale of the spurious features is 1 while for the ℓ_∞ norm, the scale is set to 3.

We observe that the clean (i.e. in distribution; $\sigma_Q = 0$) loss of the “core” model may be higher or lower than that of the “total” model, but in both cases, the total model is consistently more vulnerable to distributional shifts resulted from breaking spurious correlations.

4 Empirical Evidence

We now demonstrate increased spurious feature reliance in adversarially trained models over multiple benchmarks. We evaluate models on two backbones (ResNet18, ResNet50) adversarially trained on ImageNet [12] using two norms (ℓ_2, ℓ_∞) under five attack budgets (denoted ϵ) per norm, resulting in a $2 \times 2 \times 5 = 20$ model test suite, as well as standardly trained baselines. See appendix for details.

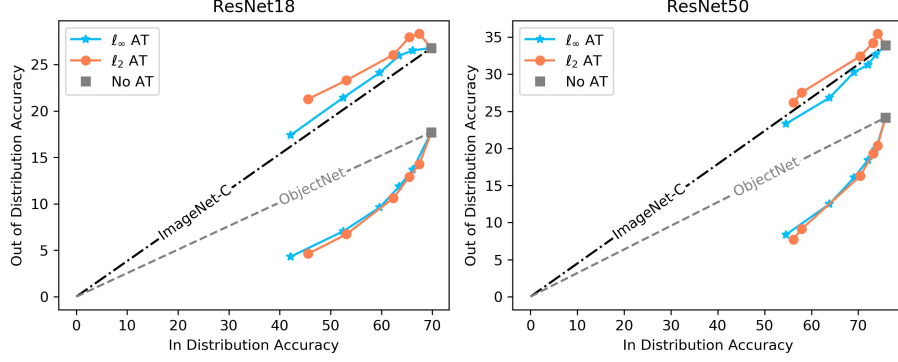


Figure 5: OOD accuracy vs standard ImageNet accuracy for adversarially trained ResNets. ImageNet-C accuracy is closely tied to standard accuracy, but for ObjectNet, where spurious correlations are broken, performance drop is more severe than a linear relation with standard accuracy would entail.

4.1 AT hurts Natural Distributional Robustness *only* when Spurious Correlations are broken

We first show reduced distributional robustness of adversarial models occurs specifically in cases where natural spurious correlations are broken. We appeal to the **ImageNet-C** [22] and **ObjectNet** [5] OOD benchmarks. ImageNet-C augments ImageNet samples with common corruptions like noise or blurring, distorting both core and spurious features equally. Crucially, these corruptions do not break spurious correlations. On the other hand, ObjectNet is formed by having workers capture images of common household objects (including samples from 113 classes of ImageNet) *in their homes*. Thus, only spurious features are affected. Namely, ObjectNet introduces distribution shifts in background, rotation, and viewpoint. We plot accuracies on these benchmarks in figure 5.

Recall that adversarially trained models have lower standard accuracy, which can confound our analysis, so we compare the drop in OOD accuracy to the drop in ImageNet accuracy across our model suite. Observe that the ratio of ImageNet-C accuracy to ImageNet accuracy is roughly constant across models. However, the ratio of ObjectNet accuracy to ImageNet accuracy is *lower* for adversarially trained models. Therefore, even after controlling for reduced standard accuracy, the distributional robustness of adversarially trained models is worse than that of standard models. Importantly, this effect does not hold for distribution shifts that maintain spurious correlations, indicating that the reduced distributional robustness is due to increased spurious feature reliance.

4.2 Reduced Core Sensitivity, and Difference in the Effect of ℓ_2 & ℓ_∞ Adversarial Training

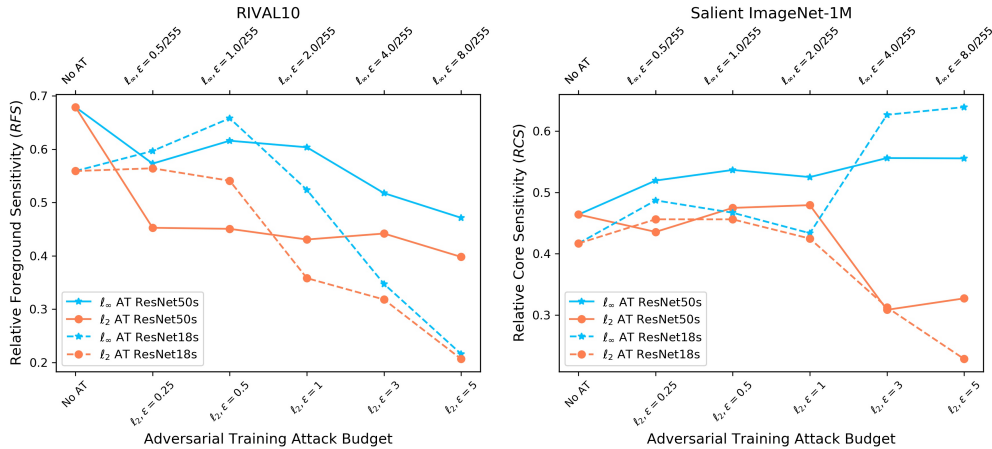


Figure 6: Noise-based evaluation of model sensitivity to foreground (RFS on RIVAL10) or core (RCS on Salient ImageNet-1M) regions. Lower values entail greater sensitivity to spurious regions.

We now directly quantify sensitivity to core features via **RIVAL10** and **Salient ImageNet-1M** datasets [41, 60]. The premise of this analysis is that model sensitivity to an input region can be quantified by the drop in accuracy due to corrupting that region [59]. [41] introduced the noise-based metric *relative foreground sensitivity* (RFS), which is the gap between accuracy drops due to background and foreground noise, normalized so to allow for comparisons across models with varying general noise robustness. RIVAL10 object segmentations allow for RFS computation. Analogously, *relative core sensitivity* (RCS) is computed using Salient ImageNet-1M’s soft segmentations of core input regions. A key distinction between the two metrics is that RCS is computed directly on pretrained models performing the original 1000-way ImageNet classification task, while RFS first requires models to be finetuned on the **much coarser** 10-way classification task of RIVAL10. Also, Salient ImageNet-1M includes *all ImageNet images*, while RIVAL10 only consists of 20 ImageNet classes.

Figure 6 shows a decrease in RFS and RCS as the attack budget ϵ seen during ℓ_2 adversarial training rises. Thus, **adversarial training reduces core feature sensitivity relative to spurious feature sensitivity**. Notably, this effect does *not* hold for models adversarially trained with attacks under the ℓ_∞ norm for RCS , though it does for RFS . Alluding to our theoretical result, we conjecture that in Salient ImageNet-1M, the scales of the spurious features are much smaller than in RIVAL10, due to the diversity of images and finer grain of classes. That is, a smaller perturbation is needed to alter a spurious feature so that it correlates with an incorrect class when there are 1000 classes than when there are only 10 classes with generally disparate backgrounds.

4.3 Adversarial Training Increases Background Reliance in Synthetic Datasets

Now, we take a closer look at the reliance of adversarially trained models on the contextual spurious feature of *backgrounds* via the synthetic datasets **ImageNet-9** [68] and **Waterbirds** [49]. Both datasets use segmentations to superimpose objects over varying backgrounds, detailed below.

