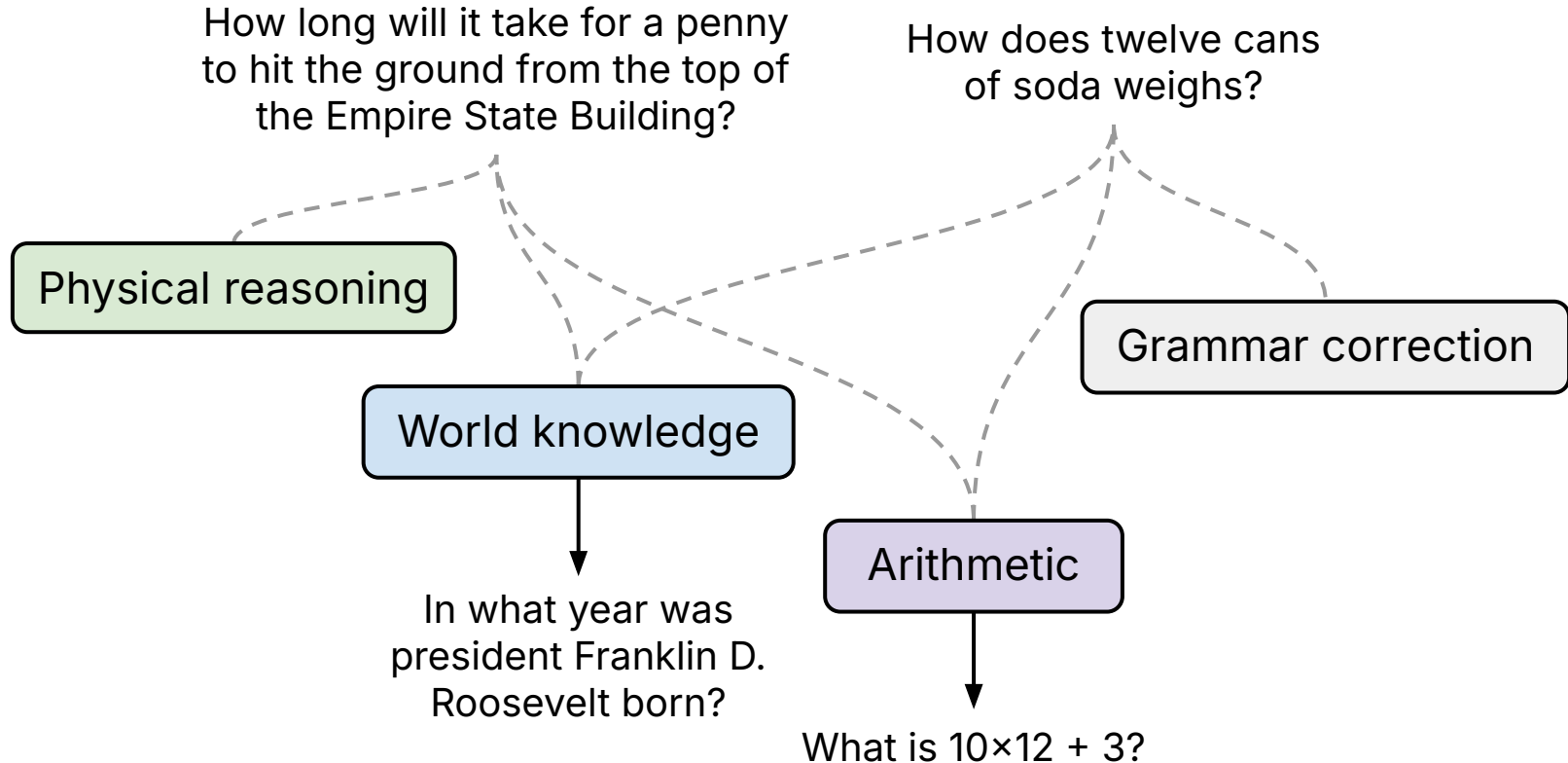


Merging and MoErging for Compositional Generalization

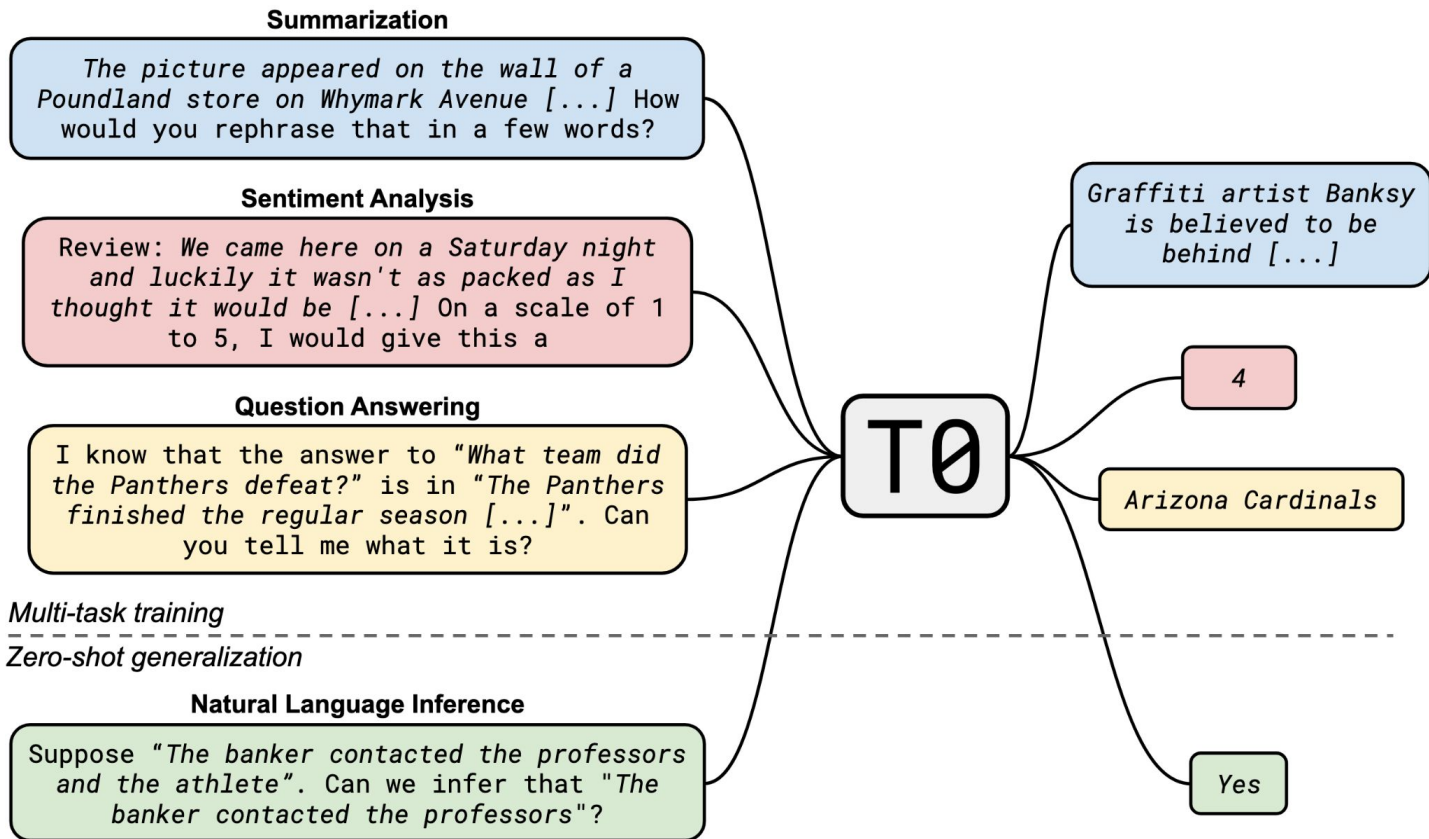
Colin Raffel



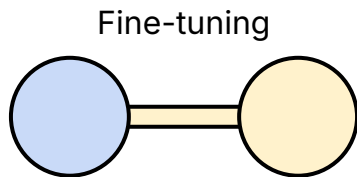
Tasks can be considered as a composition of skills



