

Do Students Debias Like Teachers? On the Distillability of Bias Mitigation Methods

Anonymous Authors¹

Abstract

Knowledge distillation (KD) is an effective method for model compression and transferring knowledge between models. However, its effect on model’s robustness against spurious correlations, shortcuts and task-irrelevant features that degrade performance on out-of-distribution data remains underexplored. This study investigates the effect of knowledge distillation on natural language inference (NLI) and image classification tasks, with a focus on the transferability of “debiasing” capabilities from teacher models to student models. Through extensive experiments, we illustrate several key findings: (i) the effect of KD on debiasing performance depends on the underlying debiasing method, the relative scale of the models involved, and the size of the training set; (ii) KD effectively transfers debiasing capabilities when teacher and student are similar in scale (number of parameters); (iii) KD may amplify the student model’s reliance on spurious features, and this effect does not diminish as the teacher model scales up; and (iv) although the overall robustness of a model may remain stable post-distillation, significant variations can occur across different types of biases; and Given the above findings, we propose three effective solutions to improve the distillability of debiasing methods: developing high quality data for augmentation, implementing iterative knowledge distillation, and initializing student models with weights obtained from teacher models.

1. Introduction

Machine learning models are susceptible to biases or spurious correlations in datasets, commonly known to as “short-

cuts” or “dataset biases”. Models that rely on shortcuts can achieve high performance on in-domain test sets or over-represented groups by exploiting superficial correlations between features and labels. However, these models suffer significant performance degradation on out of distribution or challenging test data, such as swapped subject and object (McCoy et al., 2019) in natural language understanding, or under-represented groups, such as “Male” subjects with “Blond Hair” (Liu et al., 2015) in image classification.

Despite recent advancements in bias mitigation (Guo et al., 2023; Chew et al., 2024; Li et al., 2023; Cheng & Amiri, 2024) and knowledge distillation (Stanton et al., 2021; Sultan, 2023), their integration is largely unexplored. This work studies the following research questions (RQs):

- **RQ1:** *To what extent can knowledge distillation transfer debiasing capabilities between models?*
- **RQ2:** *Can knowledge distillation train less biased models compared to standard training?*
- **RQ3:** *Do different debiasing methods show consistent patterns in task performance and debiasing effectiveness before and after knowledge distillation?*

Answering these questions will help us understand the efficacy of knowledge distillation in handling dataset biases, its underlying mechanisms, and its role in developing new training methods for bias mitigation.

We answer these questions by designing and conducting an empirical analysis on natural language understanding and image classification tasks. Our analyses show that: (i) the effect of knowledge distillation on debiasing performance depends on the underlying debiasing method, the relative scale of the models involved, and the size of the training set; (ii) knowledge distillation effectively transfers debiasing capabilities when teacher and student are similar in scale (number of parameters); (iii) knowledge distillation may amplify the student model’s reliance on spurious features, and this effect does not diminish as the teacher model scales up; and (iv) although the overall robustness of a model may remain stable post-distillation, significant variations can occur across different types of biases; and (v) consistent

¹Anonymous Institution, Anonymous City, Anonymous Region, Anonymous Country. Correspondence to: Anonymous Author <anon.email@domain.com>.

Preliminary work. Under review by the International Conference on Machine Learning (ICML). Do not distribute.

transfer patterns sometimes emerge, such as performance gap between teacher and student on out-of-distribution (OOD) data, suggesting the possibility of predictable changes in robustness after distillation. Given the above findings, we propose three effective solutions to improve the distillability of debiasing methods: developing high quality data for augmentation, implementing iterative knowledge distillation, and initializing student models with weights obtained from teacher models.

The contributions of this paper are:

- to the best of our knowledge, we present the first study to investigate the effect of knowledge distillation on model robustness against dataset biases, and analyze the distillability of bias mitigation methods across both language and vision tasks;
- our analysis reveals unique patterns in how knowledge distillation affects robustness to spurious correlations across different backbones and debiasing methods; and
- we propose three strategies to improve the distillability of debiasing methods and provide insights for future development of bias mitigation techniques.

2. Knowledge Distillation and Debiasing

2.1. Problem Formulation

We investigate the effect of knowledge distillation (KD) on debiasing methods. We define *distillability of debiasing methods* as the amount of performance maintained before and after distilling a debiased model. We define *contribution of KD* as the performance improvement gained by training a debiasing method with KD over training without KD.

2.2. Notation and Training Setup

Let f and g denote models trained without knowledge distillation and with knowledge distillation respectively. In this paper, we use subscript $\mathcal{T}$ and $\mathcal{S}$ to denote teacher and student scales respectively. As illustrated in Figure 1, we train the following models for each debiasing method M_i : (i) we train M_i from scratch for both teacher and student scales to obtain $f_{\mathcal{T}}$ and $f_{\mathcal{S}}$, see Figure 1(a). (ii) Then for every scale $\mathcal{T} > \mathcal{S}$, we distill the knowledge from $f_{\mathcal{T}}$ to $g_{\mathcal{T} \rightarrow \mathcal{S}}$, see Figure 1(b). Given a debiasing method M and the three models obtained above ($f_{\mathcal{T}}$, $f_{\mathcal{S}}$, and $g_{\mathcal{T} \rightarrow \mathcal{S}}$), we conduct the following comparisons:

- **C1: Teacher ($f_{\mathcal{T}}$) vs. Student ($g_{\mathcal{T} \rightarrow \mathcal{S}}$).** This comparison reveals if knowledge distillation can distill debiasing capability between models and if it affects model’s robustness to spurious correlations, which answers *RQ1* (§4).

- **C2: Non-KD vs. KD,** realized by comparing $f_{\mathcal{S}}$ vs. $g_{\mathcal{T} \rightarrow \mathcal{S}}$. This comparison demonstrates if training bias mitigation networks can benefit from external knowledge from teacher models, which answers *RQ2* (§5).
- **C3: Comparison between debiasing methods** (M_i vs. M_j). This comparison provides insights into differences between debiasing methods, which answers *RQ3* (§6).

We note that when $\mathcal{T} = \mathcal{S}$, C1 and C2 are essentially the same comparison. To avoid duplicate discussion, we will present results when $\mathcal{T} = \mathcal{S}$ in C2.

3. Experimental Setup

For consistency and fair comparison with previous debiasing works in NLU (Jeon et al., 2023; Reif & Schwartz, 2023) and image classification (Kirichenko et al., 2023; LaBonte et al., 2023; Li et al., 2023), we adopt commonly used experimental setups, including choice of backbone models, datasets, evaluation protocols, and debiasing methods. In addition, all experiments are repeated three times with different random seeds to account for any stochastic effect.

Backbones We conduct experiments on a series of BERT (Devlin et al., 2019; Turc et al., 2019), T5 (Tay et al., 2022), ResNet (He et al., 2016), and ViT (Dosovitskiy et al., 2021) backbones of different scales, shown in Table 1. These backbones are chosen for several reasons: BERT and ResNet are commonly employed in prior works, which enables consistent comparisons. In addition, ViT and T5 are commonly used backbones for vision and language tasks, but relatively less experimented in prior debiasing works, which allows investigating the generalizability of our findings beyond existing research. Finally, each backbone is associated with a series of publicly available pre-trained checkpoints of different scales, with consistent network architecture and pre-training data, which enables cross-scale distillation and comparisons.

Table 1. Different scales of backbones in our experiments. h and d denote number of hidden layers and size of hidden dimension respectively. T, S, M, B, L refer to Tiny, Small, Medium, Base and Large version of the backbone. See Appendix for more details.

SCALE	BERT		T5		RESNET		ViT	
	h	d	h	d	h	d	h	d
T	2	128	4	256	18	512	12	192
S	4	256	8	384	34	512	12	384
M	8	512	16	512	50	2048	12	768
B	12	768	24	768	101	2048	24	1024
L	24	1024	48	1024	152	2048	32	1280

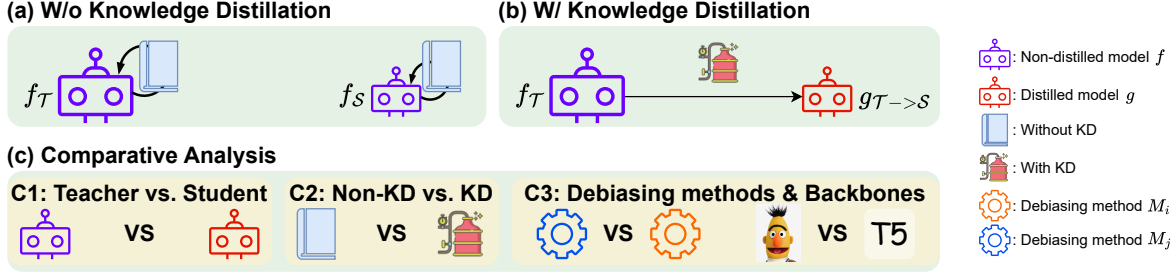


Figure 1. Framework for the analysis of distillability of debiasing methods. (a) **training from scratch**: we train a debiasing method M_i from scratch without knowledge distillation on different scales (teacher $\mathcal{T}$ and student $\mathcal{S}$ such that $\mathcal{T} > \mathcal{S}$) to obtain models $f_{\mathcal{T}}$, $f_{\mathcal{S}}$ respectively. (b) **Training with knowledge distillation**: we apply knowledge distillation to transfer knowledge from teacher ($f_{\mathcal{T}}$) to student ($g_{\mathcal{T} \rightarrow \mathcal{S}}$). (c) **Assessment**: C1 determines if knowledge distillation can transfer the debiasing capability from teacher ($f_{\mathcal{T}}$) to student ($g_{\mathcal{T} \rightarrow \mathcal{S}}$), C2 determines the contribution of knowledge distillation in training a debiased model, and C3 compares different debiasing methods and backbones under knowledge distillation.