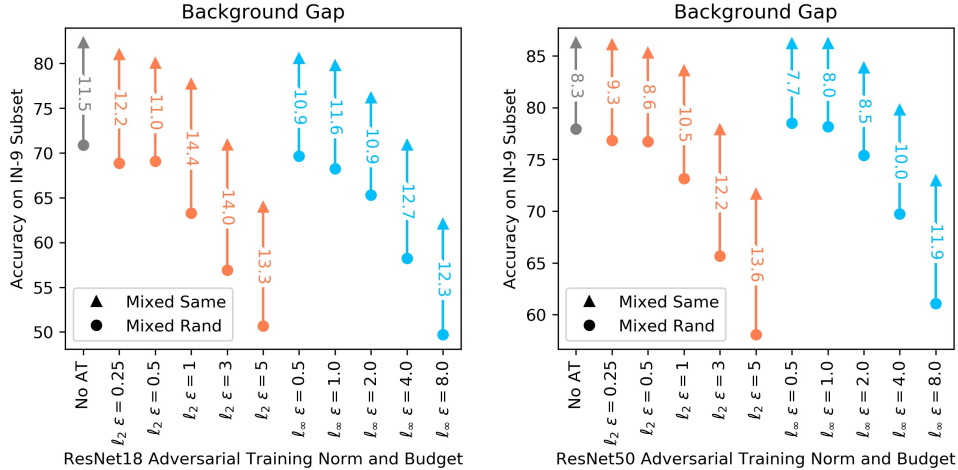


Figure 7: Background Gap (difference in accuracies on ImageNet-9 subsets MIXED-SAME and MIXED-RAND). The drop in accuracy due to background cross-class swapping (MIXED-RAND) causes larger drops in accuracy for robust models, especially ℓ_2 adversarially trained models.

ImageNet-9 (IN-9) organizes a subset of ImageNet into nine superclasses. Multiple validation sets exist for IN-9, where backgrounds or foregrounds are altered; we use MIXED-SAME and MIXED-RAND. In both sets, original backgrounds are swapped out for new ones. Crucially, in MIXED-SAME, the new backgrounds are taken from other instances *within the same class*, while MIXED-RAND uses *random backgrounds*. The metric, **Background Gap**, is the difference in model accuracy on MIXED-SAME and MIXED-RAND (i.e. drop due to breaking spurious background correlation).

Figure 7 shows that the background gap for our test suite of twenty robust ResNets and two standardly trained baselines. Nine out of the ten ℓ_2 adversarially trained models have larger background gaps than the standard baselines, while the same is true for five out of the ten ℓ_∞ adversarially trained models. When considering relative gaps (i.e. as a percent of the accuracy on MIXED-SAME), the

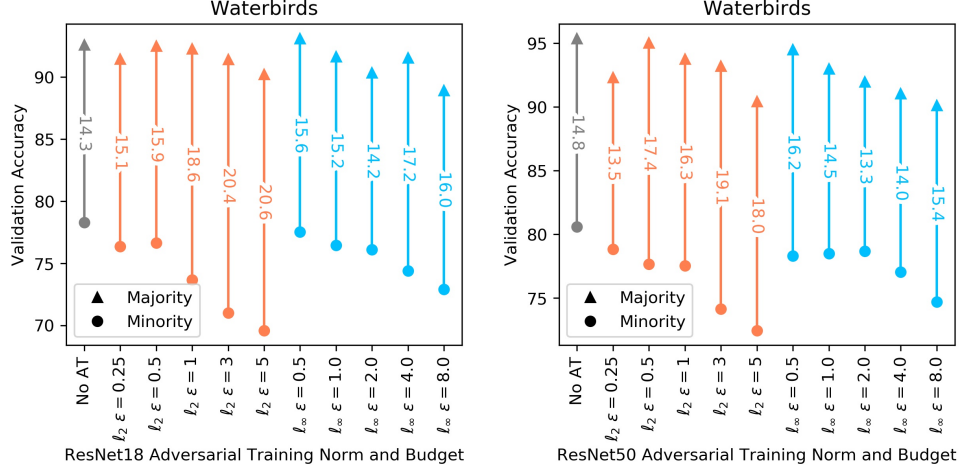


Figure 8: Accuracies on Waterbirds subsets where spurious correlation is intact (majority group; e.g. land birds on land backgrounds) and where it is broken (minority group; e.g. land birds on water backgrounds). The drop in performance due to breaking the background correlation grows larger for robust models, especially ℓ_2 adversarially trained models.

increase in gap becomes even more dramatic, with the ℓ_2 adversarially trained ResNet50 for $\epsilon = 5$ having a 18.9% relative drop, compared to the 9.7% relative drop for the standardly trained ResNet50.

Waterbirds combines foregrounds from Caltech Birds [64] and backgrounds from SUN Places [67]. The task is binary classification of landbirds and waterbirds. The majority group (95% of training samples) consists of landbirds over land backgrounds and waterbirds over water backgrounds. The minority group breaks this spurious correlation, placing landbirds over water backgrounds, and vice versa. The test set is evenly split between these groups. We train only a final linear layer atop the frozen feature extractors (so that models remain adversarially robust) for each of our models on the Waterbirds training set for ten epochs, saving the model with highest validation accuracy.

Figure 8 shows majority and minority group accuracies, and the gap between them, for our test suite of models. Again, we see increased gaps for robust models, with 100% of ℓ_2 and 60% of ℓ_∞ models respectively having larger gaps than the standardly trained baseline on the corresponding backbone.

In both benchmarks, breaking the background spurious correlation causes a more significant drop in performance for adversarially trained models than standardly trained models, indicating that adversarial training led to increased reliance on backgrounds. The observed affect is stronger for ℓ_2 adversarially trained models than ℓ_∞ ones. Further, the gaps grow near monotonically with ϵ .

4.4 Reverse Effect: Presence of Spurious Correlations Can Improve Adversarial Robustness

Finally, we show evidence that is directly at odds with the claim that spurious features lead to adversarial vulnerability. We train ResNet18s on CIFAR10 [32] with a spurious feature injected. Namely, images have all values in one color channel slightly increased. The majority group consists of red-shifted images from classes 0–4 and green shifted images from classes 5–9, while the minority group has reverse color-shifts. The parameter ρ is the ratio between majority and minority group size, controlling the strength of the spurious correlation (higher ρ means stronger spurious correlation; $\rho = 1 : 1$ means the spurious feature has no predictive power). We then evaluate the accuracy of the trained model under adversarial attack on an i.i.d. test set (**spurious feature retained**, no distribution shift).

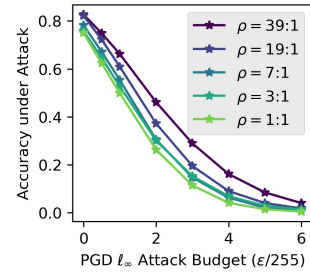


Figure 9: Injecting spurious feature improves adversarial robustness.

Figure 9 visualizes the results. Not surprisingly, the clean accuracy is higher for models trained on data with higher ρ , as the spurious feature is more predictive for

higher ρ . In fact, the added predictive influence the spurious feature leads to better accuracy *under attack*, with the gap between highest and lowest ρ values growing up to four fold compared to the baseline gap in clean accuracy. Thus, using a spurious feature can improve adversarial robustness. Despite being on a contrived example, this experiment shows that, while some spurious correlations may cause adversarial vulnerability, *others do the opposite*: the picture is more nuanced than previously assumed.

5 Acknowledgements

This project was supported in part by NSF CAREER AWARD 1942230, a grant from NIST 60NANB20D134, HR001119S0026 (GARD), ONR YIP award N00014-22-1-2271, Army Grant No. W911NF2120076 and the NSF award CCF2212458.