Multitask models can generalize to new tasks

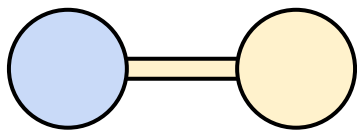


Merging models enables new paths for transferring capabilities

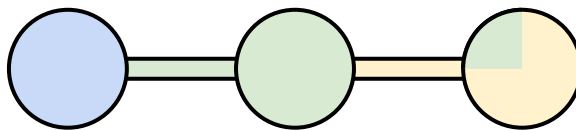


Merging models enables new paths for transferring capabilities

Fine-tuning

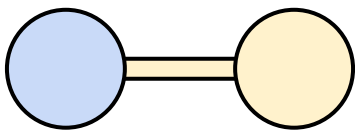


Continual learning

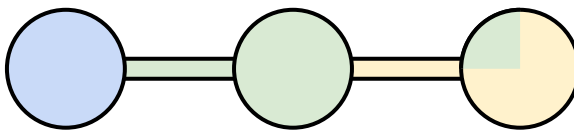


Merging models enables new paths for transferring capabilities

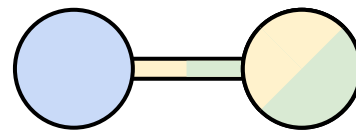
Fine-tuning



Continual learning

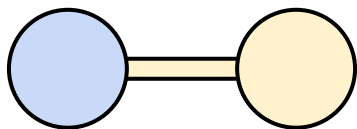


Multitask learning

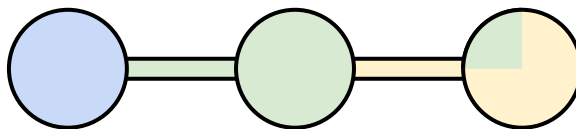


Merging models enables new paths for transferring capabilities

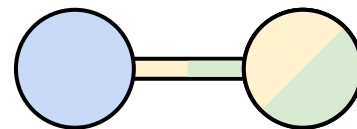
Fine-tuning



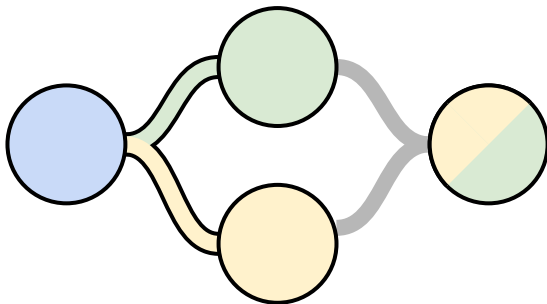
Continual learning



Multitask learning

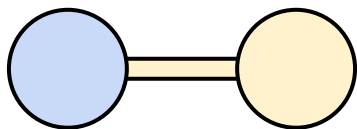


Model merging

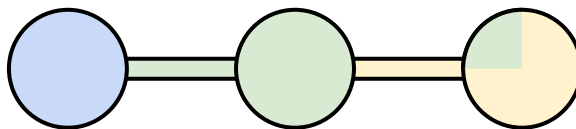


Merging models enables new paths for transferring capabilities

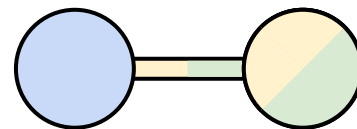
Fine-tuning



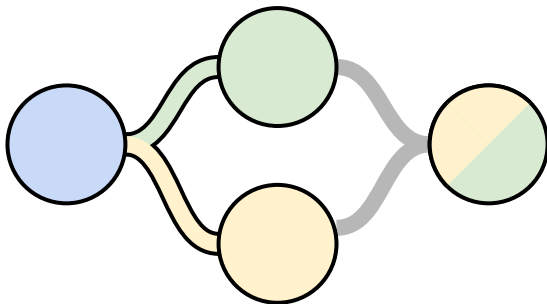
Continual learning



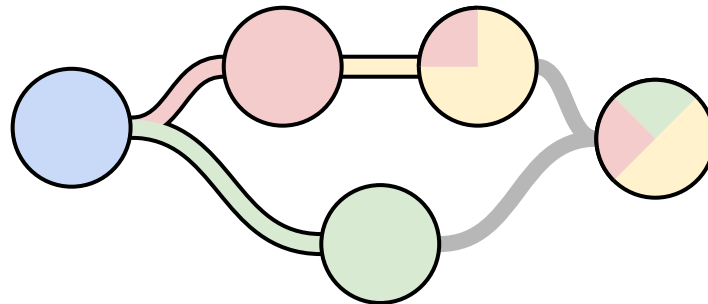
Multitask learning



Model merging

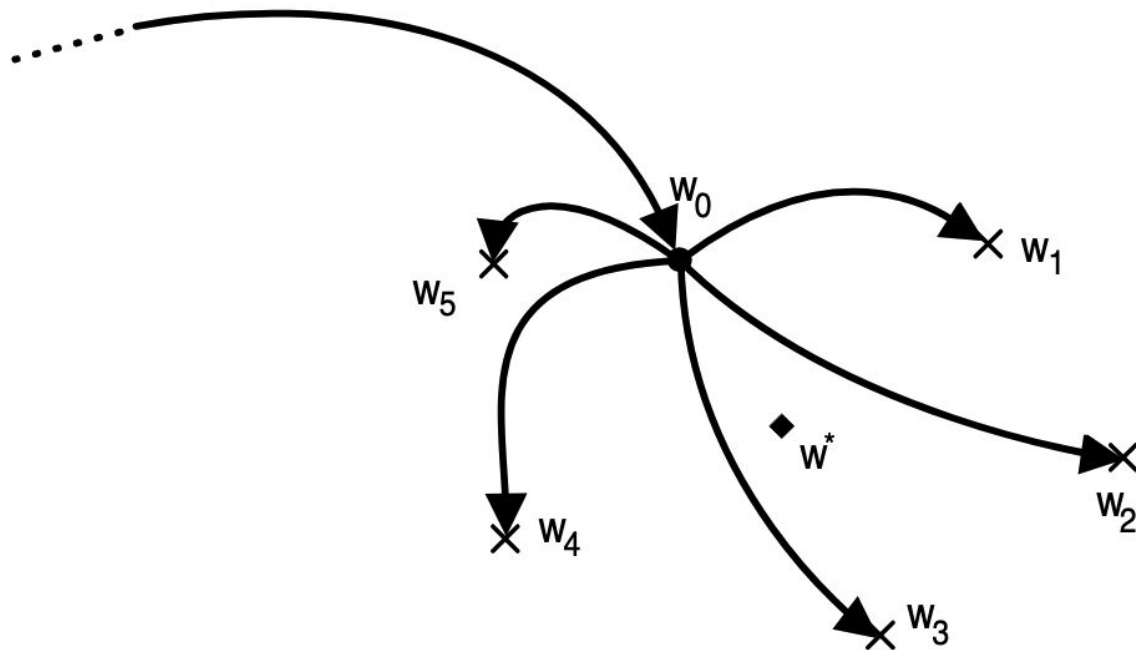


... and beyond



"Vanilla" merging is just parameter averaging

weight space



Model merging as an optimization problem

$$\arg \max_{\theta} \sum_{i=1}^M \lambda_i \log p(\theta | \mathcal{D}_i)$$

Parameter averaging makes an isotropic Gaussian assumption

$$\arg \max_{\theta} \sum_{i=1}^M \lambda_i \log p(\theta | \mathcal{D}_i)$$

$\downarrow \theta \sim \mathcal{N}(\theta_i, \mathbf{I})$

$$\theta^* = \frac{\sum_i \lambda_i \theta_i}{\sum_i \lambda_i}$$

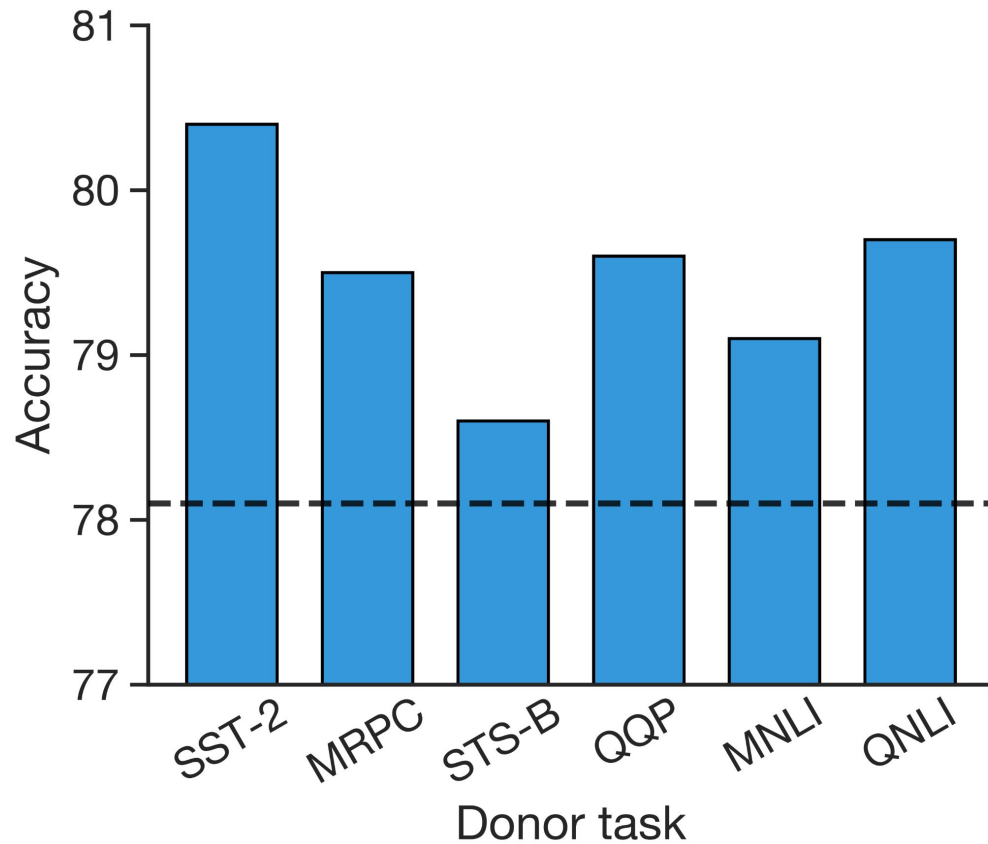
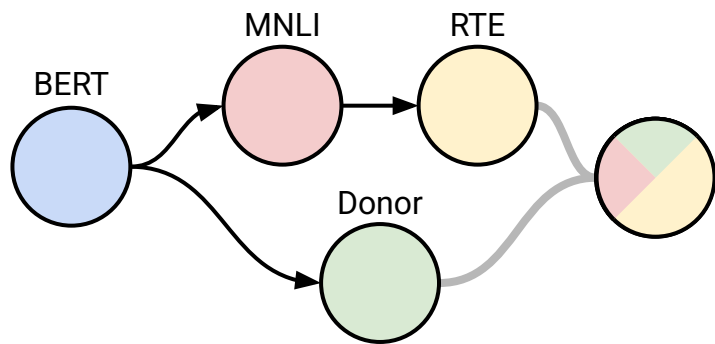
Fisher merging uses the Laplace approximation