Evaluation To provide a comprehensive evaluation of robustness against spurious correlations, we compare the teacher $f_{\mathcal{T}}$ and the student $g_{\mathcal{T} \rightarrow \mathcal{S}}$ from the following perspectives:

- **Performance on in-domain test set (ID, $\uparrow$)**: This is the average performance on in-domain test set. A robust model should achieve high performance on this set to demonstrate general capability.
- **Performance on out-of-domain test set / worst-group samples (OOD, $\downarrow$)**: For text datasets, we evaluate models on OOD test sets, comprised with specially crafted hard samples (McCoy et al., 2019). Such samples require real task-related signals to predict, where biased models fall short. For image datasets, samples are divided into groups based on their label and spurious attributes. The worst performance on all subgroups indicates the robustness to spurious correlations (Yang et al., 2023).
- **Spurious gap (Spu. Gap, $\downarrow$)**: This metric calculates the performance gap between ID and OOD, which quantifies a model’s vulnerability to spurious correlations. Ideally, a robust model should have high performances on both ID and OOD with a small spurious gap.
- **Centralized Kernel Alignment (CKA)** CKA is a commonly adopted technique to measure the similarity between activation matrices or hidden representations of neural networks (Kornblith et al., 2019; Cortes et al., 2012). Following previous work (Raghu et al., 2021; Nguyen et al., 2021), we use CKA by first probing the intermediate representations from each layer and then comparing all pairwise similarities between representations of the teacher and student models, under linear kernel CKA.

Similarly, we compare KD and Non-KD as above. We compute F1 score on QQP and accuracy on other datasets.

Datasets We use the following datasets

- **CelebA (Liu et al., 2015)** consists of 16k images of celebrity faces, where the objective is to predict “Blond_Hair” given “Male” as a spurious attribution.
- **Waterbird (Sagawa et al., 2020)** consists of synthetic images of birds from CUB dataset (Wah et al., 2011) and backgrounds (land & water) from Places (Zhou et al., 2018) dataset. The objective is to correctly infer “land bird” or “water bird,” given the background as misleading information.
- **MNLI (Williams et al., 2018)** consists of 39k natural language inference (NLI) samples from various domains, where the objective is to classify relationship between a premise and a hypothesis as “Entailment”, “Contradiction”, or “Neutral”. Previous studies discover that models are prone to negation words, lexical overlap, and sub-sequence biases in NLI task (Naik et al., 2018; Mendelson & Belinkov, 2021). We use HANS (McCoy et al., 2019) as the out-of-distribution test set (OOD) and SNLI (Bowman et al., 2015) as the transfer test set (Transfer), detailed below.
- **QQP (Sharma et al., 2019)** is a paraphrase identification (PI) dataset with 43k samples, where the objective is to predict if two questions are paraphrases of each other. Similar to MNLI, models are likely to be misled by lexical overlap between two questions. We exploit PAWS (Zhang et al., 2019) as the out-of-distribution test set (OOD) and MRPC (Dolan & Brockett, 2005) as the transfer test set (Transfer), detailed below.

Debiasing Methods Experiments are conducted on a comprehensive list of commonly used debiasing methods, each of which is designed with special formulation and assumptions. We use (a) Empirical Risk Minimization (ERM) (standard training without debiasing techniques, (b) HypothesisOnly-PoE (Karimi Mahabadi et al.,

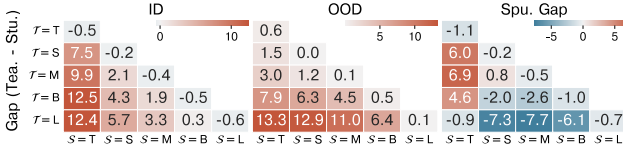


Figure 2. **C1: Teacher vs. Student:** average performance gaps between teacher and student models on ID, OOD, and Spurious Gap across text datasets. X-axis and Y-axis show the scale of student ($\mathcal{S}$) and teacher ($\mathcal{T}$) respectively: tiny (T), small (S), medium (M), big (B), and large (L). Each cell shows the performance gap between corresponding scales of a teacher and a student. See Appendix D for detailed results.

2020), (c) WeakLearner-PoE (Sanh et al., 2021), (d) KernelWhitening (Gao et al., 2022), (e) AttentionPoE (Wang et al., 2023), (f) σ -Damp (Puli et al., 2023), (g) DeepFeatReweight (Kirichenko et al., 2023), and (h) PerSampleGrad (Ahn et al., 2023). The above debiasing methods have a wide coverage of existing algorithms, ranging from auxiliary biased model-based debiasing, to disentanglement of representations. Meanwhile, they can handle multiple types of shortcuts at the same time, without overfitting to a specific bias. Details of these methods are provided in Appendix A.

4. RQ1: Distillability of Debiasing Methods

We first examine if KD can effectively distill the debiasing capability from teachers to students of different scales.

Students become more biased than teachers We observe that teachers consistently achieve better performance than their smaller scale students on ID and OOD test sets after knowledge distillation. The positive values in Figure 2 show that although KD encourages students to mimic their teachers in the logit space, it may undesirably increase student’s susceptibility to spurious correlations in datasets as the teacher’s in-domain and debiasing capabilities are not effectively transferred to the student. The prediction agreement between teacher and student models show similar trend, where the student generally aligns with the teacher on ID but often largely diverges from the teacher on OOD. Furthermore, the extent of knowledge loss after distillation varies depending on the relative scales of the teacher and student models. For example, as depicted in Figure 2, when $\mathcal{S}$ is tiny ($\mathcal{S} = \mathcal{T}$), more debiasing power is lost, shown by mostly positive values in spurious gap. When $\mathcal{T}$ is large ($\mathcal{T} = \mathcal{L}$), more ID knowledge is lost, shown by mostly negative values in spurious gap. The results show that if a teacher model learns a partially debiased representation but still retains residual biases, the student might amplify this bias rather than mitigate it.

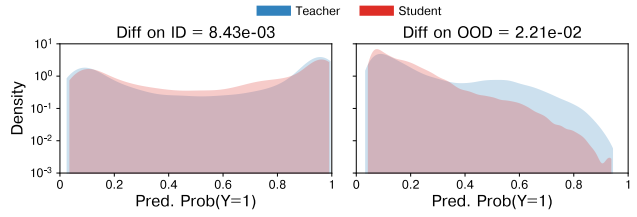


Figure 3. **C1: Teacher vs. Student:** density of predicted probability on text datasets. On OOD, students has larger deviation in prediction confidence than teachers. See Appendix D for detailed results.

Students show diverse distribution shifts in predicted probabilities To understand the influence of KD on the debiasing capabilities of students, we investigate the output probability distribution $P_C(y = 1)$. Our findings show that KD significantly alter the predicted probability distribution, despite its training objective of matching output logits. This perturbation is often larger on OOD than ID test sets, which explains the larger performance drop observed in students compared to their teachers on OOD, as illustrated in Figure 3. We also observed that teachers tend to provide slightly more confident predictions on ID while more moderate predictions on OOD. Such behavior is not successfully transferred to students through KD. Such distinct behaviors on different samples may encourage models to overfit to data distributions of the training sets or to over-represented groups, which can effectively amplify reliance on shortcuts over robust features. In addition, the training sets of teachers often contain biased examples or do not equally represent all sub-groups, which leads students to inherit and potentially amplify these biases. Consequently, students often perform worse than their teachers on OOD.

Students and teachers show different attention on ID and OOD To have a deeper understanding of the teacher-student divergence, we further probe the internal representations when making predictions on ID and OOD data. Results shows that after distillation, students try to mimic teachers on ID (left). The earlier layers of students follow earlier layers of teachers, and similarly mid and later layers. This indicates that KD can transfer knowledge of ID data from the larger teachers into smaller students. On OOD (right), however, we observe similar pattern but it is not fully preserved. In particular, it is challenging for the mid and later layers of the students to follow closely to those of the teachers, which explains the performance degradation on OOD after KD, see Figure 5.

Potential for new debiasing capabilities for students beyond teacher abilities We compare prediction agreement between teacher and student models. When $\mathcal{T}$ is large ($\mathcal{T}$

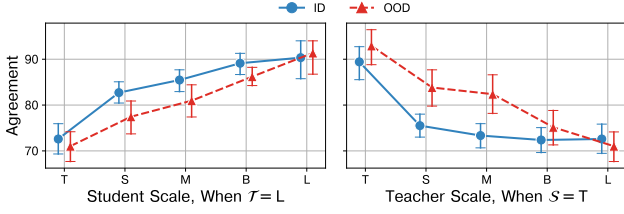


Figure 4. C1: **Teacher vs. Student:** prediction agreement on text datasets. Left: varying S (X-axis) given a fixed large teacher ($T = L$). Right: varying T (X-axis) given a fixed tiny student ($S = T$). Agreement increases as the scale of teacher and student get closer. See Appendix D for detailed results.

$= L$), we observe an increase in prediction agreement as S scales up, with consistently higher agreement on ID than OOD, as shown in the left plot in Figure 4. Conversely, when S is tiny ($S = T$), the prediction agreement diminishes as T scales up, with higher agreement on OOD than ID, the right plot in Figure 4. The imperfect agreement between teacher and student contradicts with the foundational assumptions of knowledge distillation, which assumes that students should closely mimic their teachers. However, interestingly, this unexpected behavior may not always lead to performance degradation. Sometimes it enables students to generalize to out of domain data. In particular, there are instances where students make correct predictions where their teachers do not, see the left plot in Figure 10. Students can sometimes outperform their teachers perhaps because they may learn additional patterns during the knowledge distillation process, which allows them to generalize better than their teachers. The above result suggests that students may sometimes acquire debiasing capabilities that surpass those of their teachers, which we believe is a novel avenue for robust model training.

