



Tradeoffs Between Contrastive and Supervised Learning: An Empirical Study

Ananya Karthik, Mike Wu, Noah Goodman, Alex Tamkin

{ananya23, wumike, ngoodman, atamkin}@stanford.edu

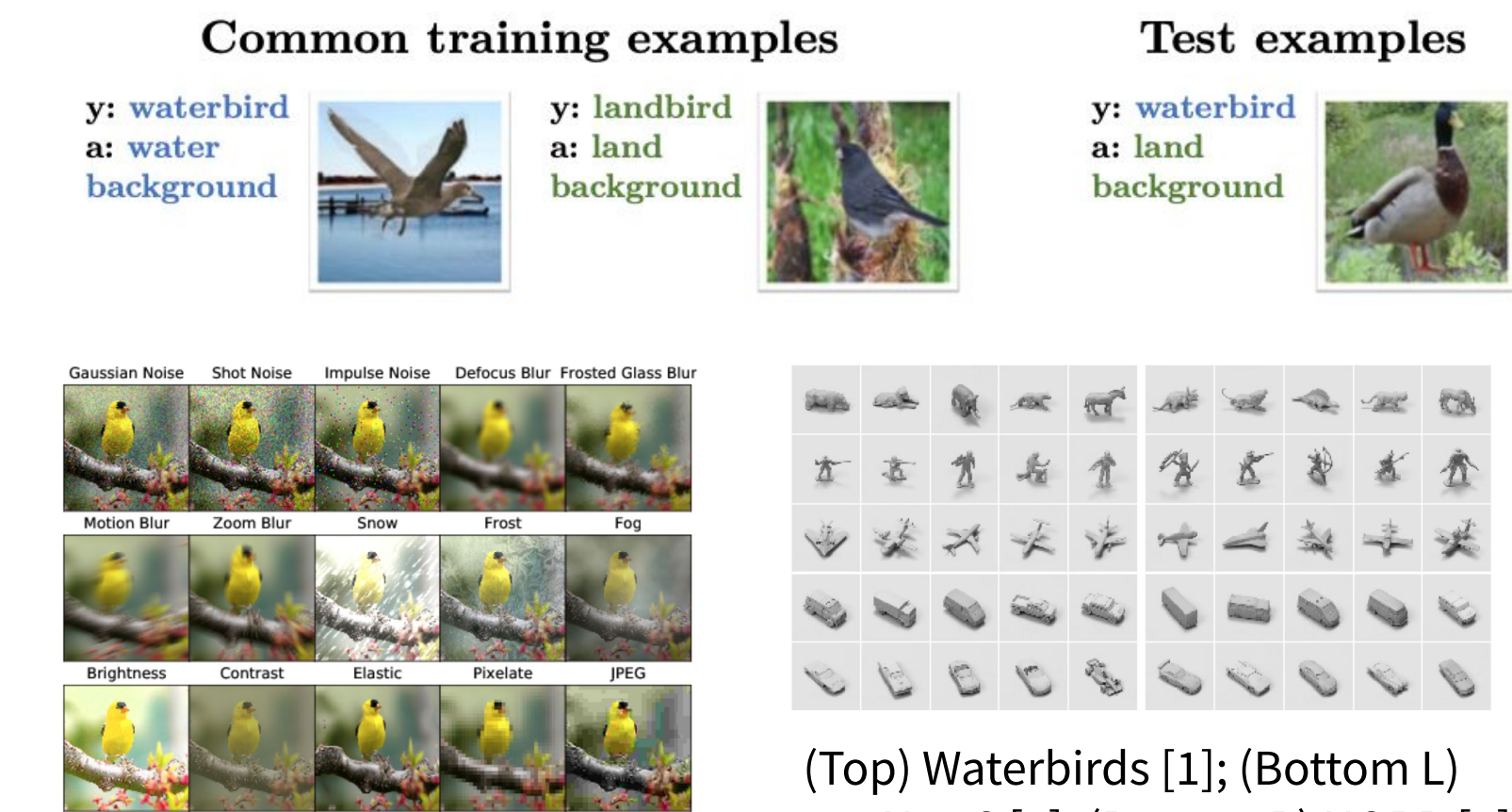
Department of Computer Science, Stanford University

Motivation

- Self-supervised pretraining for computer vision has grown in popularity recently due to the cost of annotating data
- Contrastive learning has achieved state-of-the-art results, underscoring the need for studying the real-world tradeoffs between contrastive and supervised pretraining. Specifically:
 - Is contrastive learning better **across all compute budgets**?
 - For larger compute budgets, is supervised pretraining better **on tasks where an object-centric bias is important**?