$$\arg \max_{\theta} \sum_{i=1}^M \lambda_i \log p(\theta | \mathcal{D}_i)$$

$\downarrow \theta \sim \mathcal{N}(\theta_i, \hat{F}_i^{-1})$

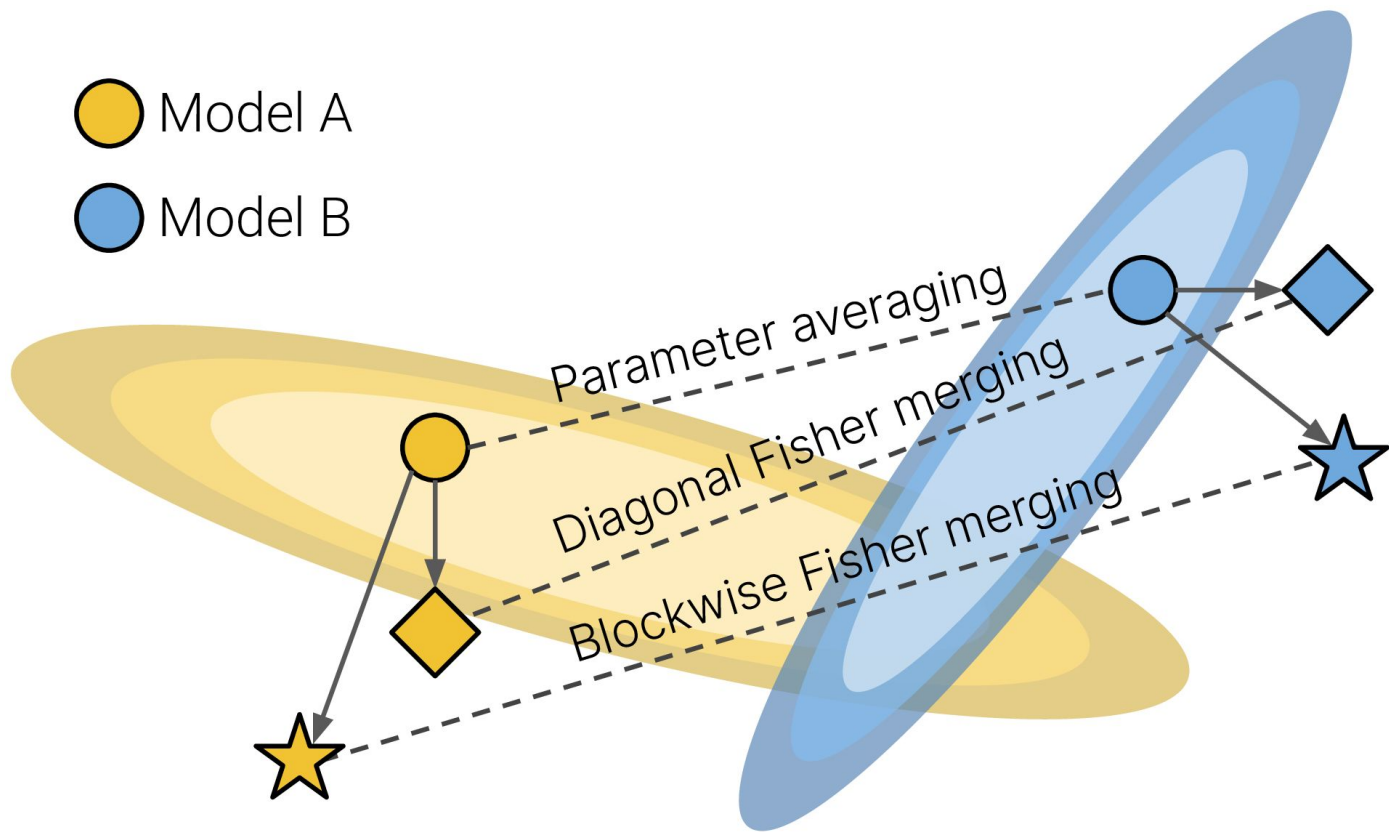
$$\theta^* = \frac{\sum_i \lambda_i \hat{F}_i \odot \theta_i}{\sum_i \lambda_i \hat{F}_i}$$

Fisher merging can combine the capabilities of different models



From "Merging Models with Fisher-Weighted Averaging" by Matena et al.

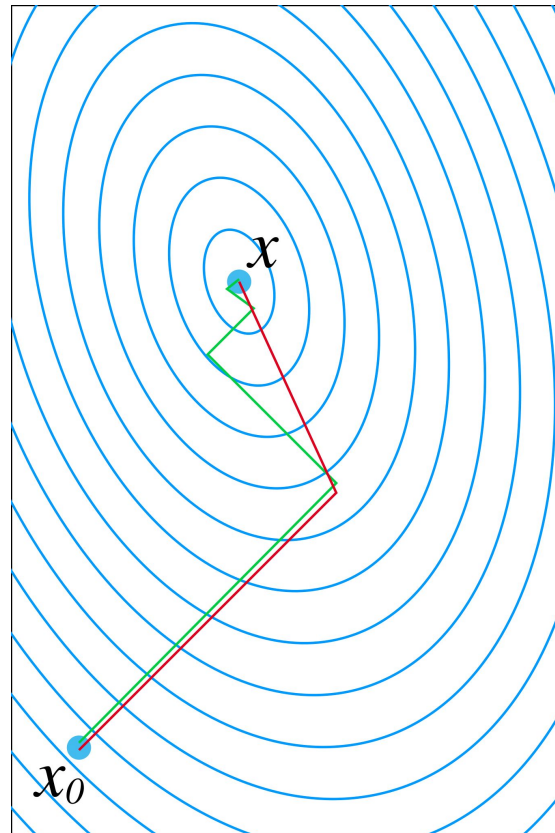
Merging models based on loss landscape curvature



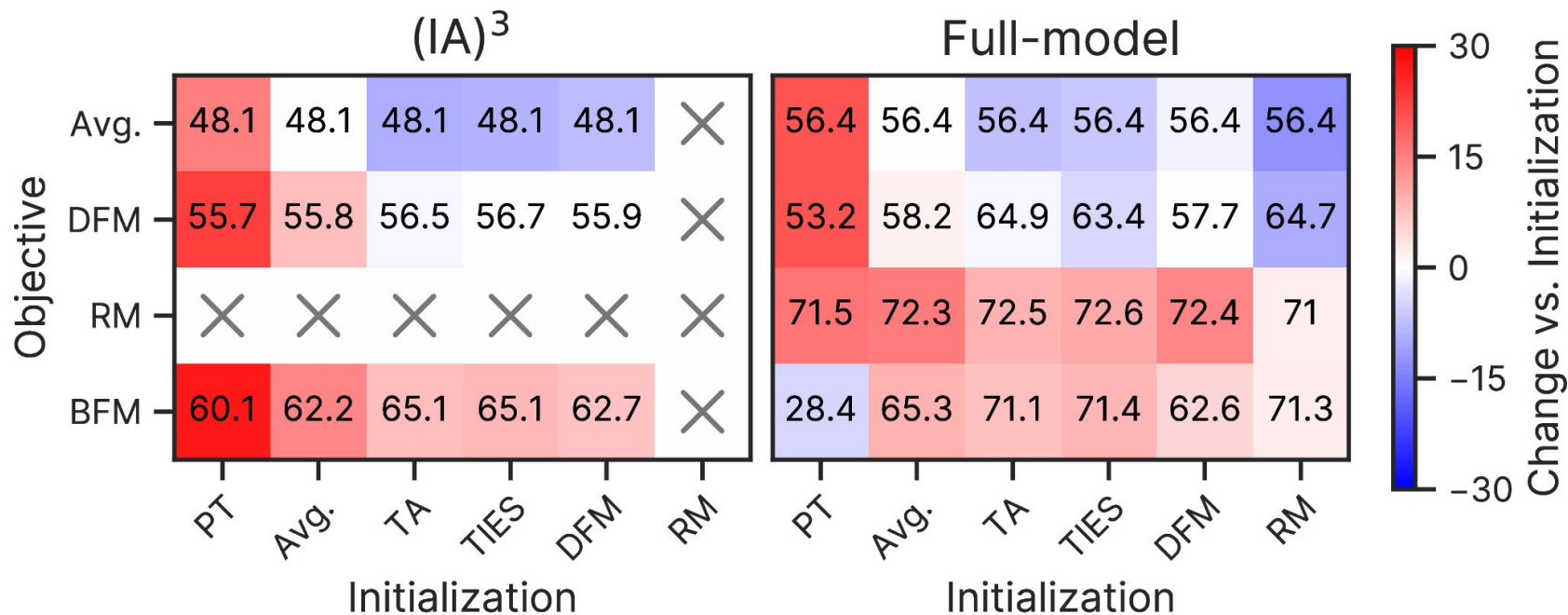
From "Merging by Matching Models in Task Subspaces" by Tam et al.