Larger teachers do not guarantee more robust students

Our findings show that a more capable teacher does not guarantee a less biased student in debiasing tasks. With a fixed student scale (as seen in the columns of Figure 2), increasing the teacher’s scale does not consistently reduce performance gap or spurious gap. Sometimes, a larger teacher may degrade the debiasing capability of the student. For example, when $S = T$, increasing the teacher scale from M to B increases the spurious gap from 6.5 to 8.1 on ERM, i.e. a more biased model. Moreover, when $S = T$, increasing T result in a drop of teacher-student agreement, indicating that the student fails to follow the teacher, see right plot in Figure 4. We attribute this finding to two reasons. Firstly, the capabilities of students are substantially bounded by their scale, and using a much larger and capable teacher may exceed the student’s capacity for effective learning (Cho & Hariharan, 2019). Secondly, training students with debias-

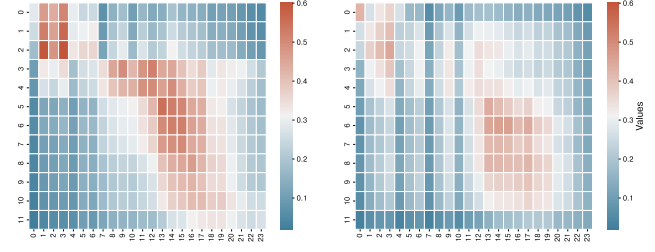


Figure 5. C1: **Teacher vs. Student:** Centered Kernel Alignment on ID (left) and OOD (right). Higher values indicate higher similarity. X-axis and Y-axis refer to the layers of teacher (T) and student (S) respectively. See Appendix D for detailed results.

ing objectives and knowledge distillation at the same time results in optimization problem, which may trap students’ parameters in local optima and affect their robustness to spurious correlations.

Students with similar scales to their teachers learn better

The effectiveness of debiasing ability transfer through distillation is greatly affected by the scale similarity between teacher and student. As the teacher and the student become similar in scale (near the diagonal cells in Figure 2), the differences on test set performance and spurious gaps decrease. However, a larger mismatch in scale (far from diagonal) results in more pronounced differences, see Figure 2. Similarly, the teacher-student agreement increases as T and S align more closely, see Figure 4. This is likely because models of similar scales have comparable expressive power and extracts similar features, which can lead to more effective knowledge transfer, better bias mitigation, and higher prediction agreement.

5. RQ2: Distillation vs. Standard Training

We assess if training with knowledge distillation (KD) can improve a model’s debiasing performance compared to standard training (Non-KD).

Non-KD is less biased than KD Our results show that debiasing models trained from scratch (Non-KD) have lower ID performance than those trained with KD. However, the Non-KD models achieve almost no changes on OOD, leading to smaller spurious gaps, see Figure 6. We hypothesize that the distillation objective of matching logits, despite effective on ID, may potentially inject additional spurious correlations and distract the model from prioritizing robust features, as the teacher is not fully unbiased.

KD does not improve generalization An interesting finding is that both Non-KD and KD have similar average prediction agreements on both ID and OOD. However,

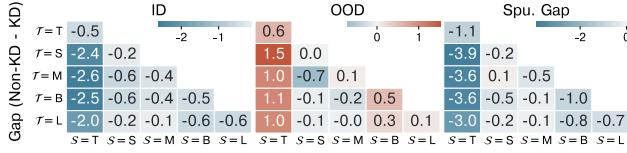


Figure 6. C2: **Non-KD vs. KD**: average performance gap between teacher and student on ID, OOD, and Spurious Gap across text datasets. X-axis and Y-axis show the scale of student ($\mathcal{S}$) and teacher ($\mathcal{T}$) respectively: : tiny (T), small (S), medium (M), big (B), and large (L). Each cell shows the performance gap of the corresponding scales of a teacher and a student. See Appendix D for detailed results.

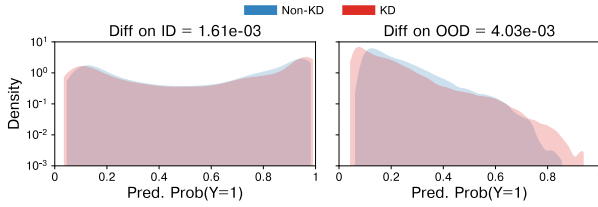


Figure 7. C2: **Non-KD vs. KD**: density of predicted probability on text datasets. On OOD, KD has larger deviation in prediction confidence than Non-KD. See Appendix D for detailed results.

the agreement on OOD varies significantly depending on dataset, debiasing method, and backbone model. This suggests that training solely with the original data (Non-KD) is sufficiently effective for debiasing, and introducing external knowledge via KD does not yield significant improvements. This result can be attributed to KD’s impact on model confidence; models trained with KD tend to produce more confident predictions than models trained without KD, see Figure 7, which is key to degenerate performance on OOD (Utama et al., 2020; Sanh et al., 2021). Such overconfidence could be a critical factor in degraded performance on OOD tasks. Moreover, such minimal contribution of KD remains unchanged even when stronger external knowledge (a larger teacher) or a more capable learner (a larger student) is used, see Figure 8.

6. RQ3: Effect of Debiasing Methods

We assess the effect of different debiasing methods and backbones on our earlier findings.

Different transfer patterns across methods Our results show that the transfer patterns are heavily influenced by the formulation of debiasing method. For example, logit-based PoE methods, such as HypothesisOnly-PoE and WeakLearner-PoE, show similar trends in performance

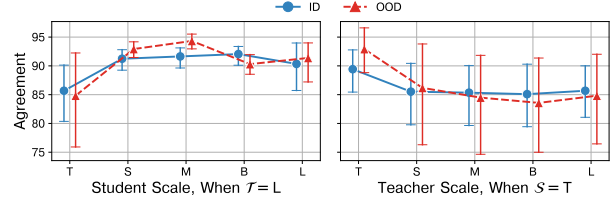


Figure 8. C2: **Non-KD vs. KD**: prediction agreement on text datasets. Left: varying $\mathcal{S}$ (X-axis) given a fixed teacher with $\mathcal{T} = \mathcal{L}$. Right: varying $\mathcal{T}$ (X-axis) given a fixed student with $\mathcal{T} = \mathcal{T}$. Agreement does not increase significantly as $\mathcal{T}$ and $\mathcal{S}$ scale up. See Appendix D for detailed results.

changes and spurious gaps, in contrast to the representation disentanglement method (KernelWhitening), see Figure 9.

Sensitivity to backbones The distillability of KD appears to varies with the architecture of the backbones and randomness in the training. KernelWhitening and WeakLearner-PoE are two methods particularly sensitive to the scale of backbone and random seeds, which controls factors such as data sampling and ordering.

Robustness to different biases transfer differently We observe that OOD and Transfer show different transfer patterns, where performance gap on Transfer exhibit much larger variations the student scales up, see Figure 9. This suggests that smaller students may outperform larger ones on OOD, indicating that during KD, larger students may become more prone to certain biases (OOD) but more resilient to others.

Universal transfer patterns in debiasing methods A number of debiasing methods show consistent changes in robustness after KD, which suggest the potential for an empirical universal transfer pattern. Specifically on text datasets, the performance gap between teacher and student models on OOD and Transfer Spurious Gap fall in the range of $[0, 5]$ and $[-5, 0]$ respectively, see Figure 9. Such change in performance is consistent across different scales of $\mathcal{T}$ and $\mathcal{S}$, which allows for predictable performance after KD. Similarly, the performance gap between models trained using Non-KD and KD remains stable on OOD, falling in range $[-1, 1]$ across different scales.

7. Potential Solutions

Based on the above analyses, we summarize the key findings on distilling debiasing models as follows:

- Training distribution significantly affect successful distillation of debiasing capabilities.

- Student models with similar scale to their teachers can better obtain debiasing knowledge from their teachers.
- The objectives of KD may introduce additional optimization challenges, especially with the presence of debiasing objectives.

To further improve the distillability of debiasing methods, we propose three solutions:

Data augmentation (DA) There is broad evidence that models become biased by relying on spurious features in the training set (Wu et al., 2022; Ahn et al., 2023), which is amplified by misrepresentation of specific classes or labeling errors. Prior studies have highlighted the important role of data in knowledge distillation (Stanton et al., 2021). Based on these prior studies and our findings, we hypothesize that providing high quality data and augmenting data size can improve the process of distilling the debiasing capability from teacher to student. For text datasets, we employ the data generated by Seq-Z filtering (Wu et al., 2022) as training set for both teacher and student models. For image datasets, we employ training and validation sets where the sub-groups are equally represented (Kirichenko et al., 2023).

Iterative knowledge distillation (IKD) Our results indicate that the transfer of debiasing capabilities is more effective between teachers and students of similar scales. Therefore, we propose to leverage Iterative Knowledge Distillation (IKD) (Liu et al., 2023): given a teacher of scale $\mathcal{S}_N$, we first distill it to a student of scale $\mathcal{S}_{N-1}$, where $\mathcal{S}_{N-1}$ is the closest neighbor of $\mathcal{S}_N$ in scale. Then the newly-distilled student acts as a teacher and transfer the knowledge to a model of smaller scale $\mathcal{S}_{N-2}$, where $\mathcal{S}_{N-2}$ is the closest neighbor of $\mathcal{S}_{N-1}$ in scale. We repeat this process iteratively by gradually decreasing student scale, such that the knowledge can be transferred smoothly from a large scale model to a small scale model. This step-wise approach enables a smooth knowledge transfer from larger to smaller scale models, and potentially improves debiasing effectiveness at each step.