References

- [1] M. A. Alcorn, Q. Li, Z. Gong, C. Wang, L. Mai, W.-S. Ku, and A. M. Nguyen. Strike (with) a pose: Neural networks are easily fooled by strange poses of familiar objects. *2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 4840–4849, 2019.
- [2] M. Arjovsky, L. Bottou, I. Gulrajani, and D. Lopez-Paz. Invariant risk minimization, 2020.
- [3] A. Athalye, N. Carlini, and D. A. Wagner. Obfuscated gradients give a false sense of security: Circumventing defenses to adversarial examples. In *ICML*, 2018.
- [4] B. Aubin, A. Slowik, M. Arjovsky, L. Bottou, and D. Lopez-Paz. Linear unit-tests for invariance discovery. *ArXiv*, abs/2102.10867, 2021.
- [5] A. Barbu, D. Mayo, J. Alverio, W. Luo, C. Wang, D. Gutfreund, J. B. Tenenbaum, and B. Katz. Objectnet: A large-scale bias-controlled dataset for pushing the limits of object recognition models. In *NeurIPS*, 2019.
- [6] S. Beery, G. V. Horn, and P. Perona. Recognition in terra incognita. *CoRR*, abs/1807.04975, 2018.
- [7] J. Buolamwini and T. Gebru. Gender shades: Intersectional accuracy disparities in commercial gender classification. In *FAT*, 2018.
- [8] N. Carlini and D. A. Wagner. Towards evaluating the robustness of neural networks. *2017 IEEE Symposium on Security and Privacy (SP)*, pages 39–57, 2017.
- [9] A. Chouldechova. Fair prediction with disparate impact: A study of bias in recidivism prediction instruments. *Big data*, 5 2:153–163, 2017.
- [10] J. M. Cohen, E. Rosenfeld, and J. Z. Kolter. Certified adversarial robustness via randomized smoothing. In *ICML*, 2019.
- [11] A. J. DeGrave, J. D. Janizek, and S.-I. Lee. Ai for radiographic covid-19 detection selects shortcuts over signal. *Nature Machine Intelligence*, 3(7):610–619, 2021.
- [12] J. Deng, W. Dong, R. Socher, L. Li, K. Li, and L. Fei-Fei. Imagenet: A large-scale hierarchical image database. *2009 IEEE Conference on Computer Vision and Pattern Recognition*, pages 248–255, 2009.
- [13] J. Djolonga, J. Yung, M. Tschannen, R. Romijnders, L. Beyer, A. Kolesnikov, J. Puigcerver, M. Minderer, A. D’Amour, D. I. Moldovan, S. Gelly, N. Houlsby, X. Zhai, and M. Lucic. On robustness and transferability of convolutional neural networks. *2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 16453–16463, 2021.
- [14] E. Dobriban, H. Hassani, D. Hong, and A. Robey. Provable tradeoffs in adversarially robust classification. *ArXiv*, abs/2006.05161, 2020.
- [15] L. Engstrom, B. Tran, D. Tsipras, L. Schmidt, and A. Madry. Exploring the landscape of spatial robustness. In *ICML*, 2019.
- [16] R. Geirhos, J. Jacobsen, C. Michaelis, R. S. Zemel, W. Brendel, M. Bethge, and F. A. Wichmann. Shortcut learning in deep neural networks. *CoRR*, abs/2004.07780, 2020.
- [17] R. Geirhos, P. Rubisch, C. Michaelis, M. Bethge, F. A. Wichmann, and W. Brendel. Imagenet-trained cnns are biased towards texture; increasing shape bias improves accuracy and robustness, 2018.

- [18] I. Goodfellow, J. Shlens, and C. Szegedy. Explaining and harnessing adversarial examples. In *International Conference on Learning Representations*, 2015.
- [19] T. B. Hashimoto, M. Srivastava, H. Namkoong, and P. Liang. Fairness without demographics in repeated loss minimization. In *ICML*, 2018.
- [20] K. He, X. Zhang, S. Ren, and J. Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2016.
- [21] D. Hendrycks, S. Basart, N. Mu, S. Kadavath, F. Wang, E. Dorundo, R. Desai, T. L. Zhu, S. Parajuli, M. Guo, D. X. Song, J. Steinhardt, and J. Gilmer. The many faces of robustness: A critical analysis of out-of-distribution generalization. *2021 IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 8320–8329, 2021.
- [22] D. Hendrycks and T. G. Dietterich. Benchmarking neural network robustness to common corruptions and perturbations. *ArXiv*, abs/1903.12261, 2019.
- [23] W. Hu, G. Niu, I. Sato, and M. Sugiyama. Does distributionally robust supervised learning give robust classifiers? In *ICML*, 2018.
- [24] G. Huang, Z. Liu, and K. Q. Weinberger. Densely connected convolutional networks. *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 2261–2269, 2017.
- [25] A. Ilyas, S. Santurkar, D. Tsipras, L. Engstrom, B. Tran, and A. Madry. Adversarial examples are not bugs, they are features. In *NeurIPS*, 2019.
- [26] A. Javanmard, M. Soltanolkotabi, and H. Hassani. Precise tradeoffs in adversarial training for linear regression. In *Conference on Learning Theory*, pages 2034–2078. PMLR, 2020.
- [27] P. Kattakinda and S. Feizi. Focus: Familiar objects in common and uncommon settings. *ArXiv*, abs/2110.03804, 2021.
- [28] F. Khani and P. Liang. Removing spurious features can hurt accuracy and affect groups disproportionately. *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 2021.
- [29] P. Kirichenko, P. Izmailov, and A. G. Wilson. Last layer re-training is sufficient for robustness to spurious correlations, 2022.
- [30] P. W. Koh, S. Sagawa, H. Marklund, S. M. Xie, M. Zhang, A. Balsubramani, W. Hu, M. Yasunaga, R. L. Phillips, S. Beery, J. Leskovec, A. Kundaje, E. Pierson, S. Levine, C. Finn, and P. Liang. Wilds: A benchmark of in-the-wild distribution shifts. In *ICML*, 2021.
- [31] A. Kolesnikov and C. H. Lampert. Improving weakly-supervised object localization by micro-annotation. *ArXiv*, abs/1605.05538, 2016.
- [32] A. Krizhevsky. Learning multiple layers of features from tiny images. 2009.
- [33] D. Krueger, E. Caballero, J.-H. Jacobsen, A. Zhang, J. Binas, R. L. Priol, and A. C. Courville. Out-of-distribution generalization via risk extrapolation (rex). In *ICML*, 2021.
- [34] C. Laidlaw and S. Feizi. Functional adversarial attacks. In *NeurIPS*, 2019.
- [35] C. Laidlaw, S. Singla, and S. Feizi. Perceptual adversarial robustness: Defense against unseen threat models. *ArXiv*, abs/2006.12655, 2021.
- [36] A. Levine and S. Feizi. Improved, deterministic smoothing for l1 certified robustness. In *ICML*, 2021.
- [37] W.-A. Lin, C. P. Lau, A. Levine, R. Chellappa, and S. Feizi. Dual manifold adversarial robustness: Defense against lp and non-lp adversarial attacks. *ArXiv*, abs/2009.02470, 2020.
- [38] E. Z. Liu, B. Haghighi, A. S. Chen, A. Raghunathan, P. W. Koh, S. Sagawa, P. Liang, and C. Finn. Just train twice: Improving group robustness without training group information. In *ICML*, 2021.
- [39] A. Madry, A. Makelov, L. Schmidt, D. Tsipras, and A. Vladu. Towards deep learning models resistant to adversarial attacks. In *6th International Conference on Learning Representations, ICLR 2018, Vancouver, BC, Canada, April 30 - May 3, 2018, Conference Track Proceedings*. OpenReview.net, 2018.