Solving a general merging problem with the conjugate gradient method

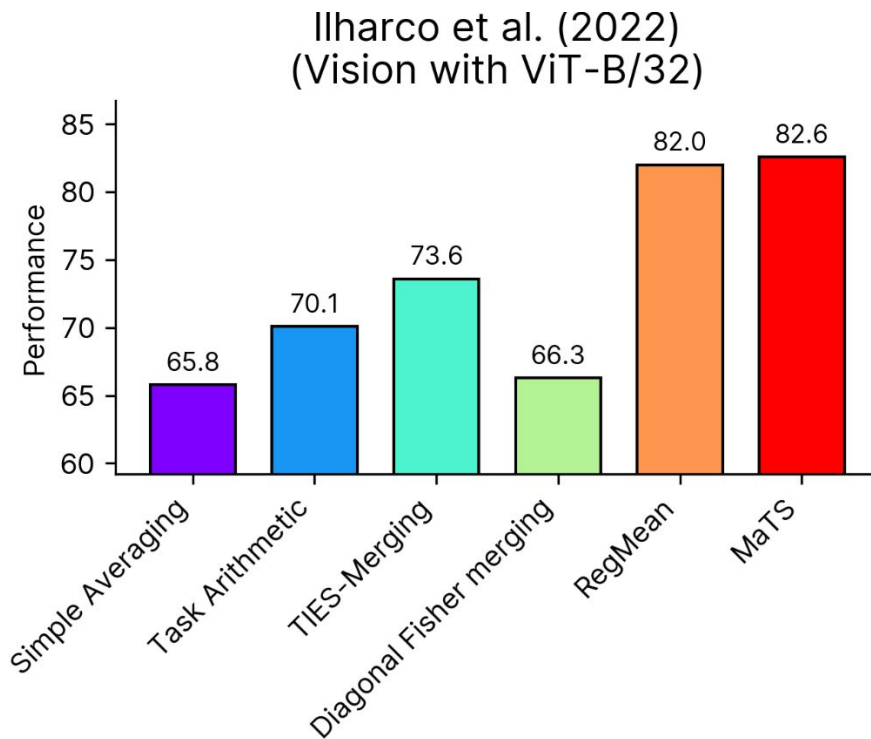
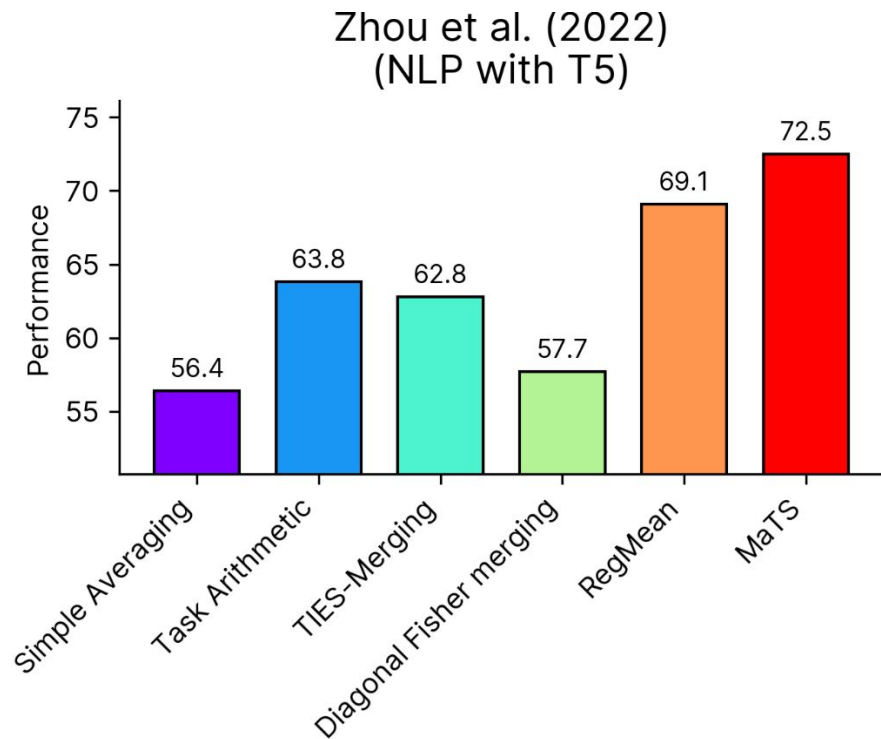
$$\boldsymbol{\theta}^* = \left(\sum_{m=1}^M \mathbf{C}_m \right)^{-1} \left(\sum_{m=1}^M \mathbf{C}_m \boldsymbol{\theta}_m \right)$$



MaTS allows flexibly combining merging objectives and initializations...