Initialize student with teacher weights (Init) Previous research by Stanton et al. (2021) has discovered that initializing a student model with the weights of its teacher can increase their centered kernel alignment (Kornblith et al., 2019) in activation space. This approach can head-start the student model with a stronger debiasing capability from the teacher. It can also help alleviate potential optimization obstacles and stuck in local optima. If the teacher and student models are of the same scale, we initialize the student with the teacher parameters. If the teacher is larger, we initialize the student with the first few layers of the teacher.

Table 2. Improvement of distillability. Vanilla refers to standard knowledge distillation, DA, IKD, and Init represent data augmentation, iterative knowledge distillation, and initialization of student with teacher weights (Init) as our three solutions to improve the distillability of debiasing methods.

DIFF IN	ID (↓)	OOD (↓)	SPU. GAP (↓)	ID (↓)	OOD (↓)	SPU. GAP (↓)
	TEACHER - STUDENT			NON KD - KD		
VANILLA	5.1	7.3	12.7	1.4	0.7	2.2
+ DA	2.3	5.4	8.2	0.2	0.2	0.5
+ IKD	3.6	5.9	9.9	1.0	0.5	1.6
+ INIT.	4.7	6.5	11.5	1.3	0.7	2.0

Results Table 2 shows that all three solutions result in improved distillability. Specifically, on spurious gap between teacher and student, data augmentation (DA), iterative knowledge distillation (IKD), and initialization with teacher weights (Init) yield performance gains of 4.5, 2.8, 1.2 absolute points compared to vanilla KD across datasets and backbones respectively. On spurious gap between Non-KD and KD, DA, IKD, and Init outperforms vanilla KD by 1.7, 0.6 and 0.2 absolute points respectively. We find that DA achieves the largest improvement, since the root cause of spurious correlations come from the underlying dataset (Chen et al., 2018). Debiasing the dataset itself can benefit all training methods including knowledge distillation. As noted by previous work (Stanton et al., 2021), Init may facilitate teacher-student agreement in activation space, but result in non-significant gains, which aligns with our findings as well.

8. Related Work

Bias mitigation in NLU: Debiasing approaches usually employ a biased model to inform the training of a robust model (Clark et al., 2019; Karimi Mahabadi et al., 2020; Sanh et al., 2021; Utama et al., 2020; Cheng et al., 2024). Other methods aim at learning debiased or robust representations (Gao et al., 2022; Lyu et al., 2022; Wang et al., 2023; Jeon et al., 2023; Reif & Schwartz, 2023), or removing bias-encoding parameters (Meissner et al., 2022; Yu et al., 2023). Other works include measurement of bias of specific words with statistical test (Gardner et al., 2021), generating non-biased samples (Wu et al., 2022), identification of bias-encoding parameters (Yu et al., 2023), when bias mitigation works (Ravichander et al., 2023), and bias transfer from other models (Jin et al., 2021).

Bias mitigation in vision: In vision, worst-group performance is measured as a sign of model robustness. Several works investigate how to learn debiased models from failure cases (Nam et al., 2020), biased representations (Bahng

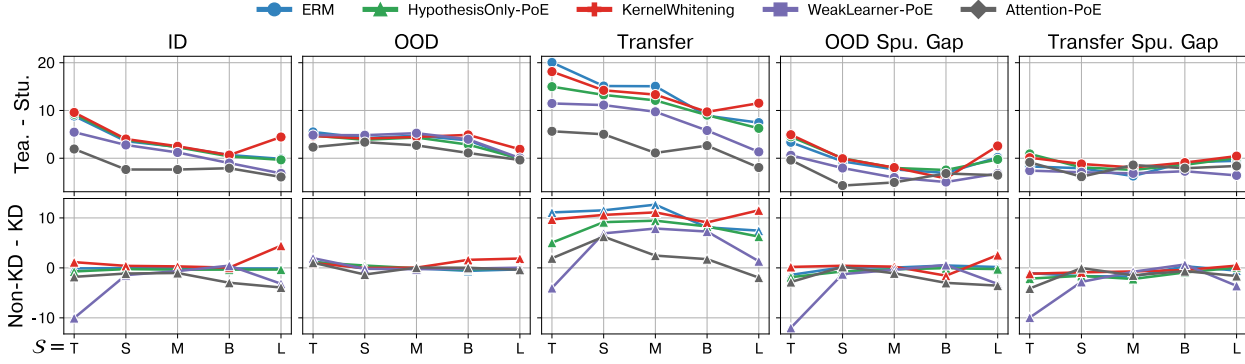


Figure 9. C3: Comparison Between Debiasing Methods: performance gap between Teacher and Student (Above), Non-KD and KD (Lower) on text datasets. Detailed results are shown in Appendix.

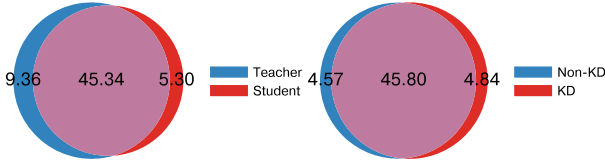


Figure 10. C1: Teacher vs. Student (Left) and C2: Non-KD vs. KD (Right): correctly predicted examples on OOD on text datasets.

et al., 2020), multiple biased models (Kim et al., 2022), and by simply re-training the last layer of a neural model (i.e. the classification layer) with additional equally represented data (Kirichenko et al., 2023; LaBonte et al., 2023). Li et al. (2023) showed that multiple spurious features can occur in a dataset, while suppressing one may inevitably boost another one. Other perspectives for debiasing include causal attention (Wang et al., 2021), building uniform margin classifiers (Puli et al., 2023), using representations from earlier layers (Tiwari et al., 2024), and neural collapse (Wang et al., 2024), where feature space collapses into a stable geometric structure that results in robustness and generalizability.

Knowledge distillation: Knowledge Distillation (KD) is initially proposed to transfer knowledge from a larger model (teacher) to a smaller model (student), by encouraging the student to follow the teacher on prediction logits (Hinton et al., 2015), learned features (Romero et al., 2015; Wang et al., 2020), attention map (Zagoruyko & Komodakis, 2017; Chen et al., 2021), activation patterns (Huang & Wang, 2017; Heo et al., 2019). Later works discovered that KD can be viewed as a special form of regularization similar to label smoothing (Szegedy et al., 2016), providing no task-specific knowledge. However, on text classification tasks, whether KD can regularize the student depends on the choice of teacher model (Sultan, 2023), which may result in opposite model confidence between teacher and student compared to label smoothing. Stanton et al. (2021) discovers that opti-

mization and dataset details are crucial to matching students to teachers, and such matching does not guarantee better generalization ability of students. Xue et al. (2023) investigates cross-modal KD, where the teacher functions on a different modality or extra modalities than student. The authors proposes modality fusing hypothesis, which claims that modality decisive features are critical for the effectiveness of cross-modal KD. However, despite briefly discussed (Cho & Hariharan, 2019; Tiwari et al., 2024), the potential of knowledge distillation to transfer debiasing capabilities across different modalities and backbone models remains underexplored and poorly understood in existing work.

9. Conclusion

We present the first study on the distillability of debiasing capabilities between neural models, and the extent of bias transfer through knowledge distillation (KD). We evaluate eight popular debiasing methods and five scales of backbones on four datasets. Extensive experiments show that vanilla KD does not consistently preserve debiasing capabilities; in many cases, student models become more reliant on spurious correlations than their teachers; the effectiveness of debiasing transfer depends on model scale similarity—distillation works best when teacher and student models are comparable in complexity; and larger teachers do not always yield more robust students, which indicates the need for targeted debiasing strategies in KD. We propose three solutions—data augmentation, iterative KD, and student initialization—which significantly improve the distillability of debiasing methods and contribution of KD on debiasing.

In future we will investigate self-distilled debiasing, where the student iteratively distills knowledge from itself rather than relying on a fixed teacher. A potential improvement is to explicitly guide the student using counterfactual data augmentation during distillation.

Impact Statement

Our research focuses on mitigating dataset biases in text and vision datasets, and understanding why debiasing methods may fail under knowledge distillation. The broader impacts of our work are in advancing dataset fairness and potentially enhancing decision-making based on data. Our work contributes to improving the accuracy and reliability of NLP and vision models, as well as their trust and adoption.