- [40] M. Moayeri and S. Feizi. Sample efficient detection and classification of adversarial attacks via self-supervised embeddings. *2021 IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 7657–7666, 2021.
- [41] M. Moayeri, P. Pope, Y. Balaji, and S. Feizi. A comprehensive study of image classification model sensitivity to foregrounds, backgrounds, and visual attributes, 2022.
- [42] S.-M. Moosavi-Dezfooli, A. Fawzi, and P. Frossard. Deepfool: A simple and accurate method to fool deep neural networks. *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 2574–2582, 2016.
- [43] N. Papernot, P. Mcdaniel, X. Wu, S. Jha, and A. Swami. Distillation as a defense to adversarial perturbations against deep neural networks. *2016 IEEE Symposium on Security and Privacy (SP)*, pages 582–597, 2016.
- [44] J. Peters, P. Buhlmann, and N. Meinshausen. Causal inference by using invariant prediction: identification and confidence intervals. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)*, 78, 2015.
- [45] M. Ren, W. Zeng, B. Yang, and R. Urtasun. Learning to reweight examples for robust deep learning. In *ICML*, 2018.
- [46] A. Robey, G. J. Pappas, and H. Hassani. Model-based domain generalization. In *NeurIPS*, 2021.
- [47] A. Rosenfeld, R. S. Zemel, and J. K. Tsotsos. The elephant in the room. *ArXiv*, abs/1808.03305, 2018.
- [48] S. Sagawa, P. W. Koh, T. B. Hashimoto, and P. Liang. Distributionally robust neural networks for group shifts: On the importance of regularization for worst-case generalization. *ArXiv*, abs/1911.08731, 2019.
- [49] S. Sagawa*, P. W. Koh*, T. B. Hashimoto, and P. Liang. Distributionally robust neural networks. In *International Conference on Learning Representations*, 2020.
- [50] H. Salman, A. Ilyas, L. Engstrom, A. Kapoor, and A. Madry. Do adversarially robust imagenet models transfer better? In *ArXiv preprint arXiv:2007.08489*, 2020.
- [51] H. Salman, G. Yang, J. Li, P. Zhang, H. Zhang, I. P. Razenshteyn, and S. Bubeck. Provably robust deep learning via adversarially trained smoothed classifiers. *ArXiv*, abs/1906.04584, 2019.
- [52] M. Sandler, A. G. Howard, M. Zhu, A. Zhmoginov, and L.-C. Chen. Mobilenetv2: Inverted residuals and linear bottlenecks. *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 4510–4520, 2018.
- [53] S. Santurkar, A. Ilyas, D. Tsipras, L. Engstrom, B. Tran, and A. Madry. Image synthesis with a single (robust) classifier. In *NeurIPS*, 2019.
- [54] A. Shafahi, M. Najibi, A. Ghiasi, Z. Xu, J. P. Dickerson, C. Studer, L. S. Davis, G. Taylor, and T. Goldstein. Adversarial training for free! In *NeurIPS*, 2019.
- [55] H. Shah, K. Tamuly, A. Raghunathan, P. Jain, and P. Netrapalli. The pitfalls of simplicity bias in neural networks. *ArXiv*, abs/2006.07710, 2020.
- [56] R. Shetty, B. Schiele, and M. Fritz. Not using the car to see the sidewalk — quantifying and controlling the effects of context in classification and segmentation. *2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 8210–8218, 2019.
- [57] K. Simonyan and A. Zisserman. Very deep convolutional networks for large-scale image recognition. *CoRR*, abs/1409.1556, 2015.
- [58] S. Singla and S. Feizi. Second-order provable defenses against adversarial attacks. *ArXiv*, abs/2006.00731, 2020.
- [59] S. Singla and S. Feizi. Salient imagenet: How to discover spurious features in deep learning?, 2021.
- [60] S. Singla, M. Moayeri, and S. Feizi. Core risk minimization using salient imagenet, 2022.
- [61] S. Singla, B. Nushi, S. Shah, E. Kamar, and E. Horvitz. Understanding failures of deep networks via robust feature extraction. *2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 12848–12857, 2021.

- [62] C. Szegedy, W. Zaremba, I. Sutskever, J. Bruna, D. Erhan, I. Goodfellow, and R. Fergus. Intriguing properties of neural networks. In *International Conference on Learning Representations*, 2014.
- [63] H. Wang, S. Ge, Z. Lipton, and E. P. Xing. Learning robust global representations by penalizing local predictive power. In *Advances in Neural Information Processing Systems*, pages 10506–10518, 2019.
- [64] P. Welinder, S. Branson, T. Mita, C. Wah, F. Schroff, S. Belongie, and P. Perona. Caltech-UCSD Birds 200. Technical Report CNS-TR-2010-001, California Institute of Technology, 2010.
- [65] E. Wong, L. Rice, and J. Z. Kolter. Fast is better than free: Revisiting adversarial training. *CoRR*, abs/2001.03994, 2020.
- [66] E. Wong, S. Santurkar, and A. Madry. Leveraging sparse linear layers for debuggable deep networks. *ArXiv*, abs/2105.04857, 2021.
- [67] J. Xiao, J. Hays, K. A. Ehinger, A. Oliva, and A. Torralba. Sun database: Large-scale scene recognition from abbey to zoo. *2010 IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pages 3485–3492, 2010.
- [68] K. Y. Xiao, L. Engstrom, A. Ilyas, and A. Madry. Noise or signal: The role of image backgrounds in object recognition. In *International Conference on Learning Representations*, 2021.
- [69] S. Xie, R. B. Girshick, P. Dollár, Z. Tu, and K. He. Aggregated residual transformations for deep neural networks. *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 5987–5995, 2017.
- [70] N. Ye, K. Li, H. Bai, R. Yu, L. Hong, F. Zhou, Z. Li, and J. Zhu. Ood-bench: Quantifying and understanding two dimensions of out-of-distribution generalization. 2021.
- [71] J. R. Zech, M. A. Badgeley, M. Liu, A. B. Costa, J. J. Titano, and E. K. Oermann. Variable generalization performance of a deep learning model to detect pneumonia in chest radiographs: A cross-sectional study. *PLoS Medicine*, 15, 2018.
- [72] C. H. Zhang, K. Zhang, and Y. Li. A causal view on robustness of neural networks. *ArXiv*, abs/2005.01095, 2020.
- [73] H. R. Zhang, Y. Yu, J. Jiao, E. P. Xing, L. E. Ghaoui, and M. I. Jordan. Theoretically principled trade-off between robustness and accuracy. In *ICML*, 2019.
- [74] J. Zhang, A. K. Menon, A. Veit, S. Bhojanapalli, S. Kumar, and S. Sra. Coping with label shift via distributionally robust optimisation. *ArXiv*, abs/2010.12230, 2021.
- [75] X. Zhang, X. Zhou, M. Lin, and J. Sun. Shufflenet: An extremely efficient convolutional neural network for mobile devices. *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 6848–6856, 2018.
- [76] Y. Zhang, M. Gong, T. Liu, G. Niu, X. Tian, B. Han, B. Scholkopf, and K. Zhung. Causaladv: Adversarial robustness through the lens of causality. In *ICLR*, 2022.

A Limitations of Theorem 1

While Theorem 1 provides an understanding of the tradeoff between adversarial and natural distributional robustness, there are some limitations. Firstly, the results consider a setting where the core and spurious features are completely disentangled, i.e, they each represent different parts of the input. In practice, spurious features may be entangled with the core features (e.g., the color of an image may represent a spurious feature.) In addition, our results mainly consider the *goal* of adversarial training as we focus on the expected loss $L_{p,\epsilon}(\theta)$, rather than its finite-sample variant. This is because even for an ℓ_2 adversary, characterizing the finite-sample behaviour of adversarial training is difficult and requires careful assumptions on the asymptotic behaviour of the parameters (e.g., see Theorem 3.3 in [26]). We leave exploring these directions to future work. Even so, we believe our theoretical results are of interest to the community since disjoint features already capture a wide variety of spurious correlations, e.g., background correlations, as well as examples where a spurious object is present in the image. The main goal of our theoretical analysis is to show the existence of explicit tradeoffs between adversarial and distributional robustness and build practical insights using those results.