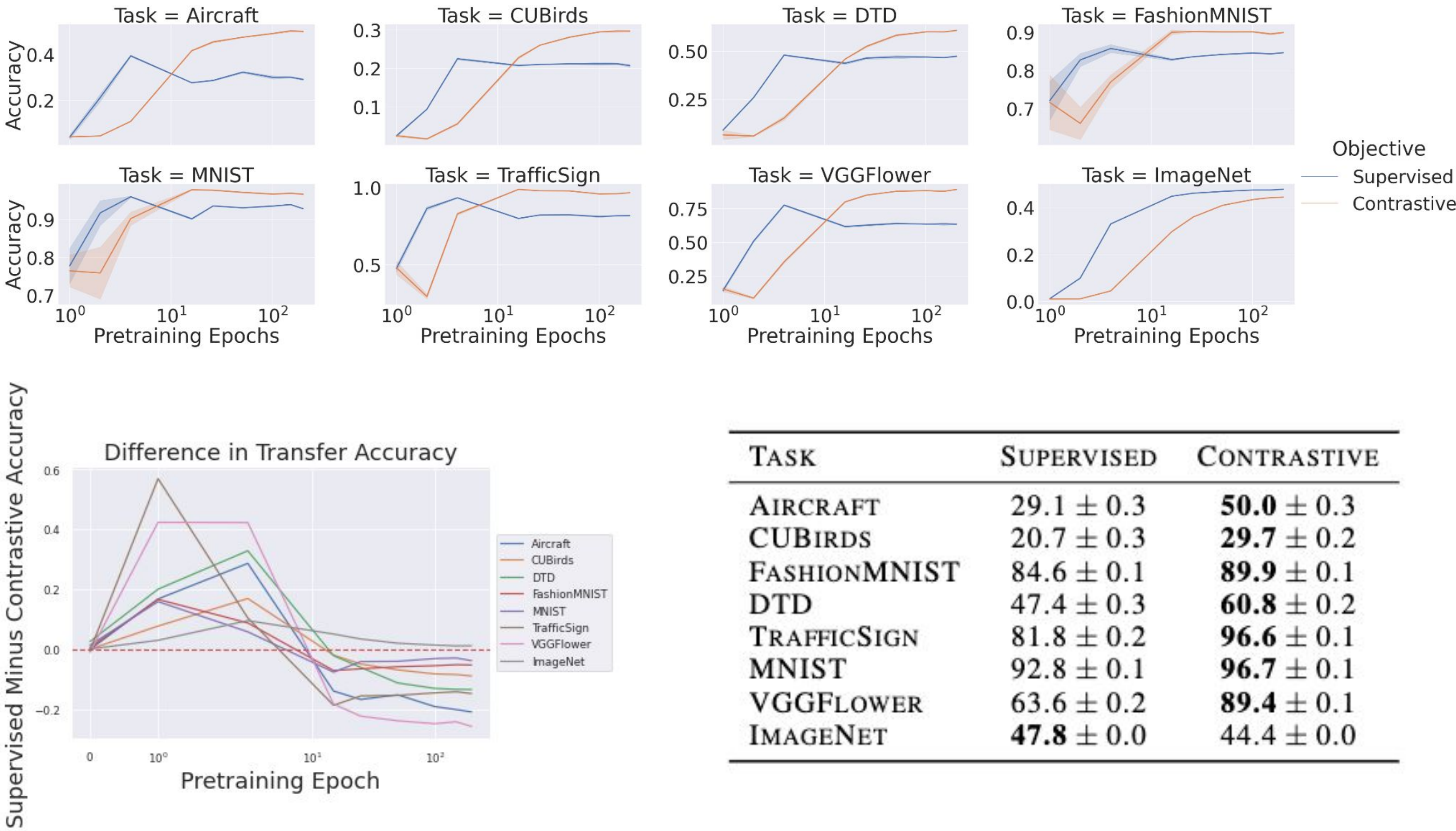
Methodology

- Experimental Settings
 - Pretraining: 2 ResNet-18 models on ImageNet (200 epochs)
 - Standard cross entropy loss for supervised, InfoNCE objective for the contrastive model
 - Transfer: Linear evaluation protocol (100 epochs)
 - Datasets:
 - Q1 (Transfer across compute budgets): Aircraft, CUBirds, FashionMNIST, DTD, TrafficSign, MNIST, VGGFlower, ImageNet
 - Q2 (Object-centric bias): Waterbirds, Norb, ImageNet-C (below)