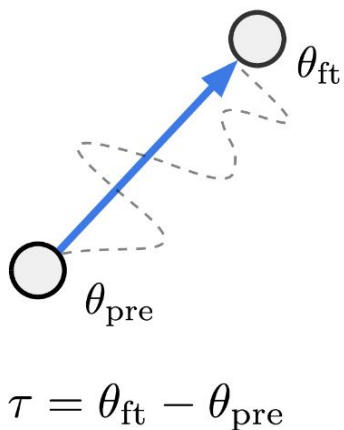


... and achieved state-of-the-art resulted across settings

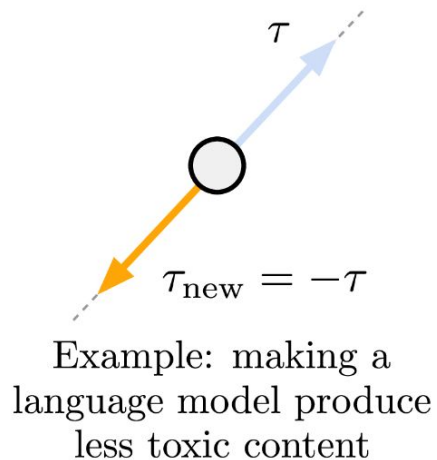


Merging models with task vectors

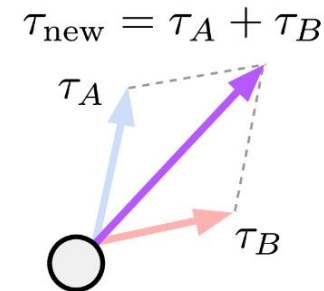
a) Task vectors



b) Forgetting via negation

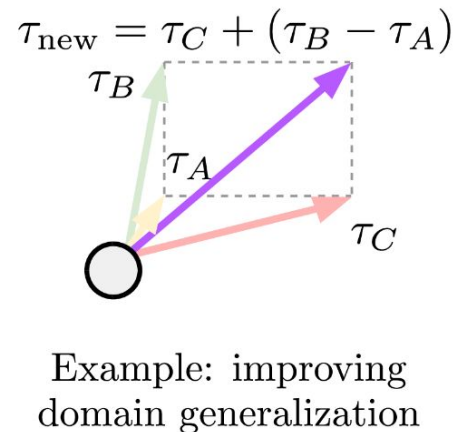


c) Learning via addition

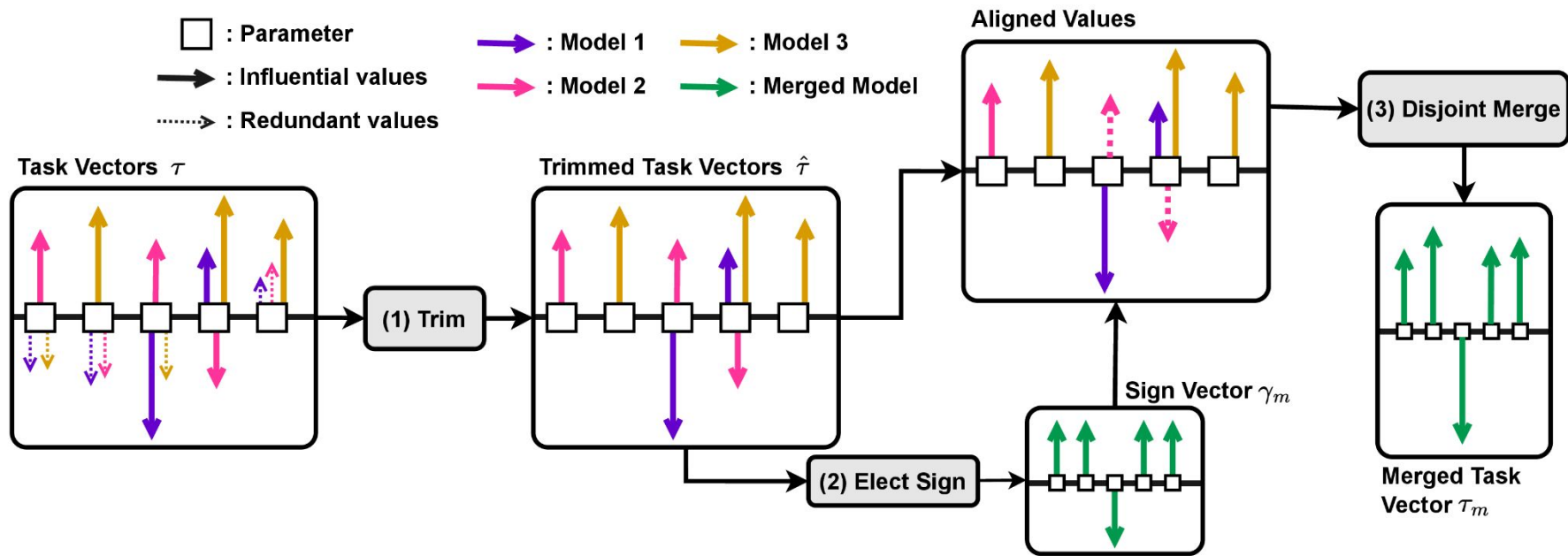


Example: building a multi-task model

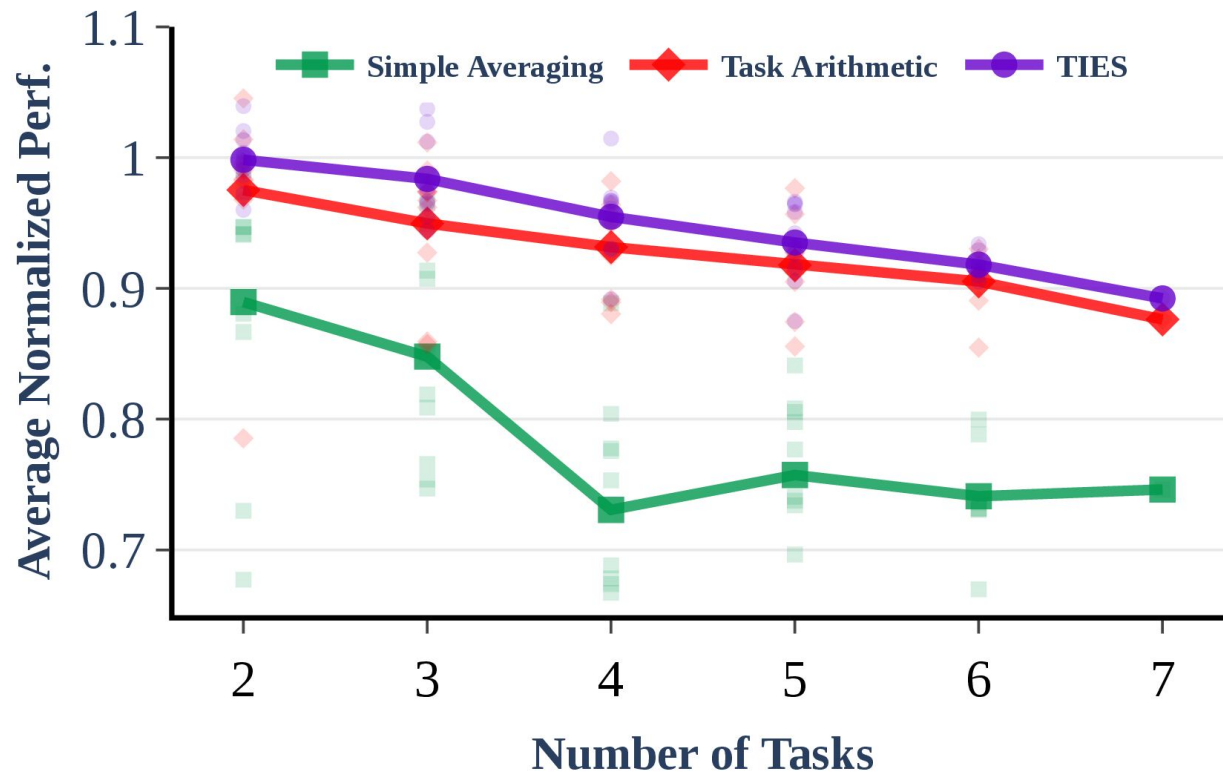
d) Task analogies



TIES Merging resolves interference between task vectors

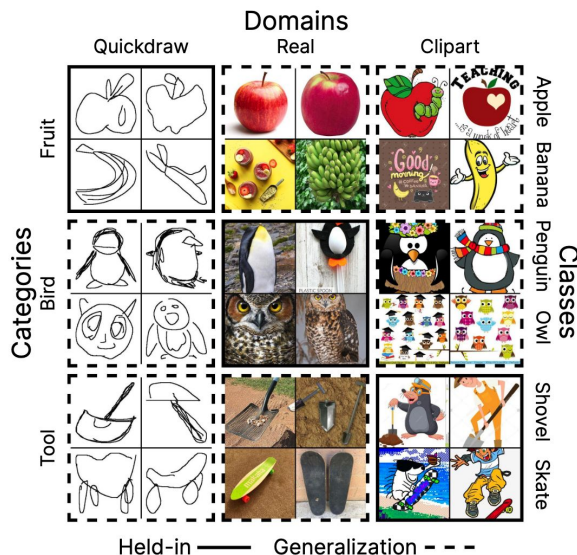


TIES helps retain specialist model performance



From "Resolving Interference When Merging Models" by Yadav et al.

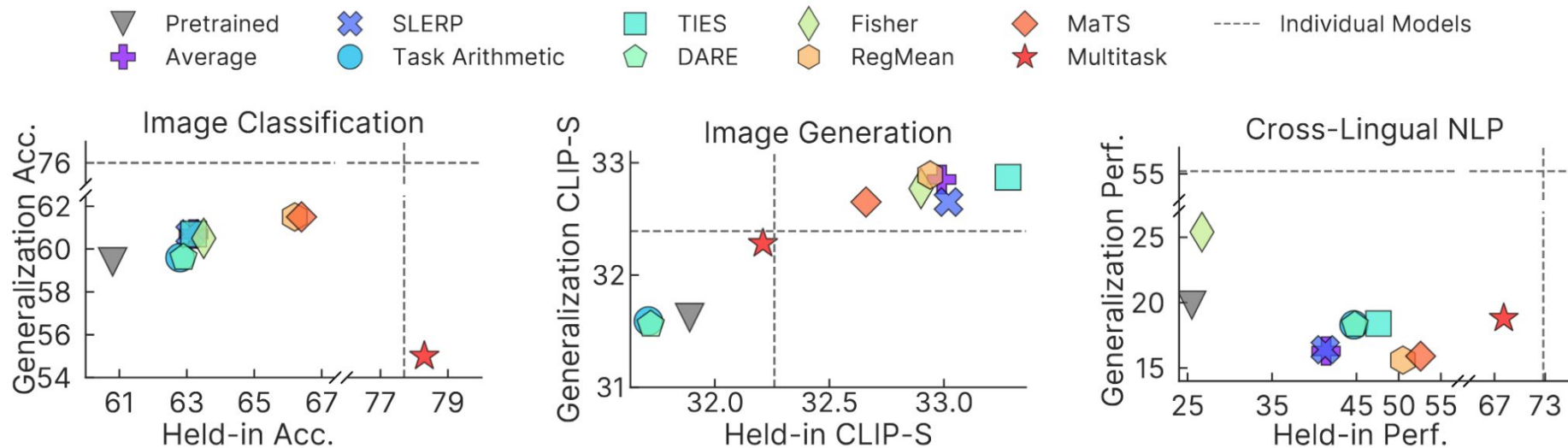
Is merging actually doing what we want?



NLP Task ↓ / Language →	English	Arabic	Thai	German	Korean
Question-Answering (SQuAD/XQuAD)					
Natural Language Inference (XNLI)					
Summarization (WikiLingua)					
Word Sense Disambiguation (WiC/XLWiC)					
Is question answerable? (TyDiQA)					

From "Realistic Evaluation of Model Merging for Compositional Generalization" by Tam et al.

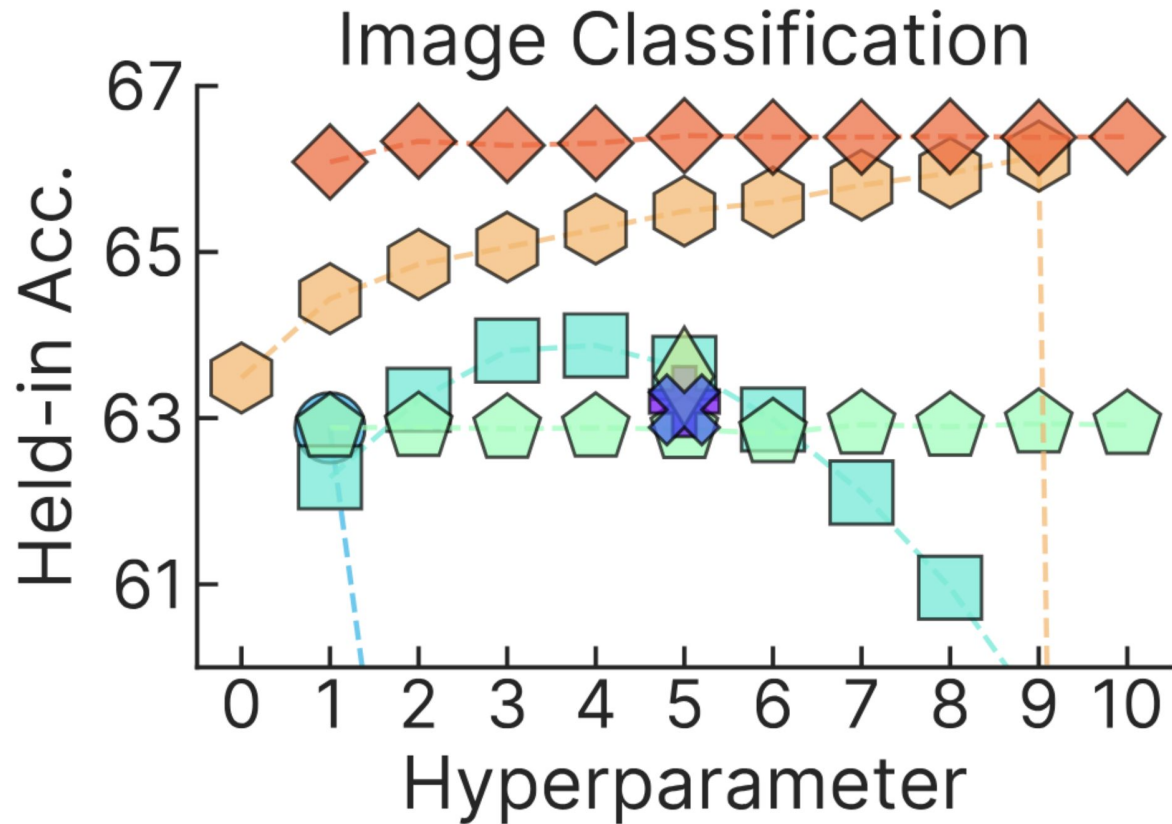
... not really.



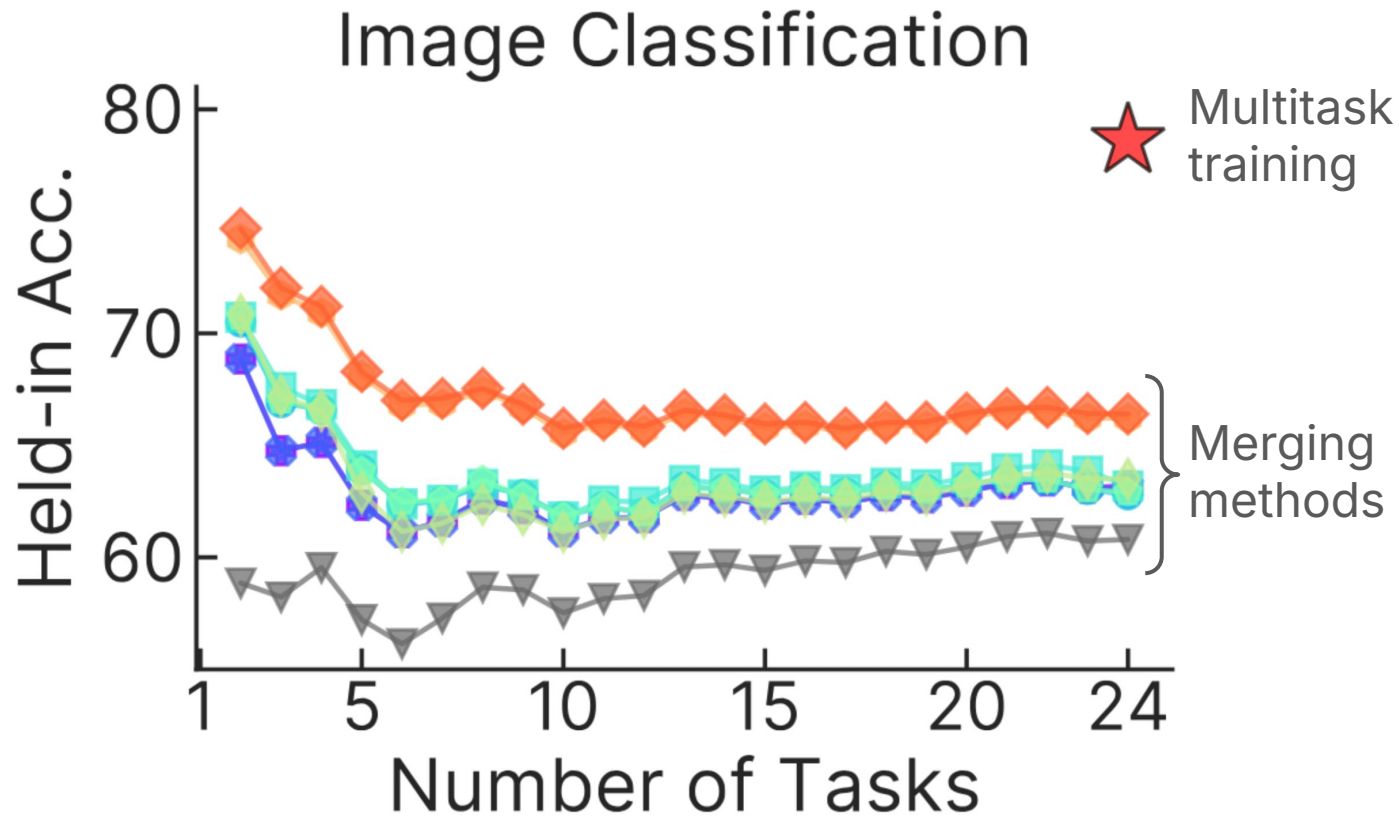
Merging methods also have different practical requirements...