B Societal Impact

Our work touches on two important notions of robustness for the safe and fair deployment of deep models in the wild. We hope our results lead to careful analysis of all modes of robustness, and the interplay between them, before deep models are used in sensitive applications. While our results create tension with some previous works [72, 76, 25], we stress that we do not wish to diminish their work; instead, we hope our work reveals the vast nuance associated with spurious correlations, which can help and hurt models in various ways. Lastly, we release all code to encourage future work.

C Varying the Number of Core and Spurious Features

In this section, we further analyze the plateauing behaviour of the performance of the linear model observed in Figure 2 for different values of core and spurious features with the ℓ_1 and ℓ_2 adversarial training.

We first focus on the ℓ_1 case. We consider different values for the number of core features c and total features m and measure NFS for different values of adversarial budget ϵ as in Figure 2. The matrix Σ is constructed using Equation (4) as before, with modified number of rows and columns based on the values of c, p . Similarly, θ^{opt} is constructed as before, with the core coordinates set to 1 and the spurious coordinates set to 0. The value of η is fixed at 0.5. The results are shown in Figure 10.

As shown in the Figure, when using m total features and c core features, NFS plateaus at $\frac{m-c}{m}$ for large values of ϵ . Intuitively, this is because of the structure of the optimization problem (2). Recall that when using the ℓ_1 norm, the value of q in (2) equals ∞ . As such, adversarial training tries to find a parameter θ that has a low ℓ_∞ norm and is “close” (as measured by σ_θ) to θ^{opt} . The ℓ_∞ penalty encourages values of θ that are uniform across the coordinates. Since there are $m - c$ spurious features and m total features, this leads to models that have an NFS value of $\frac{m-c}{m}$.

We further repeat the above experiment with ℓ_2 norm. The results are shown in Figure 11. As seen in the figure, we see the same qualitative results as in Figure 2, with higher values of NFS when increasing numbers of spurious features.

D Proof of Theorem 1

Proof. We first claim that

$$\max_{\|\delta\| \leq \epsilon} (Y - \langle X + \delta, \theta \rangle)^2 = \left(|Y - \langle X, \theta \rangle| + \epsilon \cdot \|\theta\|_q \right)^2$$

To see why this holds, note that for all δ satisfying $\|\delta\|_p \leq \epsilon$,

$$\begin{aligned} |Y - \langle X + \delta, \theta \rangle| &\stackrel{(a)}{\leq} |Y - \langle X, \theta \rangle| + |\langle \delta, \theta \rangle| \\ &\stackrel{(b)}{\leq} |Y - \langle X, \theta \rangle| + \epsilon \cdot \|\theta\|_q, \end{aligned}$$

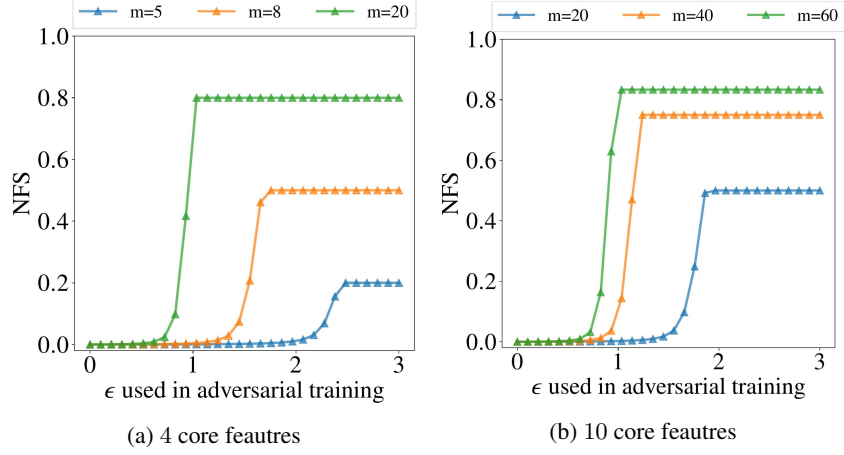


Figure 10: Analysis of NFS for the linear model when using the ℓ_1 norm in adversarial training. Each figure measures the reliance of the model on spurious features (measured by NFS) while varying the adversarial training budget ϵ , using different number of *total* features m . The number of core features is kept constant and set to 4 in Figure (a) and to 10 in Figure (b).

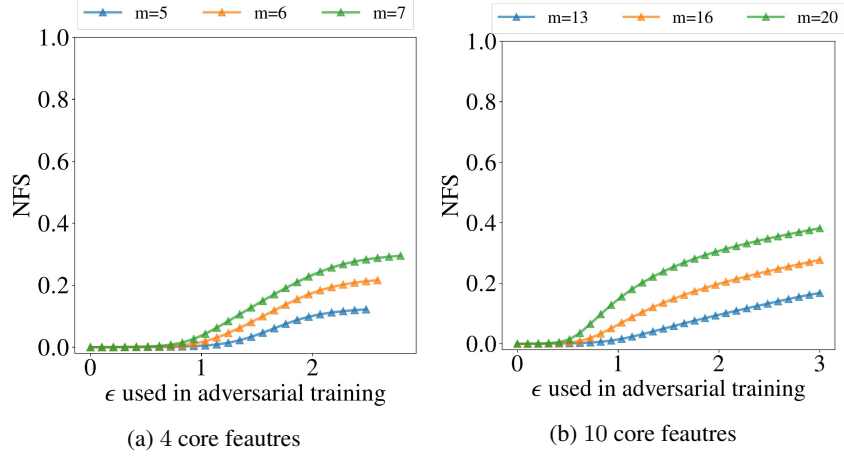


Figure 11: Analysis of NFS for the linear model when using the ℓ_2 norm in adversarial training. Each figure measures the reliance of the model on spurious features (measured by NFS) while varying the adversarial training budget ϵ , using different number of *total* features m . The number of core features is kept constant and set to 4 in Figure (a) and to 10 in Figure (b).

where (a) follows from the triangle inequality and (b) follows from Hölder's inequality. With a suitable choice of δ , we can achieve equality for (b). As $\|\theta\|_q = \|-\theta\|_q$, at least one of $\{\delta, -\delta\}$ would further achieve equality for (a). As maximizing $|\cdot|$ is equivalent to maximizing $(\cdot)^2$, (3) is proved.

Given (3), we can rewrite (1) as

$$\begin{aligned}
L_{p,\epsilon} &= \mathbb{E} \left[\left(|Y - \langle X, \theta \rangle| + \epsilon \cdot \|\theta\|_q \right)^2 \right] \\
&= \mathbb{E} \left[(Y - \langle X, \theta \rangle)^2 \right] + \epsilon^2 \cdot \|\theta\|_q^2 + 2 \cdot \epsilon \cdot \|\theta\|_q \cdot \mathbb{E} [|Y - \langle X, \theta \rangle|] \\
&\stackrel{(a)}{=} \mathbb{E} \left[(\langle X, \theta - \theta^{\text{opt}} \rangle + W)^2 \right] + \epsilon^2 \cdot \|\theta\|_q^2 + 2 \cdot \epsilon \cdot \|\theta\|_q \cdot \mathbb{E} [\langle X, \theta - \theta^{\text{opt}} \rangle + W],
\end{aligned}$$

Where for (a) we have used the fact that $Y = \langle X, \theta \rangle + W$.