References

- Ahn, S., Kim, S., and Yun, S.-Y. Mitigating dataset bias by using per-sample gradient. In *The Eleventh International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=7mgUec-7GMv>.
- Bahng, H., Chun, S., Yun, S., Choo, J., and Oh, S. J. Learning de-biased representations with biased representations. In III, H. D. and Singh, A. (eds.), *Proceedings of the 37th International Conference on Machine Learning*, volume 119 of *Proceedings of Machine Learning Research*, pp. 528–539. PMLR, 13–18 Jul 2020. URL <https://proceedings.mlr.press/v119/bahng20a.html>.
- Bowman, S. R., Angeli, G., Potts, C., and Manning, C. D. A large annotated corpus for learning natural language inference. In Màrquez, L., Callison-Burch, C., and Su, J. (eds.), *Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing*, pp. 632–642, Lisbon, Portugal, September 2015. Association for Computational Linguistics. doi: 10.18653/v1/D15-1075. URL <https://aclanthology.org/D15-1075/>.
- Chen, D., Mei, J.-P., Zhang, Y., Wang, C., Wang, Z., Feng, Y., and Chen, C. Cross-layer distillation with semantic calibration. *Proceedings of the AAAI Conference on Artificial Intelligence*, 35(8):7028–7036, May 2021. doi: 10.1609/aaai.v35i8.16865. URL <https://ojs.aaai.org/index.php/AAAI/article/view/16865>.
- Chen, I., Johansson, F. D., and Sontag, D. Why is my classifier discriminatory? *Advances in neural information processing systems*, 31, 2018.
- Cheng, J. and Amiri, H. FairFlow: Mitigating dataset biases through undecided learning for natural language understanding. In Al-Onaizan, Y., Bansal, M., and Chen, Y.-N. (eds.), *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing*, pp. 21960–21975, Miami, Florida, USA, November 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.emnlp-main.1225. URL <https://aclanthology.org/2024.emnlp-main.1225/>.
- Cheng, J., Elgaar, M., Vakil, N., and Amiri, H. Cognitive: Multimodal and multilingual fusion networks for mild cognitive impairment assessment from spontaneous speech. *arXiv preprint arXiv:2407.13660*, 2024.
- Chew, O., Lin, H.-T., Chang, K.-W., and Huang, K.-H. Understanding and mitigating spurious correlations in text classification with neighborhood analysis. In Graham, Y. and Purver, M. (eds.), *Findings of the Association for Computational Linguistics: EACL 2024*, pp. 1013–1025, St. Julian’s, Malta, March 2024. Association for Computational Linguistics. URL <https://aclanthology.org/2024.findings-eacl.68/>.
- Cho, J. H. and Hariharan, B. On the efficacy of knowledge distillation. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, October 2019.
- Clark, C., Yatskar, M., and Zettlemoyer, L. Don’t take the easy way out: Ensemble based methods for avoiding known dataset biases. In Inui, K., Jiang, J., Ng, V., and Wan, X. (eds.), *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pp. 4069–4082, Hong Kong, China, November 2019. Association for Computational Linguistics. doi: 10.18653/v1/D19-1418. URL <https://aclanthology.org/D19-1418/>.
- Cortes, C., Mohri, M., and Rostamizadeh, A. Algorithms for learning kernels based on centered alignment. *Journal of Machine Learning Research*, 13(28): 795–828, 2012. URL <http://jmlr.org/papers/v13/cortes12a.html>.
- Devlin, J., Chang, M.-W., Lee, K., and Toutanova, K. BERT: Pre-training of deep bidirectional transformers for language understanding. In Burstein, J., Doran, C., and Solorio, T. (eds.), *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pp. 4171–4186, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics. doi: 10.18653/v1/N19-1423. URL <https://aclanthology.org/N19-1423/>.
- Dolan, W. B. and Brockett, C. Automatically constructing a corpus of sentential paraphrases. In *Proceedings of the Third International Workshop on Paraphrasing (IWP2005)*, 2005. URL <https://aclanthology.org/I05-5002/>.
- Dosovitskiy, A., Beyer, L., Kolesnikov, A., Weissenborn, D., Zhai, X., Unterthiner, T., Dehghani, M., Minderer, M., Heigold, G., Gelly, S., Uszkoreit, J., and Houlsby,

- N. An image is worth 16x16 words: Transformers for image recognition at scale. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=YicbFdNTTy>.
- Gao, S., Dou, S., Zhang, Q., and Huang, X. Kernel-whitening: Overcome dataset bias with isotropic sentence embedding. In Goldberg, Y., Kozareva, Z., and Zhang, Y. (eds.), *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pp. 4112–4122, Abu Dhabi, United Arab Emirates, December 2022. Association for Computational Linguistics. doi: 10.18653/v1/2022.emnlp-main.275. URL <https://aclanthology.org/2022.emnlp-main.275/>.
- Gardner, M., Merrill, W., Dodge, J., Peters, M., Ross, A., Singh, S., and Smith, N. A. Competency problems: On finding and removing artifacts in language data. In Moens, M.-F., Huang, X., Specia, L., and Yih, S. W.-t. (eds.), *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pp. 1801–1813, Online and Punta Cana, Dominican Republic, November 2021. Association for Computational Linguistics. doi: 10.18653/v1/2021.emnlp-main.135. URL <https://aclanthology.org/2021.emnlp-main.135/>.
- Guo, Q., Tang, Y., Ouyang, Y., Wu, Z., and Dai, X. Debias NLU datasets via training-free perturbations. In Bouamor, H., Pino, J., and Bali, K. (eds.), *Findings of the Association for Computational Linguistics: EMNLP 2023*, pp. 10886–10901, Singapore, December 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.findings-emnlp.726. URL <https://aclanthology.org/2023.findings-emnlp.726/>.
- He, K., Zhang, X., Ren, S., and Sun, J. Deep residual learning for image recognition. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2016.
- Heo, B., Lee, M., Yun, S., and Choi, J. Y. Knowledge transfer via distillation of activation boundaries formed by hidden neurons. *Proceedings of the AAAI Conference on Artificial Intelligence*, 33(01): 3779–3787, Jul. 2019. doi: 10.1609/aaai.v33i01.33013779. URL <https://ojs.aaai.org/index.php/AAAI/article/view/4264>.
- Hinton, G., Vinyals, O., and Dean, J. Distilling the knowledge in a neural network. *arXiv preprint arXiv:1503.02531*, 2015.
- Hinton, G. E. Training products of experts by minimizing contrastive divergence. *Neural Comput.*, 14(8): 1771–1800, aug 2002. ISSN 0899-7667. doi: 10.1162/089976602760128018. URL <https://doi.org/10.1162/089976602760128018>.
- Huang, Z. and Wang, N. Like what you like: Knowledge distill via neuron selectivity transfer. *arXiv preprint arXiv:1707.01219*, 2017.
- Jeon, E., Lee, M., Park, J., Kim, Y., Mok, W.-L., and Lee, S. Improving bias mitigation through bias experts in natural language understanding. In Bouamor, H., Pino, J., and Bali, K. (eds.), *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pp. 11053–11066, Singapore, December 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.emnlp-main.681. URL <https://aclanthology.org/2023.emnlp-main.681/>.
- Jin, X., Barbieri, F., Kennedy, B., Mostafazadeh Davani, A., Neves, L., and Ren, X. On transferability of bias mitigation effects in language model fine-tuning. In Toutanova, K., Rumshisky, A., Zettlemoyer, L., Hakkani-Tur, D., Beltagy, I., Bethard, S., Cotterell, R., Chakraborty, T., and Zhou, Y. (eds.), *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pp. 3770–3783, Online, June 2021. Association for Computational Linguistics. doi: 10.18653/v1/2021.naacl-main.296. URL <https://aclanthology.org/2021.naacl-main.296/>.
- Karimi Mahabadi, R., Belinkov, Y., and Henderson, J. End-to-end bias mitigation by modelling biases in corpora. In Jurafsky, D., Chai, J., Schluter, N., and Tetreault, J. (eds.), *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pp. 8706–8716, Online, July 2020. Association for Computational Linguistics. doi: 10.18653/v1/2020.acl-main.769. URL <https://aclanthology.org/2020.acl-main.769/>.
- Kim, N., HWANG, S., Ahn, S., Park, J., and Kwak, S. Learning debiased classifier with biased committee. In Koyejo, S., Mohamed, S., Agarwal, A., Belgrave, D., Cho, K., and Oh, A. (eds.), *Advances in Neural Information Processing Systems*, volume 35, pp. 18403–18415. Curran Associates, Inc., 2022. URL https://proceedings.neurips.cc/paper_files/paper/2022/file/750046157471c56235a781f2eff6e226-Paper-Conference.pdf.
- Kirichenko, P., Izmailov, P., and Wilson, A. G. Last layer re-training is sufficient for robustness to spurious correlations. In *The Eleventh International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=Zb6c8A-Fghk>.