	Prerequisites			Hparams?	Computational cost (FLOPs)	
	θ_p	Stats	Data		Merging	Statistics
Average					Mdk	
SLERP					$\mathcal{O}((5M - 2)dk)$	
Task Arith.	$\times$		$\times$	$\checkmark$	$(2M + 1)dk$	
DARE	$\times$		$\times$	$\checkmark$	$(6M + 1)dk$	
TIES	$\times$		$\times^*$	$\checkmark^*$	$(4M + 1)dk$	$\mathcal{O}(MKdk)$
Fisher		$\times^*$	$\times^*$		$(3M - 1)dk$	$4MTd^2k$
RegMean		$\times^*$	$\times^*$	$\checkmark$	$\mathcal{O}((M + 2)d^2k)$	MTd^2k
MaTS		$\times^*$	$\times^*$	$\checkmark$	$\mathcal{O}((M + N)d^2k)$	$4MTd^2k$

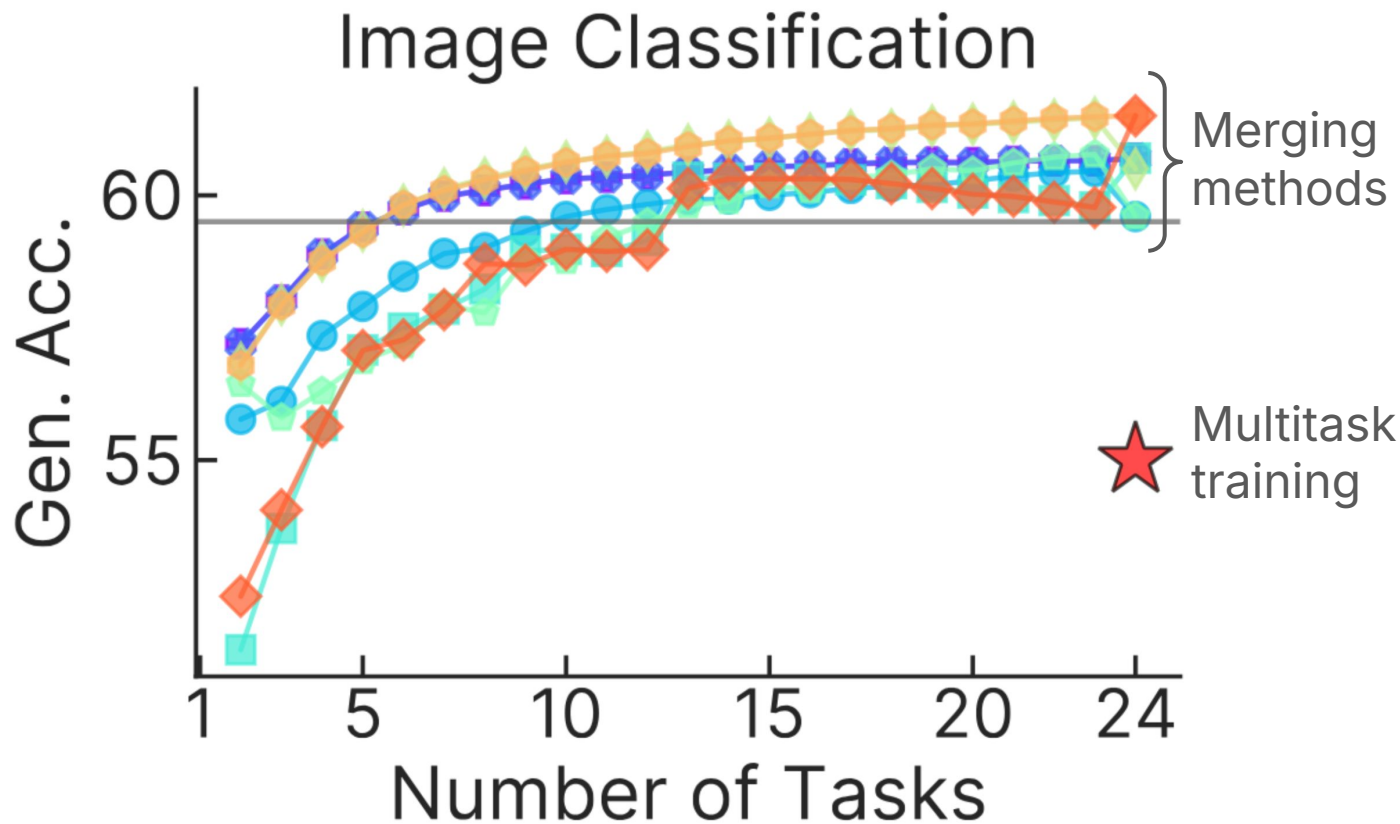
... and different hyperparameter sensitivity



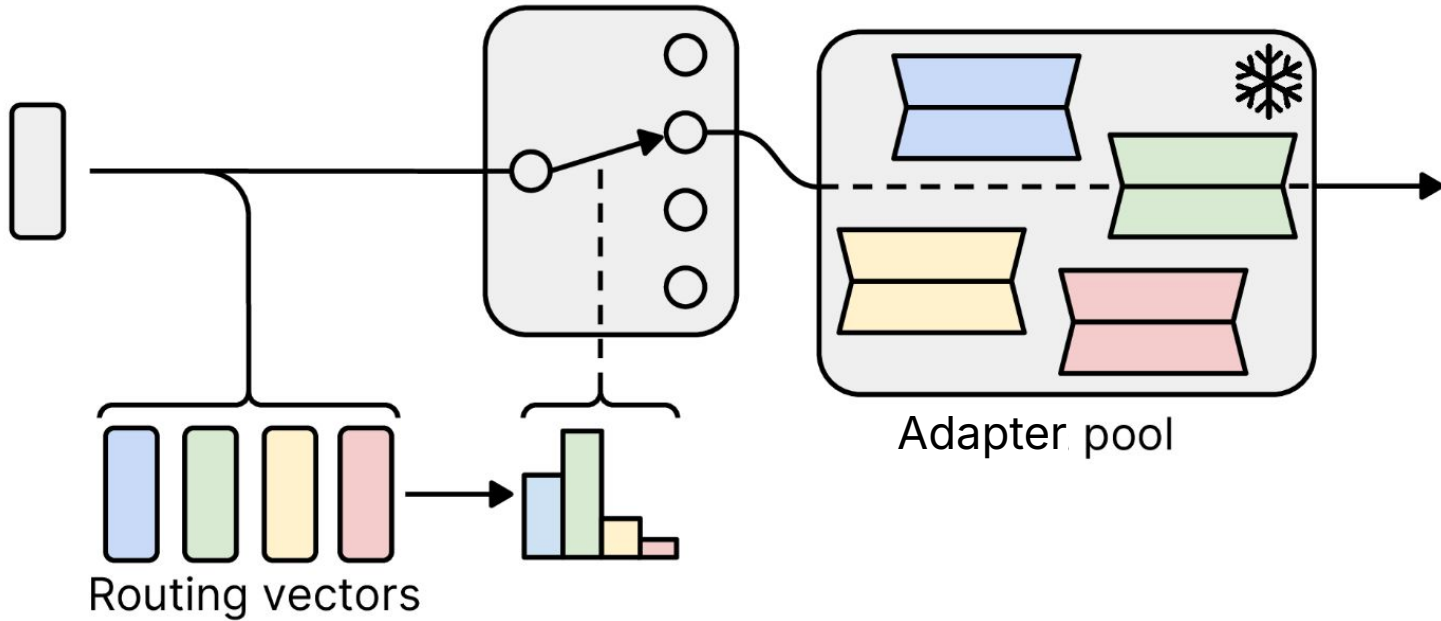
Multitask performance is still poor as you the number of models...



... but the picture for generalization is better.