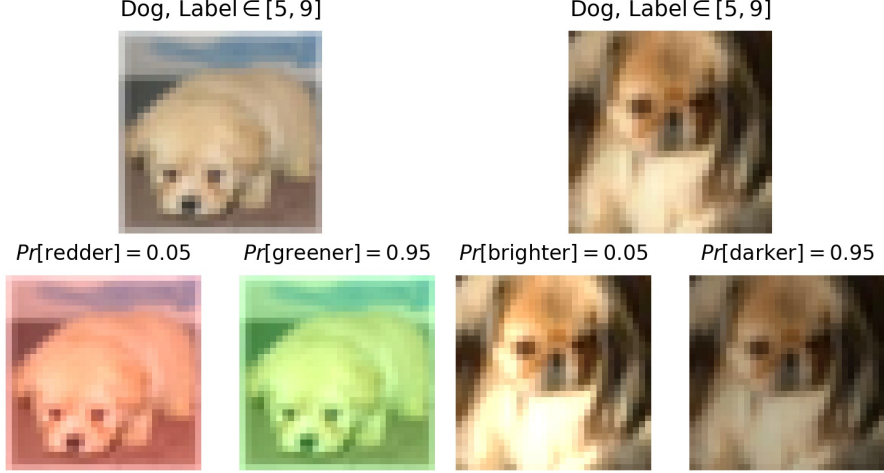


Figure 12: Color Shift, $\rho = 19 : 1$

Figure 13: Lighting Shift, $\rho = 19 : 1$

Define v_θ as $\langle X, \theta - \theta^{\text{opt}} \rangle + W$. As X was assumed to be sampled from $N(0, \Sigma)$, v_θ is distributed as $N(0, \sigma_\theta^2)$. It follows that

$$\begin{aligned}
 L_{p,\epsilon} &= \mathbb{E} \left[\left(\langle X, \theta - \theta^{\text{opt}} \rangle + W \right)^2 \right] + \epsilon^2 \cdot \|\theta\|_q^2 + 2 \cdot \epsilon \cdot \|\theta\|_q \cdot \mathbb{E} [\langle X, \theta - \theta^{\text{opt}} \rangle + W] , \\
 &= \mathbb{E} [v_\theta^2] + \epsilon^2 \cdot \|\theta\|_q^2 + 2 \cdot \epsilon \cdot \|\theta\|_q \cdot \mathbb{E} [v_\theta] \\
 &\stackrel{(a)}{=} \sigma_\theta^2 + \epsilon^2 \cdot \|\theta\|_q^2 + 2 \cdot c_1 \cdot \epsilon \cdot \|\theta\|_q \cdot \sigma_\theta \\
 &= (c_1^2 + c_2) \cdot \sigma_\theta^2 + \epsilon^2 \cdot \|\theta\|_q^2 + 2 \cdot c_1 \cdot \epsilon \cdot \|\theta\|_q \cdot \sigma_\theta \\
 &= c_2 \cdot \sigma_\theta^2 + (c_1 \sigma_\theta + \epsilon \cdot \|\theta\|_q)^2
 \end{aligned}$$

where for (a) we have used the fact that $\mathbb{E} [|N(0, \sigma^2)|] = c_1 \cdot \sigma$. This proves (2) as claimed.

As for convexity, σ_θ is convex in θ since it can be written as $\left\| [\Sigma^{\frac{1}{2}}(\theta - \theta^{\text{opt}}), \sigma_w] \right\|_2$ where $[\cdot, \cdot]$ denotes the vector stacking operation. As $c_1 \sigma_\theta + \epsilon \cdot \|\theta\|_q$ is always positive and $x \rightarrow x^2$ is convex and increasing for $x \geq 0$, this implies that $(c_1 \sigma_\theta + \epsilon \cdot \|\theta\|_q)^2$ is convex as well. Finally $c_2 \sigma_\theta^2$ is convex as $c_2 > 0$ and therefore (1) is convex in θ . $\square$

E Additional Details on Reverse Effect (Section 4.3)

Our final empirical observation is that the presence of a spurious feature (in both training and test distributions) can lead to increased adversarial robustness. This more directly creates tension with claims that adversarial vulnerability is born out of spurious feature reliance. We refer to this as the ‘reverse effect’, in relation to our primary empirical and theoretical finding that adversarial training increases spurious feature reliance. We now elaborate on the experimental setup discussed in Section 4.3, reproduce the results with a different spurious feature, and finally appeal to ImageNet-9 to demonstrate this effect using a more realistic spurious feature (i.e. backgrounds).

E.1 Experimental Setup

Overview. We inject spurious correlations to the CIFAR10 dataset. Based on the class label, we adjust half the images (i.e. with class label from 5 to 9) to shift in one direction with high probability. For example, a dog image is made greener with probability 0.95, corresponding to a majority-to-minority group ratio of $\rho = 19 : 1$. With probability 0.05, we shift in the other direction (e.g. make redder). We then standardly train a ResNet18 from scratch on the dataset with the spurious feature injected for the 10-way CIFAR classification task. Importantly, we evaluate robust accuracy with the spurious feature retained, and then compare adversarial robustness of models trained under data with different strengths of the injected spurious correlation.

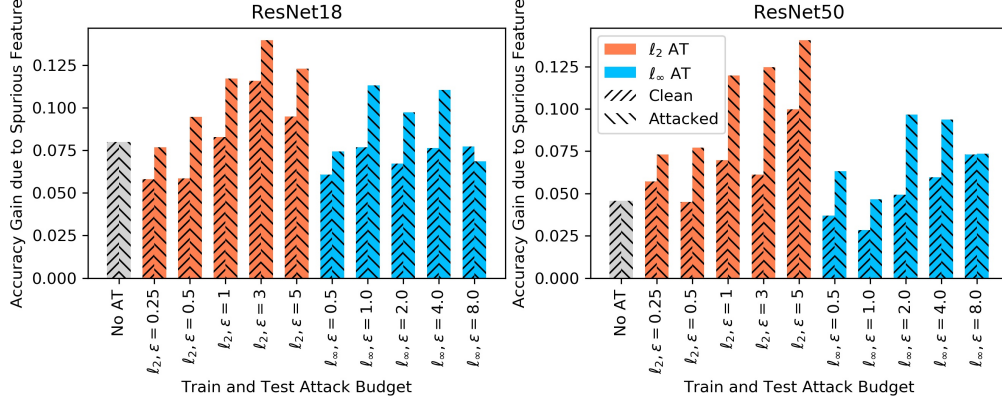


Figure 15: Background Gap (difference in accuracy on MIXED-SAME and MIXED-RAND) for clean and adversarially attacked images. Across models, background gap is larger when considering accuracy under attack, suggesting that the presence of a spurious correlation in training data makes the model more adversarially robust over the same distribution.

Figure 9 and 14 show that for two distinct spurious features (color and lighting), robust accuracy is higher when the spurious correlation is stronger. Notably, the gain is larger than the gain in standard accuracy. Intuitively, we see that relying on the predictive power of the spurious feature is helpful for standard accuracy, and especially for accuracy under adversarial attack. Despite being irrelevant to the true labeling function, the spurious feature can improve model performance, and indeed even lead to better adversarial robustness.

Details. Color shift is achieved by increasing all pixel intensities along one channel by 0.25. Lighting shift is achieved by simply scaling an input by 1.25 to make brighter or 0.75 to make darker. All images are clamped to remain in the $[0, 1]$ pixel range after spurious feature injection. Models are trained for 20 epochs using an Adam optimizer with a learning rate of 0.001 and weight decay of $1e - 4$.

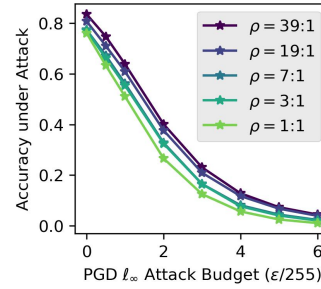


Figure 14: Reverse effect using spurious feature of *lighting*. Main text figure uses color as spurious feature.