- Kornblith, S., Norouzi, M., Lee, H., and Hinton, G. Similarity of neural network representations revisited. In Chaudhuri, K. and Salakhutdinov, R. (eds.), *Proceedings of the 36th International Conference on Machine Learning*, volume 97 of *Proceedings of Machine Learning Research*, pp. 3519–3529. PMLR, 09–15 Jun 2019. URL <https://proceedings.mlr.press/v97/kornblith19a.html>.
- LaBonte, T., Muthukumar, V., and Kumar, A. Towards last-layer retraining for group robustness with fewer annotations. In Oh, A., Naumann, T., Globerson, A., Saenko, K., Hardt, M., and Levine, S. (eds.), *Advances in Neural Information Processing Systems*, volume 36, pp. 11552–11579. Curran Associates, Inc., 2023. URL https://proceedings.neurips.cc/paper_files/paper/2023/file/265bee74aee86df77e8e36d25e786ab5-Paper-Conference-main.pdf.
- Li, Z., Evtimov, I., Gordo, A., Hazirbas, C., Hassner, T., Ferrer, C. C., Xu, C., and Ibrahim, M. A whac-a-mole dilemma: Shortcuts come in multiples where mitigating one amplifies others. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 20071–20082, June 2023.
- Liu, J., Wang, P., Shang, Z., and Wu, C. Iterde: An iterative knowledge distillation framework for knowledge graph embeddings. *Proceedings of the AAAI Conference on Artificial Intelligence*, 37:4488–4496, Jun. 2023. doi: 10.1609/aaai.v37i4.25570. URL <https://ojs.aaai.org/index.php/AAAI/article/view/25570>.
- Liu, Z., Luo, P., Wang, X., and Tang, X. Deep learning face attributes in the wild. In *Proceedings of the IEEE international conference on computer vision*, pp. 3730–3738, 2015.
- Lyu, Y., Li, P., Yang, Y., de Rijke, M., Ren, P., Zhao, Y., Yin, D., and Ren, Z. Feature-level debiased natural language understanding. In *Proceedings of the AAAI Conference on Artificial Intelligence*, 2022.
- McCoy, T., Pavlick, E., and Linzen, T. Right for the wrong reasons: Diagnosing syntactic heuristics in natural language inference. In Korhonen, A., Traum, D., and Màrquez, L. (eds.), *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pp. 3428–3448, Florence, Italy, July 2019. Association for Computational Linguistics. doi: 10.18653/v1/P19-1334. URL <https://aclanthology.org/P19-1334/>.
- Meissner, J. M., Sugawara, S., and Aizawa, A. De-biasing masks: A new framework for shortcut mitigation in NLU. In Goldberg, Y., Kozareva, Z., and Zhang, Y. (eds.), *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pp. 7607–7613, Abu Dhabi, United Arab Emirates, December 2022. Association for Computational Linguistics. doi: 10.18653/v1/2022.emnlp-main.517. URL <https://aclanthology.org/2022.emnlp-main.517/>.
- Mendelson, M. and Belinkov, Y. Debiasing methods in natural language understanding make bias more accessible. In Moens, M.-F., Huang, X., Specia, L., and Yih, S. W.-t. (eds.), *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pp. 1545–1557, Online and Punta Cana, Dominican Republic, November 2021. Association for Computational Linguistics. doi: 10.18653/v1/2021.emnlp-main.116. URL <https://aclanthology.org/2021.emnlp-main.116/>.
- Naik, A., Ravichander, A., Sadeh, N., Rose, C., and Neubig, G. Stress test evaluation for natural language inference. In Bender, E. M., Derczynski, L., and Isabelle, P. (eds.), *Proceedings of the 27th International Conference on Computational Linguistics*, pp. 2340–2353, Santa Fe, New Mexico, USA, August 2018. Association for Computational Linguistics. URL <https://aclanthology.org/C18-1198/>.
- Nam, J., Cha, H., Ahn, S., Lee, J., and Shin, J. Learning from failure: De-biasing classifier from biased classifier. In Larochelle, H., Ranzato, M., Hadsell, R., Balcan, M., and Lin, H. (eds.), *Advances in Neural Information Processing Systems*, volume 33, pp. 20673–20684. Curran Associates, Inc., 2020. URL https://proceedings.neurips.cc/paper_files/paper/2020/file/eddc3427c5d77843c2253f1e799fe933-Paper-Conference-main.pdf.
- Nguyen, T., Raghu, M., and Kornblith, S. Do wide and deep networks learn the same things? uncovering how neural network representations vary with width and depth. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=KJNCaKY8tY4>.
- Puli, A. M., Zhang, L., Wald, Y., and Ranganath, R. Don’t blame dataset shift! shortcut learning due to gradients and cross entropy. In Oh, A., Naumann, T., Globerson, A., Saenko, K., Hardt, M., and Levine, S. (eds.), *Advances in Neural Information Processing Systems*, volume 36, pp. 71874–71910. Curran Associates, Inc., 2023. URL https://proceedings.neurips.cc/paper_files/paper/2023/file/e35460304fdf6df523f068a59aaf8829-Paper-Conference-main.pdf.

- Raghu, M., Unterthiner, T., Kornblith, S., Zhang, C., and Dosovitskiy, A. Do vision transformers see like convolutional neural networks? In Ranzato, M., Beygelzimer, A., Dauphin, Y., Liang, P., and Vaughan, J. W. (eds.), *Advances in Neural Information Processing Systems*, volume 34, pp. 12116–12128. Curran Associates, Inc., 2021. URL https://proceedings.neurips.cc/paper_files/paper/2021/file/652cf38361a209088302ba2b8b7f51e0-Paper.pdf.
- Ravichander, A., Stacey, J., and Rei, M. When and why does bias mitigation work? In Bouamor, H., Pino, J., and Bali, K. (eds.), *Findings of the Association for Computational Linguistics: EMNLP 2023*, pp. 9233–9247, Singapore, December 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.findings-emnlp.619. URL <https://aclanthology.org/2023.findings-emnlp.619/>.
- Reif, Y. and Schwartz, R. Fighting bias with bias: Promoting model robustness by amplifying dataset biases. In Rogers, A., Boyd-Graber, J., and Okazaki, N. (eds.), *Findings of the Association for Computational Linguistics: ACL 2023*, pp. 13169–13189, Toronto, Canada, July 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.findings-acl.833. URL <https://aclanthology.org/2023.findings-acl.833/>.
- Romero, A., Ballas, N., Kahou, S. E., Chassang, A., Gatta, C., and Bengio, Y. Fitnets: Hints for thin deep nets. In *International Conference on Learning Representations*, 2015.
- Sagawa, S., Koh, P. W., Hashimoto, T. B., and Liang, P. Distributionally robust neural networks. In *International Conference on Learning Representations*, 2020. URL <https://openreview.net/forum?id=ryxGuJrFvS>.
- Sanh, V., Wolf, T., Belinkov, Y., and Rush, A. M. Learning from others’ mistakes: Avoiding dataset biases without modeling them. In *International Conference on Learning Representations*, 2021. URL <https://openreview.net/forum?id=Hf3qXoiNkR>.
- Sharma, L., Graesser, L., Nangia, N., and Evci, U. Natural language understanding with the quora question pairs dataset. *arXiv e-prints*, 2019. URL <https://arxiv.org/abs/1907.01041>.
- Stanton, S. D., Izmailov, P., Kirichenko, P., Alemi, A. A., and Wilson, A. G. Does knowledge distillation really work? In Beygelzimer, A., Dauphin, Y., Liang, P., and Vaughan, J. W. (eds.), *Advances in Neural Information Processing Systems*, 2021. URL <https://openreview.net/forum?id=7J-fKoXiReA>.
- Sultan, M. Knowledge distillation $\approx$ label smoothing: Fact or fallacy? In Bouamor, H., Pino, J., and Bali, K. (eds.), *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pp. 4469–4477, Singapore, December 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.emnlp-main.271. URL <https://aclanthology.org/2023.emnlp-main.271/>.
- Szegedy, C., Vanhoucke, V., Ioffe, S., Shlens, J., and Wojna, Z. Rethinking the inception architecture for computer vision. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2016.
- Tay, Y., Dehghani, M., Rao, J., Fedus, W., Abnar, S., Chung, H. W., Narang, S., Yogatama, D., Vaswani, A., and Metzler, D. Scale efficiently: Insights from pretraining and finetuning transformers. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=f2OYVDyfiB>.
- Tiwari, R., Sivasubramanian, D., Mekala, A., Ramakrishnan, G., and Shenoy, P. Using early readouts to mediate featural bias in distillation. In *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)*, pp. 2638–2647, January 2024.
- Turc, I., Chang, M.-W., Lee, K., and Toutanova, K. Well-read students learn better: On the importance of pre-training compact models, 2019.
- Utama, P. A., Moosavi, N. S., and Gurevych, I. Towards debiasing NLU models from unknown biases. In Webber, B., Cohn, T., He, Y., and Liu, Y. (eds.), *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pp. 7597–7610, Online, November 2020. Association for Computational Linguistics. doi: 10.18653/v1/2020.emnlp-main.613. URL <https://aclanthology.org/2020.emnlp-main.613/>.
- Wah, C., Branson, S., Welinder, P., Perona, P., and Belongie, S. Caltech-ucsd birds-200-2011 (cub-200-2011). Technical report, California Institute of Technology, 2011.
- Wang, F., Huang, J. Y., Yan, T., Zhou, W., and Chen, M. Robust natural language understanding with residual attention debiasing. In Rogers, A., Boyd-Graber, J., and Okazaki, N. (eds.), *Findings of the Association for Computational Linguistics: ACL 2023*, pp. 504–519, Toronto, Canada, July 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.findings-acl.32. URL <https://aclanthology.org/2023.findings-acl.32/>.

