

During the rebuttal process we added several experimental and theoretical results that showed further benefits of our methods:

1. We showed our methods significantly outperform an additional baseline suggested by the reviewers and achieve zero-shot generalization in several new settings.
2. We demonstrated GL-Net can be trained **entirely on publicly available finetunes** and achieve **excellent test accuracy on industry-relevant tasks**.
3. We proved that our theoretical results extend naturally to mixed-rank input settings.

We included the second result in our new submission. Our rebuttals are also visible in the reviews PDF.

Most importantly, we believe we fully addressed the concerns of reviewer bUbu, the lone reviewer to rate our work negatively. We demonstrated strong generalization across different LoRA structures and data sources. We also ran several new ablations that showed our method performs robustly across hyperparameter settings and developed new directions for future work.