E.2 Leveraging ImageNet-9

We now demonstrate the observed reverse effect on the higher resolution ImageNet-9 dataset, leveraging the natural and ubiquitous spurious feature of backgrounds. We finetune pretrained models on MIXED-SAME and MIXED-RAND separately, and evaluate each model’s accuracy under attack on the same split that they were trained over. Further, we leverage the adversarially trained models from test suite in this experiment. This way, accuracy under attack is more informative, as the models are trained to expect attacks (i.e. we are not imposing any distribution shifts that would lead to unexpected model behavior). Along this vain, we attack each backbone with the same norm and ϵ that it was pretrained over.

Figure 15 shows the gain in accuracy for the models trained and evaluated on MIXED-SAME compared to those using MIXED-RAND. We see that the presence of background correlations increases both standard and robust accuracy for all models (i.e. gains are positive). Further, gains in accuracy under attack are larger than gains in standard accuracy in nearly all cases. Thus, it seems like the added predictive power of the spurious background feature has a significantly nontrivial impact on improving adversarial robustness, contradicting many existing arguments on the link between spurious correlations and adversarial vulnerability.

AT Norm	ϵ	ResNet18	ResNet50	Wide ResNet50 (2x)
No Adv Training		69.79	75.80	76.97
ℓ_2	0.25	67.43	74.14	76.21
ℓ_∞	0.5/255	66.13	73.73	75.82
ℓ_2	0.5	65.49	73.16	75.11
ℓ_∞	1/255	63.46	72.05	74.65
ℓ_2	1	62.32	70.43	73.41
ℓ_∞	2/255	59.63	69.10	72.35
ℓ_2	3	53.12	62.83	66.90
ℓ_∞	4/255	52.49	63.86	68.41
ℓ_2	5	45.59	56.13	60.94
ℓ_∞	8/255	42.11	54.53	60.82

Table 1: Clean ImageNet accuracy for test suite of ℓ_2 and ℓ_∞ adversarially trained ResNets over varying ϵ . Observe that the i^{th} ℓ_2 AT model has similar clean accuracy to the i^{th} ℓ_∞ AT model.

	ShuffleNet	MobileNet	VGG	DenseNet	ResNeXt
No AT	64.25	65.26	73.66	77.37	77.38
ℓ_2 AT, $\epsilon = 3$	43.32	50.40	57.19	66.98	66.25

Table 2: Clean ImageNet accuracy for five additional architectures considered.

F Adversarially Robust Model Test Suite (Section 3)

F.1 Model Details

We utilize the treasure trove of open-source adversarially trained models, contributed by [50], accessible at <https://github.com/Microsoft/robust-models-transfer>. For completeness, we now provide details on the models we use, though we refer readers to Appendix A.1 of the original text, where the information we share now is sourced.

Training All models were trained on ImageNet in batches of 512 samples, using SGD optimizer with momentum of 0.9 and weight decay of $1e-4$, for a total of 90 epochs, with learning rate dropping by a factor of 10 every 30 epochs. The standard procedure of [39] was performed to adversarially train models, using 3 projected gradient descent steps with a step size $\frac{2}{3}\epsilon$ for the attack budget ϵ .

Selected Models We focus our empirical study on the ResNet architecture [20] because of its wide spread popularity. Specifically, we study ResNet18s and ResNet50s that are adversarially trained under the ℓ_2 norm, for $\epsilon \in \{0.25, 0.5, 1, 3, 5\}$, and ℓ_∞ norm, for $\epsilon \in \{0.5/255, 1/255, 2/255, 4/255, 8/255\}$, as well as standardly trained baselines.

Table 1 shows the standard accuracies for these models. Note that we at times compare between the ℓ_2 and ℓ_∞ adversarially trained models (e.g. figure 6). We acknowledge that direct comparisons are challenging because the threat model under which adversarial robustness is optimized for are different. However, we note that standard accuracies of the i^{th} ℓ_2 AT model is roughly the same as that of the i^{th} ℓ_∞ AT model, suggesting that those models lie in similar points of the accuracy-robustness tradeoff.

Additional Models. We extend our analysis to other architectures. We replicate all pretrained-model experiments on the Wide ResNet50 (2x) backbones, for which we have checkpoints for each of the five ϵ values for both ℓ_2 and ℓ_∞ norms. We also inspect MobileNetv2 [52], DenseNet161 [24], ResNeXt50_32x4d [69], ShuffleNet [75], and VGG16_bn [57]. For each of these five architectures, we compare an ℓ_2 adversarially trained model with $\epsilon = 3$ to a standardly trained baseline.

F.2 Experimental Details

ObjectNet and ImageNet-C [5, 22]. We report raw accuracies under noise, blur, and digital corruption types for ImageNet-C, as opposed to relative corruption error. For ObjectNet, we map ImageNet predictions to the set of 113 overlapping classes in ObjectNet. **RIVAL10 (RFS) and**

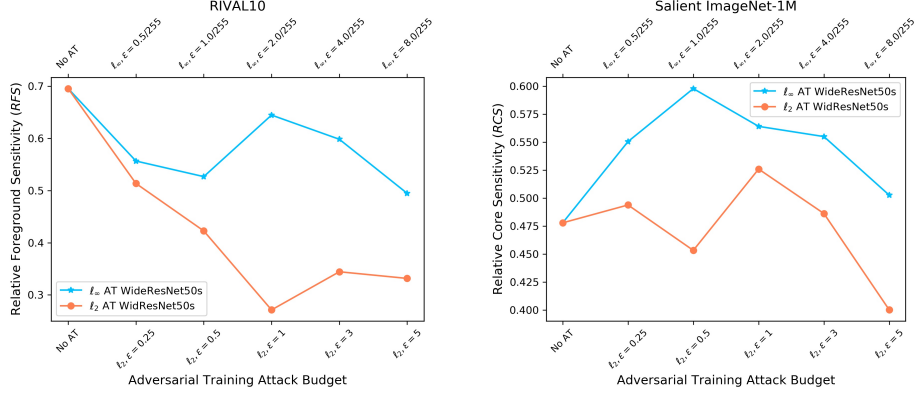


Figure 17: *RFS* and *RCS* for WideResNet50s. Sensitivity to core and foreground regions are reduced for higher ϵ , especially for ℓ_2 AT models and for *RFS*, computed over the RIVAL10 dataset, where background correlations are stronger.

Salient ImageNet-1M (*RCS*) [41, 60]. *RFS* computation requires finetuning a final linear layer over fixed features for the coarse-grained ten way RIVAL10 classification. *RCS* operates on models off the shelf, directly inspecting accuracies over ImageNet classes (and samples, with region-based noise corruption). **ImageNet-9 and Waterbirds** [68, 49]. ImageNet-9 accuracies are obtained by mapping off-the-shelf model predictions to the nine coarse labels deterministically. Waterbirds requires finetuning, which we do over fixed features. For RIVAL10 and Waterbirds finetuning, we use Adam with learning rate of $1e-4$ and weight decay of $1e-5$ for 20 and 15 epochs respectively.

F.3 Results on Extended Model Test Suite

We now corroborate all our empirical findings on new backbones, expanding our analysis to 21 new models (including 10 AT WideResNet50s over both ℓ_2 and ℓ_∞ norms) over six architectures.

WideResNets. We corroborate all our empirical findings on ResNet18s and ResNet50s on the WideResNet50 (2x) architecture. Figure 16 shows that accuracy drop in AT models is more severe on distribution shifts that break spurious correlations (ObjectNet), unlike the accuracy drop due to corruption of both core and spurious features (ImageNet-C), which can likely be explained by the reduced standard accuracy of AT models.