- Wang, T., Zhou, C., Sun, Q., and Zhang, H. Causal attention for unbiased visual recognition. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, pp. 3091–3100, October 2021.
- Wang, X., Fu, T., Liao, S., Wang, S., Lei, Z., and Mei, T. Exclusivity-consistency regularized knowledge distillation for face recognition. In Vedaldi, A., Bischof, H., Brox, T., and Frahm, J.-M. (eds.), *Computer Vision – ECCV 2020*, pp. 325–342, Cham, 2020. Springer International Publishing. ISBN 978-3-030-58586-0.
- Wang, Y., Sun, J., Wang, C., Zhang, M., and Yang, M. Navigate beyond shortcuts: Debaised learning through the lens of neural collapse. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 12322–12331, June 2024.
- Williams, A., Nangia, N., and Bowman, S. A broad-coverage challenge corpus for sentence understanding through inference. In Walker, M., Ji, H., and Stent, A. (eds.), *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*, pp. 1112–1122, New Orleans, Louisiana, June 2018. Association for Computational Linguistics. doi: 10.18653/v1/N18-1101. URL <https://aclanthology.org/N18-1101/>.
- Wu, Y., Gardner, M., Stenetorp, P., and Dasigi, P. Generating data to mitigate spurious correlations in natural language inference datasets. In Muresan, S., Nakov, P., and Villavicencio, A. (eds.), *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pp. 2660–2676, Dublin, Ireland, May 2022. Association for Computational Linguistics. doi: 10.18653/v1/2022.acl-long.190. URL <https://aclanthology.org/2022.acl-long.190/>.
- Xu, Z., Jin, R., Shen, B., and Zhu, S. Nystrom approximation for sparse kernel methods: Theoretical analysis and empirical evaluation. *Proceedings of the AAAI Conference on Artificial Intelligence*, 29(1), Feb. 2015. doi: 10.1609/aaai.v29i1.9626. URL <https://ojs.aaai.org/index.php/AAAI/article/view/9626>.
- Xue, Z., Gao, Z., Ren, S., and Zhao, H. The modality focusing hypothesis: Towards understanding cross-modal knowledge distillation. In *The Eleventh International Conference on Learning Representations*, 2023. URL <https://openreview.net/forum?id=w0QXrZ3N-s>.
- Yang, Y., Zhang, H., Katabi, D., and Ghassemi, M. Change is hard: A closer look at subpopulation shift. In Krause, A., Brunskill, E., Cho, K., Engelhardt, B., Sabato, S., and Scarlett, J. (eds.), *Proceedings of the 40th International Conference on Machine Learning*, volume 202 of *Proceedings of Machine Learning Research*, pp. 39584–39622. PMLR, 23–29 Jul 2023. URL <https://proceedings.mlr.press/v202/yang23s.html>.
- Yu, C., Jeoung, S., Kasi, A., Yu, P., and Ji, H. Unlearning bias in language models by partitioning gradients. In Rogers, A., Boyd-Graber, J., and Okazaki, N. (eds.), *Findings of the Association for Computational Linguistics: ACL 2023*, pp. 6032–6048, Toronto, Canada, July 2023. Association for Computational Linguistics. doi: 10.18653/v1/2023.findings-acl.375. URL <https://aclanthology.org/2023.findings-acl.375/>.
- Zagoruyko, S. and Komodakis, N. Paying more attention to attention: Improving the performance of convolutional neural networks via attention transfer. In *International Conference on Learning Representations*, 2017. URL https://openreview.net/forum?id=Sks9_ajex.
- Zhang, Y., Baldridge, J., and He, L. PAWS: Paraphrase adversaries from word scrambling. In Burstein, J., Doran, C., and Solorio, T. (eds.), *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long and Short Papers)*, pp. 1298–1308, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics. doi: 10.18653/v1/N19-1131. URL <https://aclanthology.org/N19-1131/>.
- Zhou, B., Lapedriza, A., Khosla, A., Oliva, A., and Torralba, A. Places: A 10 million image database for scene recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 40(6):1452–1464, 2018. doi: 10.1109/TPAMI.2017.2723009.

A. Details on Debiasing Methods

Experiments are conducted on a comprehensive list of commonly used debiasing methods, each of which is designed with special formulation and assumptions.

- **Empirical Risk Minimization (ERM)** is the standard training method that minimizes the empirical risk on a dataset. This is akin to fine-tuning a pre-trained model on a dataset using cross-entropy loss with no debiasing strategy, which works for both image and text datasets.
- **HypothesisOnly-PoE** (Karimi Mahabadi et al., 2020) assumes the hypothesis part of NLI datasets contains biases. It trains a hypothesis-only (biased) model to measure the bias of each sample, and uses Product-of-Experts (PoE) (Hinton, 2002) to adjust the confidence of the debiased model according to the confidence of the biased model. This approach is evaluated on text datasets.
- **WeakLearner-PoE** (Sanh et al., 2021) leverages weak learners to capture and model bias, including bias of unknown type. It trains a 2-layer BERT as a biased model and exploits PoE to train the debiased model. This approach is evaluated on text datasets.
- **KernelWhitening** (Gao et al., 2022) aims at learning isotropic sentence embeddings with disentangled robust and spurious representations, with Nyström kernel (Xu et al., 2015). This approach is evaluated on text datasets.
- **AttentionPoE** (Wang et al., 2023) assumes that the attention to [CLS] token in text classification is biased and introduces PoE on attention weights to learn robust attention patterns for bias mitigation. This approach is evaluated on text datasets.
- **σ -Damp** (Puli et al., 2023) assuming the standard cross-entropy loss encourages models to prioritize shortcuts over robust features, this model proposes to scale the loss by a temperature. This approach is evaluated on image datasets.
- **DeepFeatReweight** (Kirichenko et al., 2023) discovers that simply retraining the last layer of a neural model—the classification layer in supervised tasks—on top of the existing biased feature extractor is good strategy for bias mitigation. This approach is evaluated on image datasets.
- **PerSampleGrad** (Ahn et al., 2023) trains a debiased model with non-uniform sampling probability, obtained from per-sample gradient norm of a biased model. This approach is evaluated on image datasets.

B. Implementation details

We follow previous debiasing works for implementation details. For text datasets, we train each debiasing method with Adam optimizer, learning rate $5e-5$, 5 epochs, both KD and Non-KD. For image datasets, we train each debiasing method with Adam optimizer, learning rate $4e-5$, 100 epochs, both KD and Non-KD. For all other hyperparameters, we follow each debiasing method’s best-performing setting.

C. Results on Image Datasets

On image datasets, we observe similar results on text datasets. Specifically, we see that KD fall short on distilling the debiasing capabilities. Such ability is transferred more smoothly as teacher and student get similar in scale.

D. Detailed Results on Debiasing Methods and Backbones

We present the detailed results of individual debiasing method and backbone below.

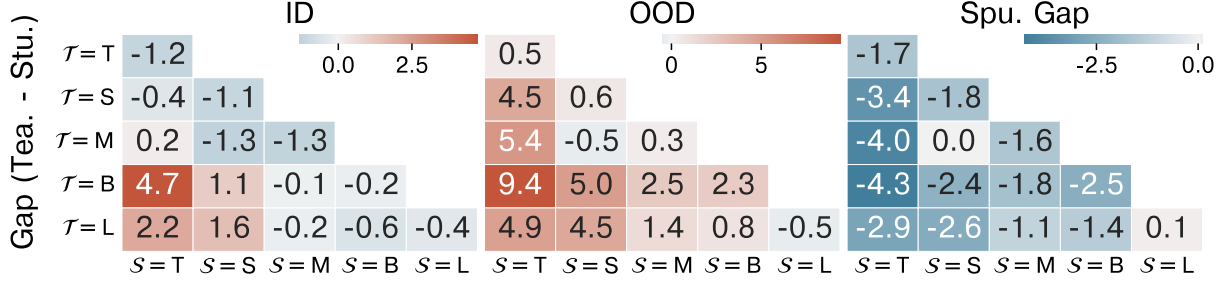


Figure 11. C1: **Teacher vs. Student**: average performance gaps between teacher and student models on ID, OOD, and Spurious Gap across image datasets. X-axis and Y-axis show the scale of student (S) and teacher (T) respectively. Each cell shows the performance gap between corresponding scales of a teacher and a student.

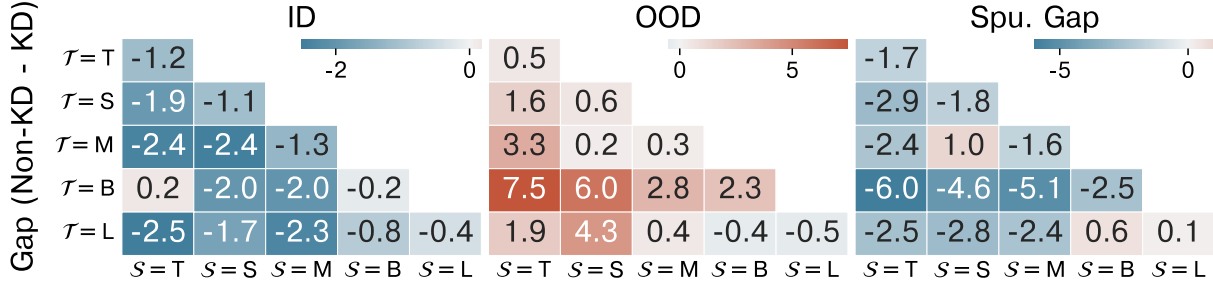


Figure 12. C2: **Non-KD vs. KD**: average performance gaps between Non-KD and KD models on ID, OOD, and Spurious Gap across image datasets. X-axis and Y-axis show the scale of student (S) and teacher (T) respectively. Each cell shows the performance gap between corresponding scales of a teacher and a student.

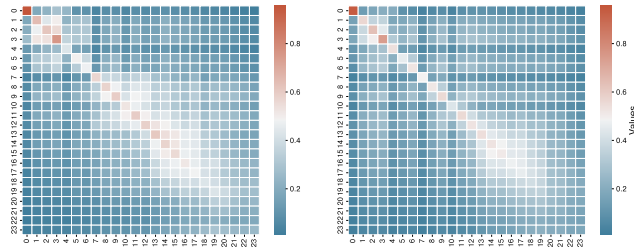


Figure 13. C2: **KD vs. Non-KD**: Centralized Kernel Alignment. Higher values indicate higher similarity. X-axis and Y-axis refer to the layers of KD (S) and Non-KD (f_S) respectively.