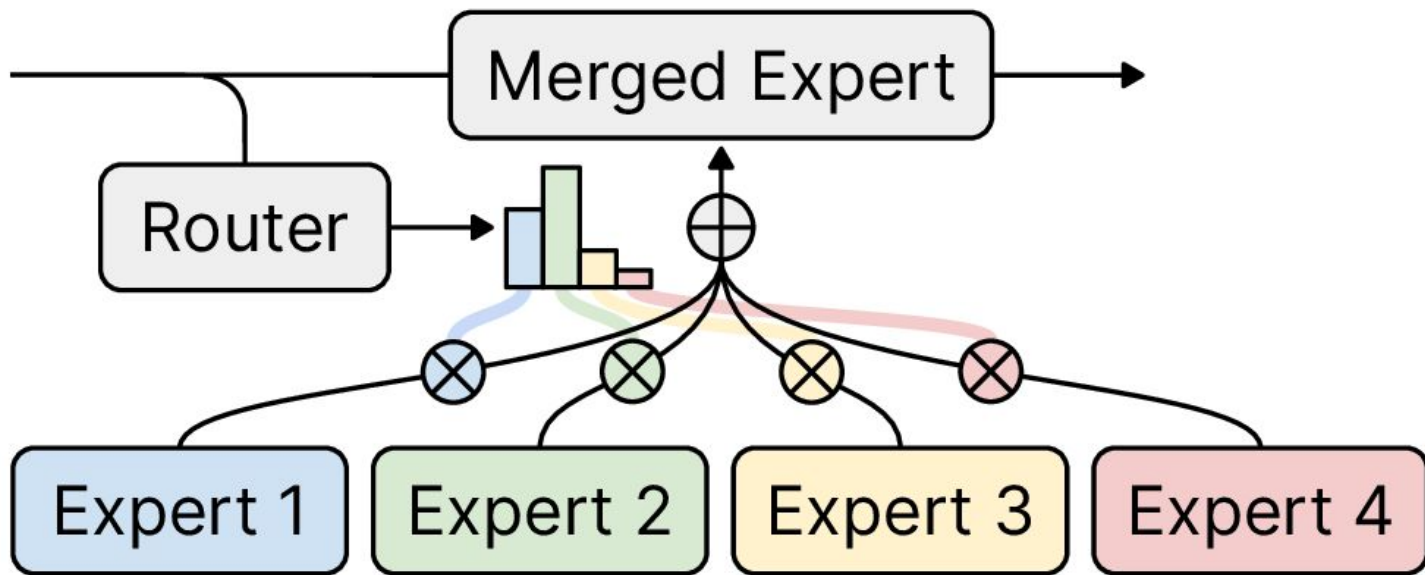


An alternative: MoErging?



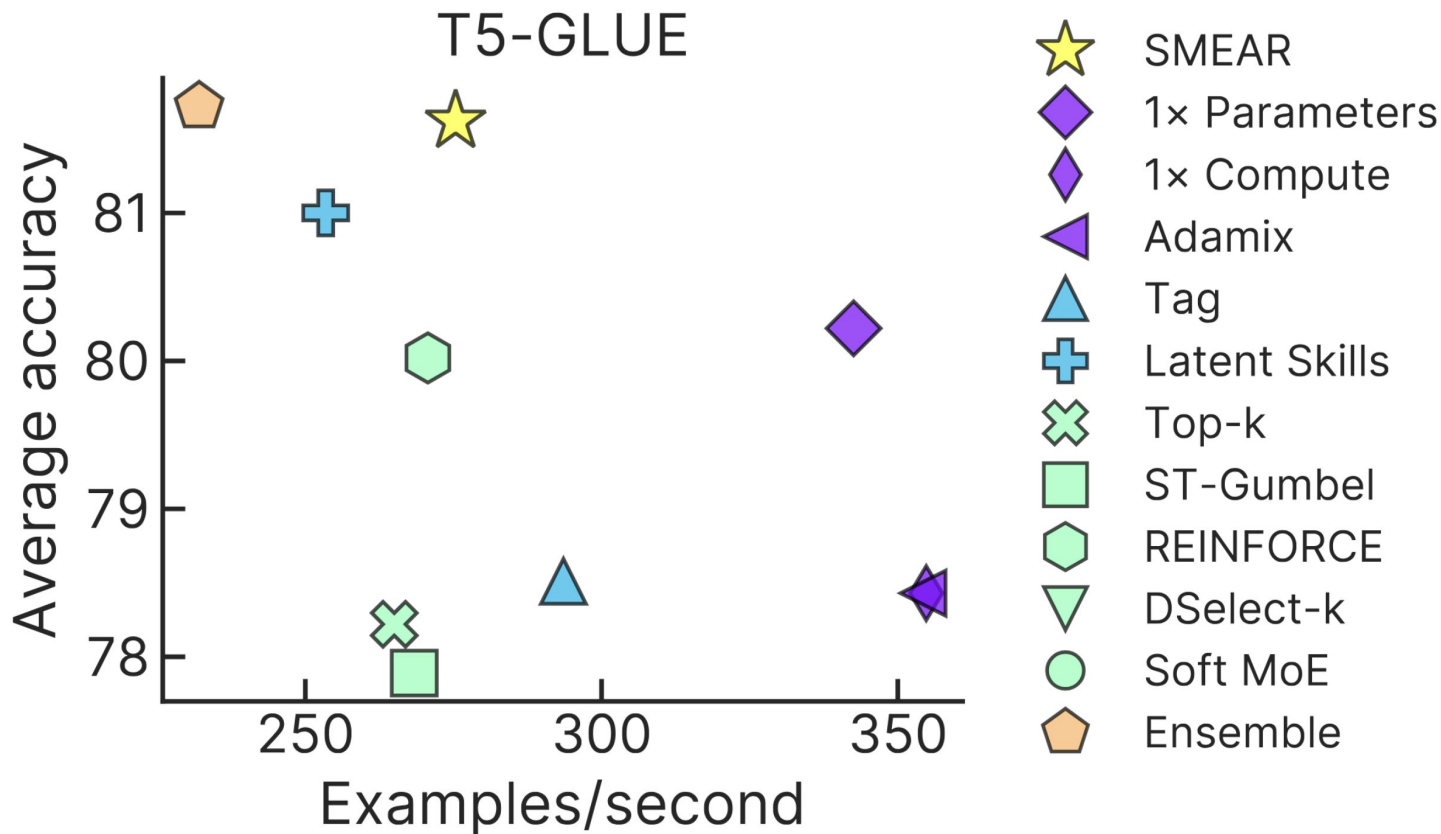
From "Learning to Route Among Specialized Experts for Zero-Shot Generalization" by Muqeeth et al.

Differentiable routing between specialist submodels with SMEAR

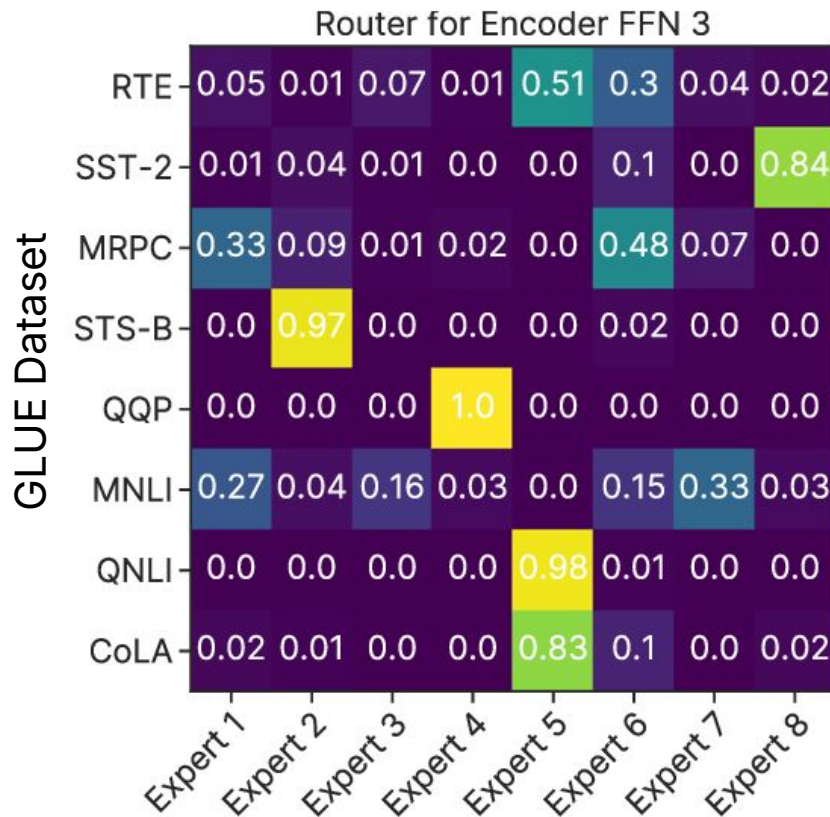


From "Soft Merging of Experts with Adaptive Routing" by Muqeeth et al.