Figure 17 shows reduced sensitivity to core and foreground regions for AT models. Again, the effect is more pronounced for ℓ_2 adversarial training and for larger ϵ . Also, we again see that decrease in *RCS* is less consistent than the drop in *RFS*. We conjecture that the diversity and fine-grain Salient ImageNet classification task reduces the strength of spurious correlations present in the data, thus diluting our observed effects of adversarial training on spurious feature reliance.

Lastly, figure 18 shows the drop in accuracy due to breaking spurious background correlations is larger for AT models. Indeed, the absolute background gap (IN-9) for the WideResNet50 under ℓ_2 AT with $\epsilon = 5$ is 50% larger than the gap for the standardly trained baseline. We note that the absolute gaps are smaller in some cases. We believe the lower standard accuracy of AT models may contribute to this, as there is less accuracy to drop from. Nonetheless, it is intriguing that in some cases, ℓ_∞ adversarial training seems to reduce spurious feature reliance; while our theory explains how a spurious feature can be completely ignored under ℓ_∞ training, it does not explain cases where spurious feature reliance is reduced compared to standard training. We believe this is an interesting direction for future work.

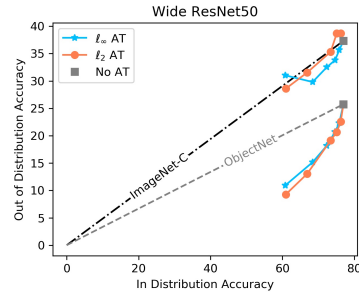


Figure 16: ObjectNet, ImageNet-C, and ImageNet accuracies for WideResNet50s.

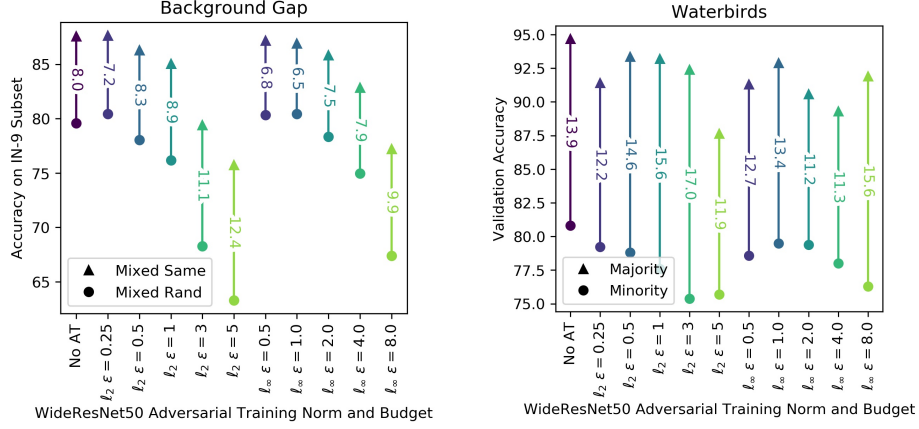
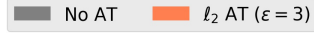
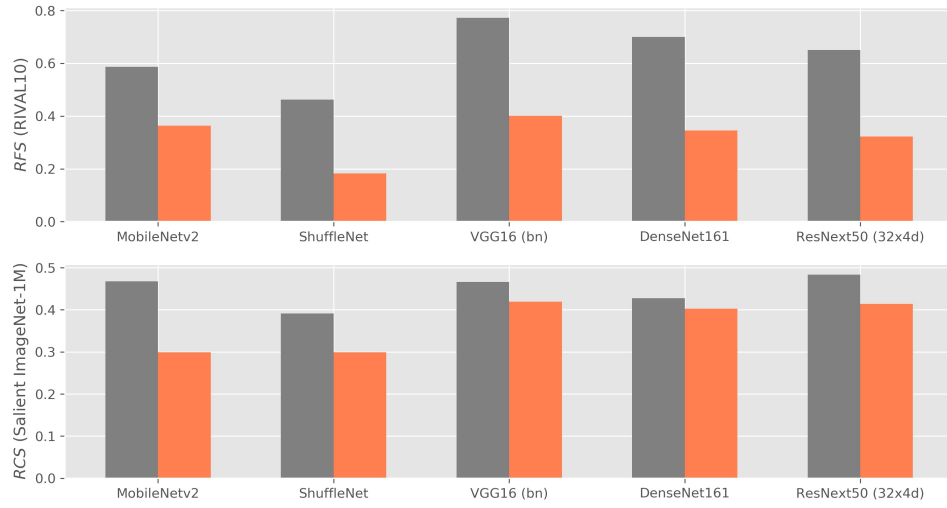


Figure 18: Background Gap (IN-9) and Waterbirds gap for WideResNet50s. AT models, especially under ℓ_2 norm, see larger accuracy drops when spurious correlations are broken.

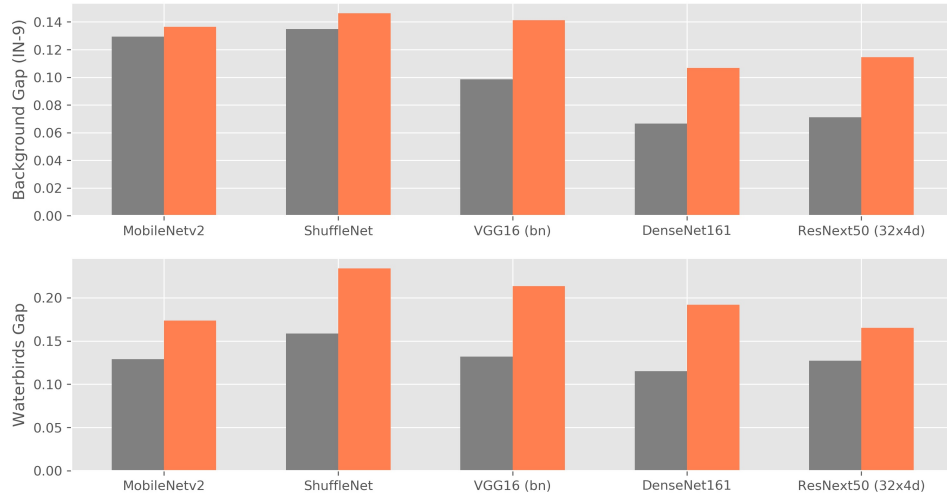
Other backbones. We now show results for ten other models, half of which are ℓ_2 adversarially trained with $\epsilon = 3$, while the others are standardly trained. Figure 19 summarizes our results, corroborating each of our empirical findings.



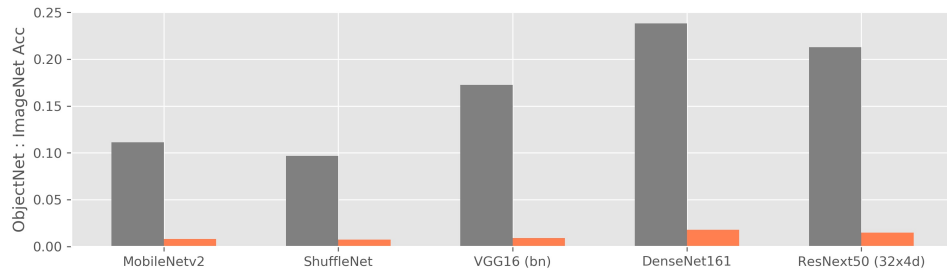
(a) Legend. We compare ℓ_2 adversarially trained models to standardly trained baselines for five new backbones.



(b) Lower *RFS* (*RCS*) entails Lower Foreground (Core Feature) Sensitivity



(c) Higher Gap entails Greater Background/Spurious Sensitivity



(d) Lower Ratio entails Lower Natural Distributional Robustness

Figure 19: Corroborating findings on additional backbones.