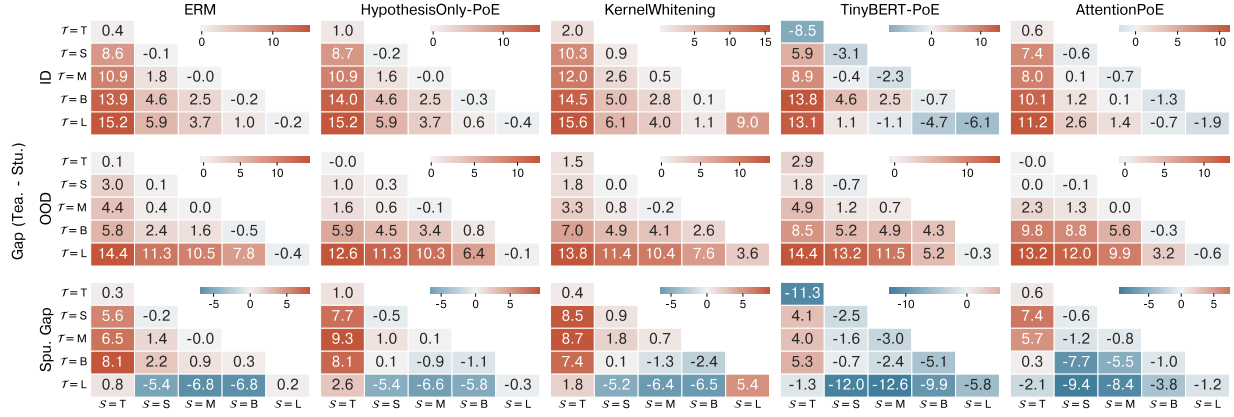


Figure 14. C1: Teacher vs. Student: average performance gaps between teacher and student models on ID, OOD, and Spurious Gap on BERT.

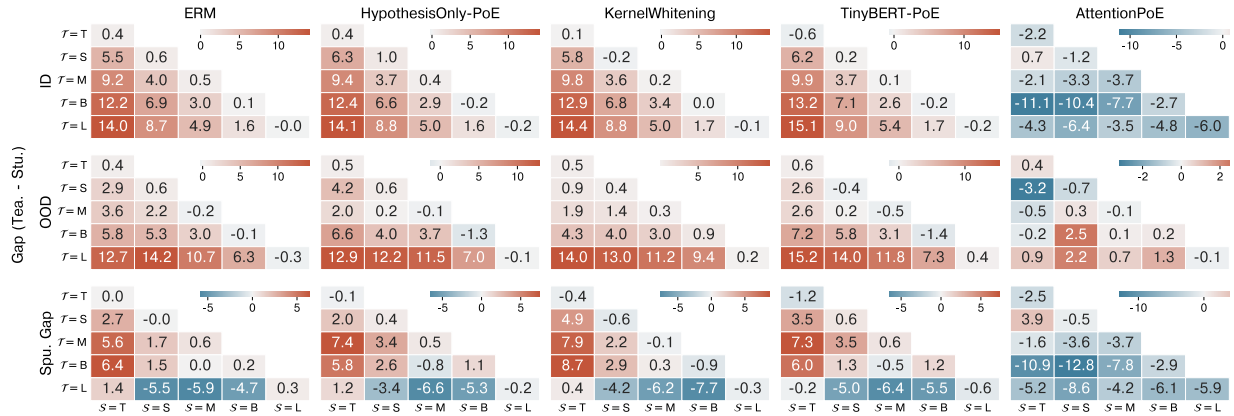


Figure 15. C1: Teacher vs. Student: average performance gaps between teacher and student models on ID, OOD, and Spurious Gap on T5.

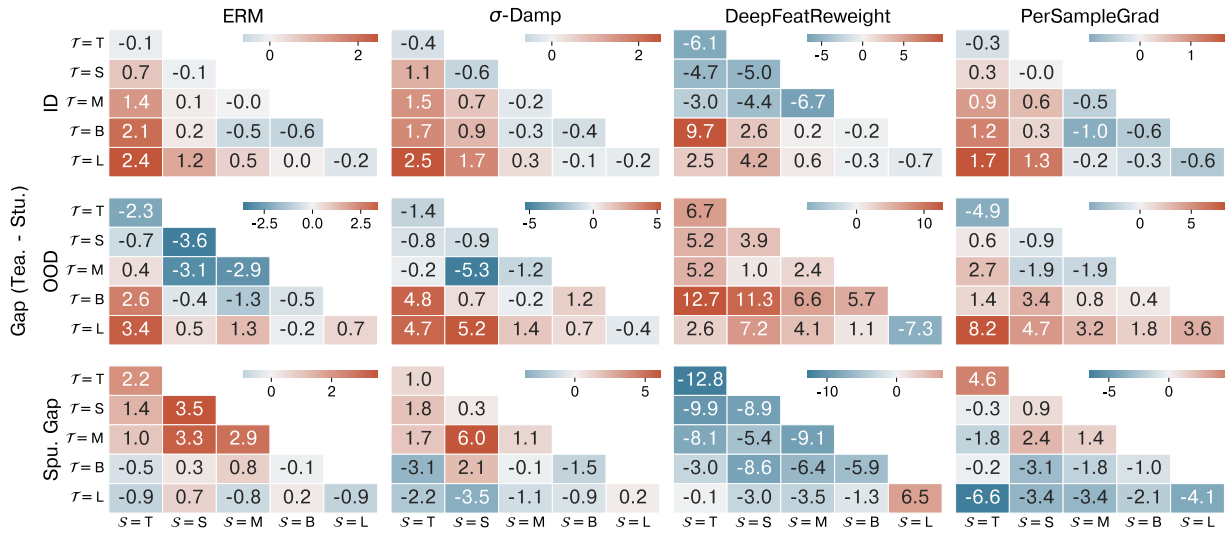


Figure 16. C1: Teacher vs. Student: average performance gaps between teacher and student models on ID, OOD, and Spurious Gap on ResNet.

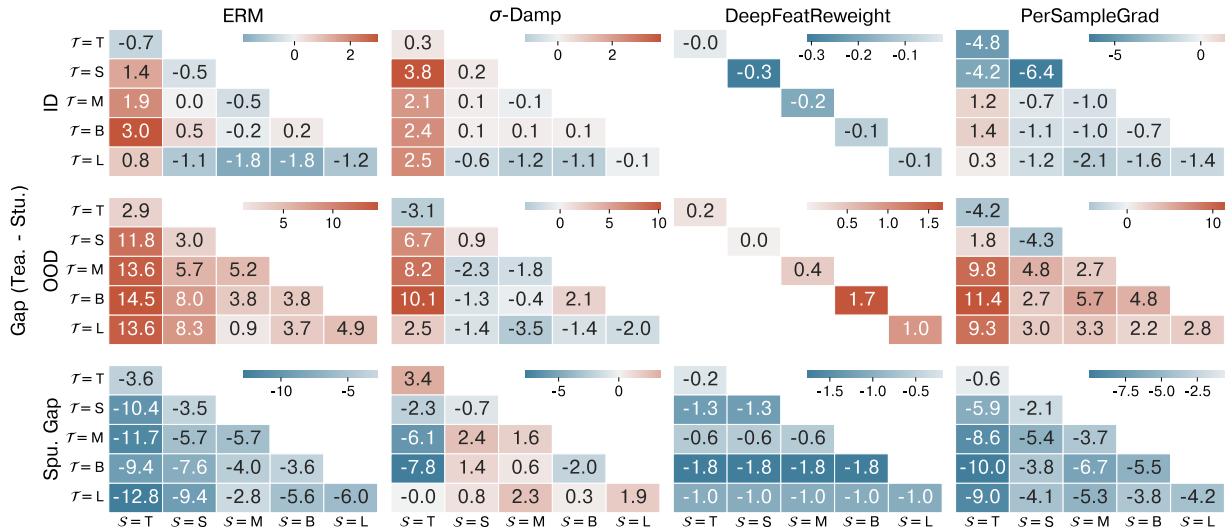


Figure 17. C1: Teacher vs. Student: average performance gaps between teacher and student models on ID, OOD, and Spurious Gap on ViT.

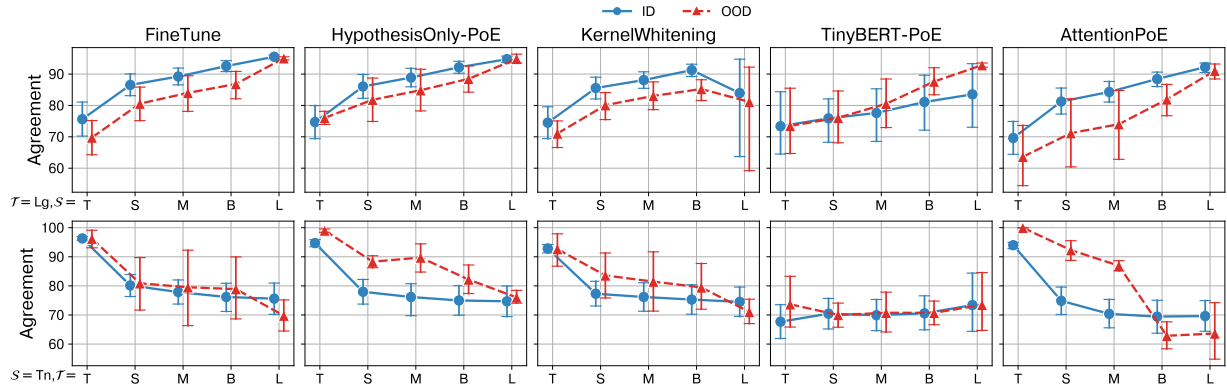


Figure 18. C1: Teacher vs. Student: prediction agreement on BERT. Left: varying S (X-axis) given a fixed teacher with $\mathcal{T} = L$. Right: varying $\mathcal{T}$ (X-axis) given a fixed student with $\mathcal{T} = T$. Agreement increases as the scale of teacher and student get closer.