SMEAR is pareto-optimal across different routing strategies

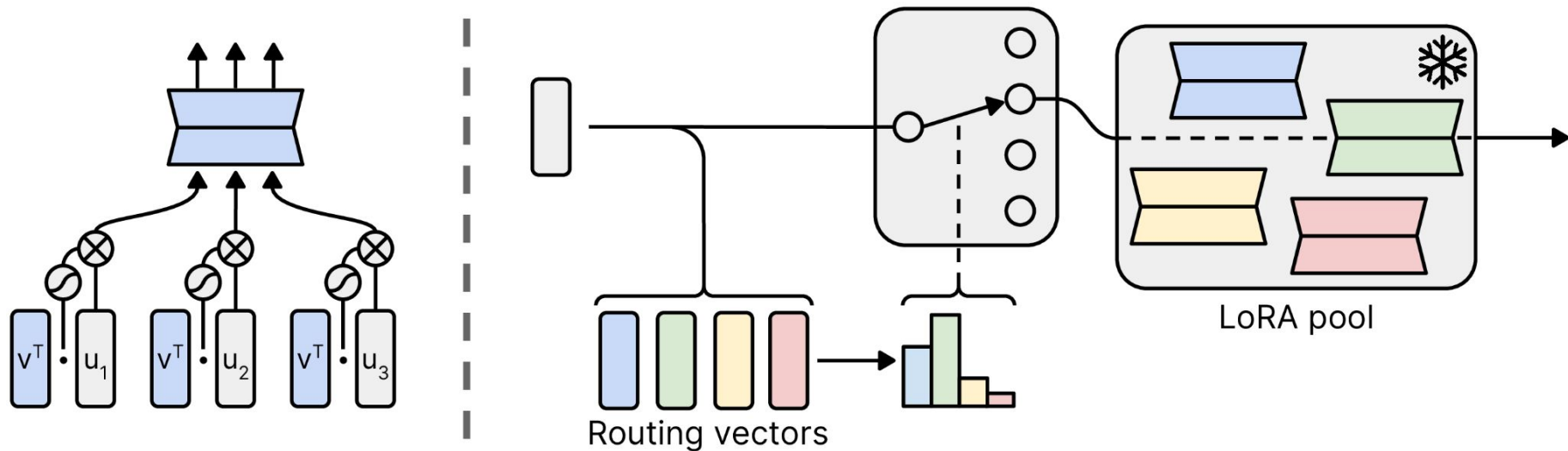


Experts specialize and are shared across different datasets

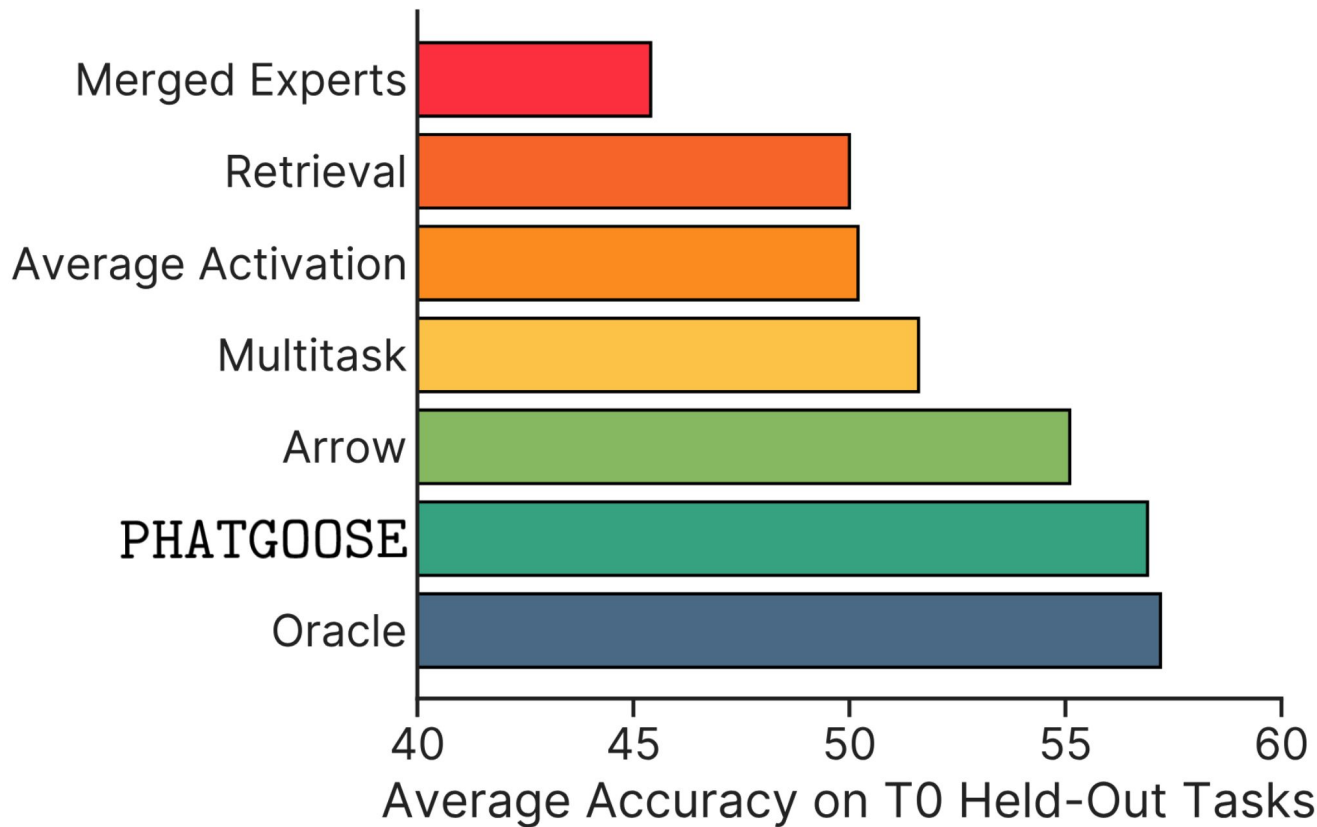


From "Soft Merging of Experts with Adaptive Routing" by Muqeeth et al.

Post-Hoc Adaptive Tokenwise Gating Over an Ocean of Specialized Experts



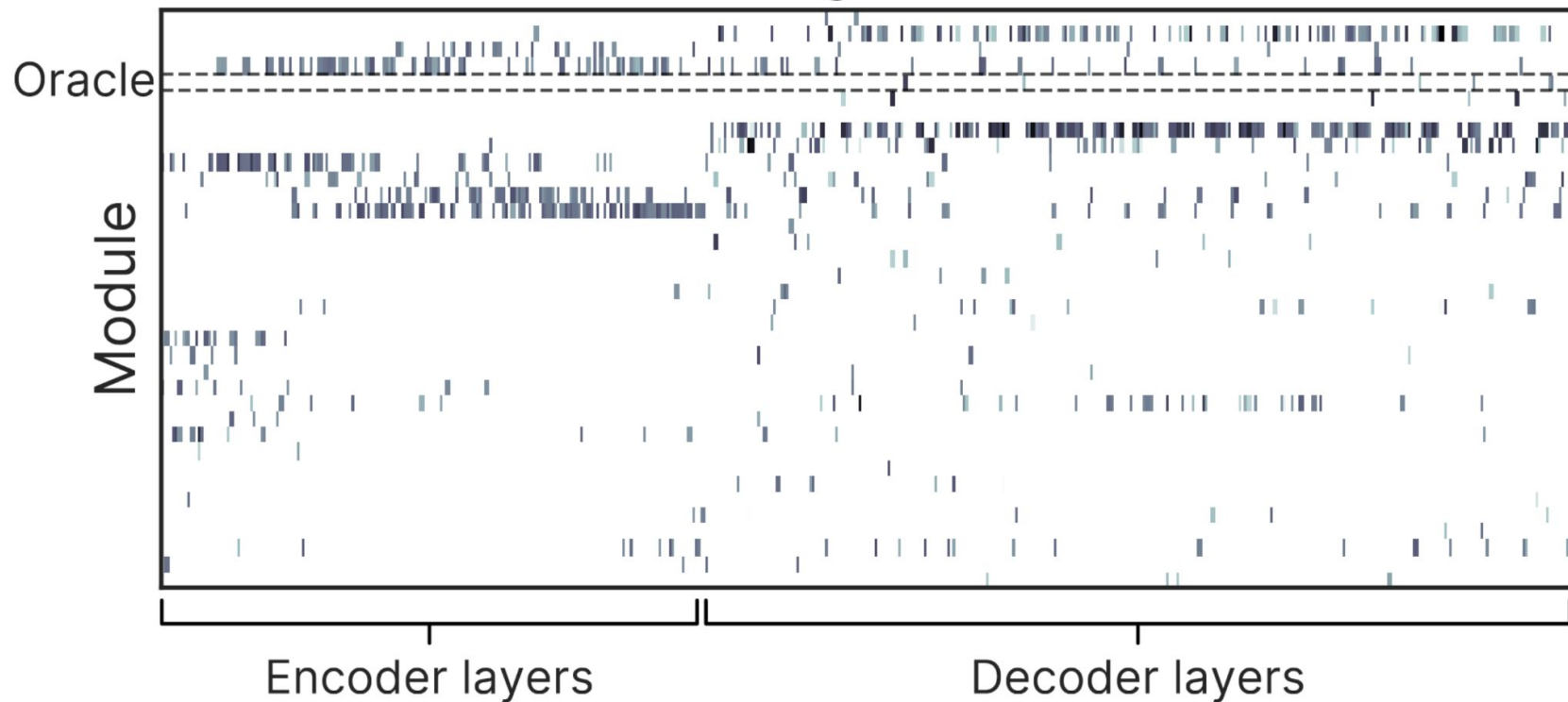
PHATGOOSE outperforms prior routing methods and multitask training



From "Learning to Route Among Specialized Experts for Zero-Shot Generalization" by Muqeeth et al.

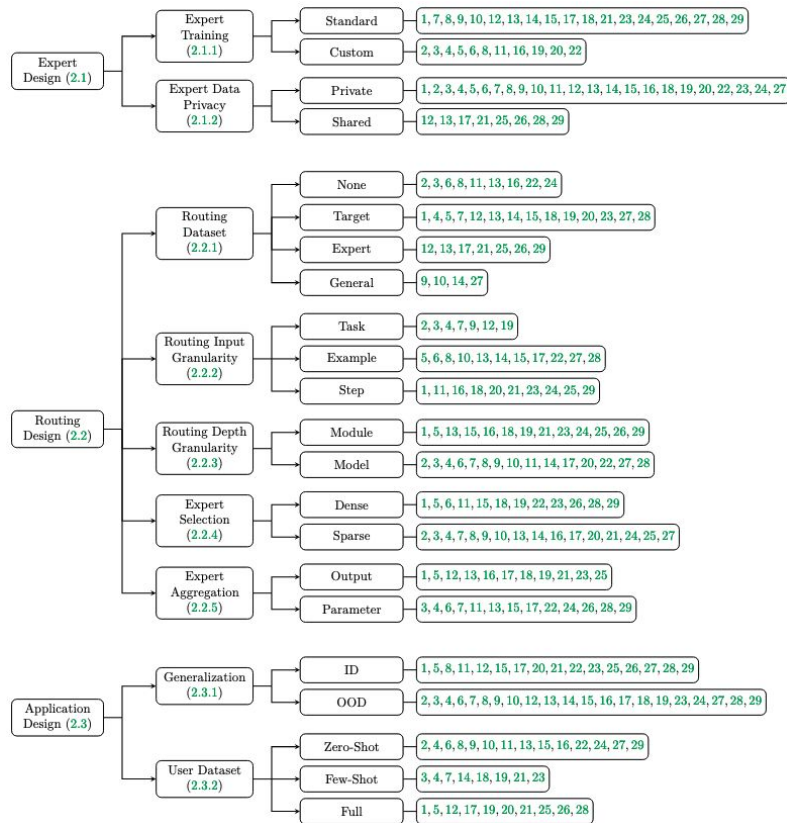
PHATGOOSE learns nontrivial routing strategies

CB Routing Distribution



From "Learning to Route Among Specialized Experts for Zero-Shot Generalization" by Muqeeth et al.

The MoErging design space



Thanks.

Please give me feedback:

<http://bit.ly/colin-talk-feedback>

craffel@gmail.com