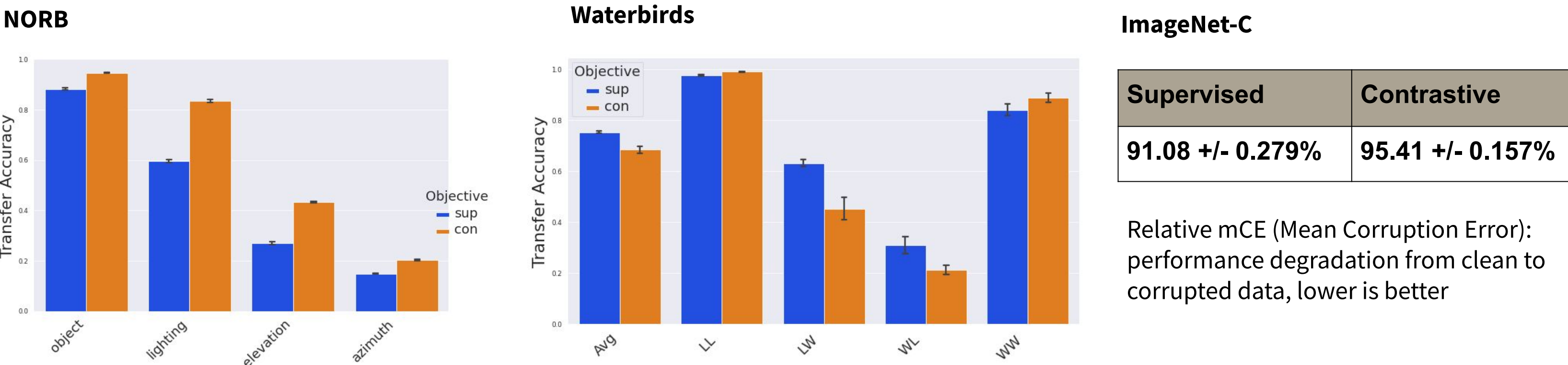
Results

1. Learning Dynamics Across Compute Budgets



Downstream accuracy of contrastive and supervised models on 8 transfer tasks for different pretraining budgets

2. Downstream Effects of Biases Acquired During Pretraining



Does the object-centric bias of supervised learning improve downstream performance on transfer tasks?
We find strong effects for Waterbirds and ImageNet-C, but weaker effects for NORB.

Conclusions

- Contrastive learning is **not necessarily better across all compute budgets**: different pretraining algorithms produce better representations at different budgets
 - Transfer performance does not increase monotonically across pretraining → potential misalignment between representations learned for pretraining vs transfer
 - While the contrastive model eventually achieves higher performance, for the first 10-15 epochs the supervised model yields better representations for downstream tasks → potential differences in the two representation learning processes
 - We encourage developers of new pretraining techniques to release learning dynamics curves
- Contrastive learning is **not necessarily better across all tasks**: the supervised model eventually achieves worse downstream accuracy on most tasks, but the object-centric bias of ImageNet pretraining aids transfer on some tasks, especially WaterBirds (reliance on spurious correlations) and ImageNet-C (robustness to common corruptions)

Future Work

- Investigating whether these conclusions hold across a wide range of architectures, hyperparameters, datasets, and training objectives
- Exploring other dimensions along which pretraining algorithms differ (e.g. Cole et al. 2021 and Horn et al. 2021 find that supervised learning tends to perform better on fine-grained classification tasks)
- Studying how pretraining objectives shape the behavior of models in ambiguous scenarios

[1] Sagawa et al. 2019

[2] Hendrycks and Dietterich 2019

[3] LeCun et al. 2004