

MAP: Model Merging with Amortized Pareto Fronts Using Limited Computation

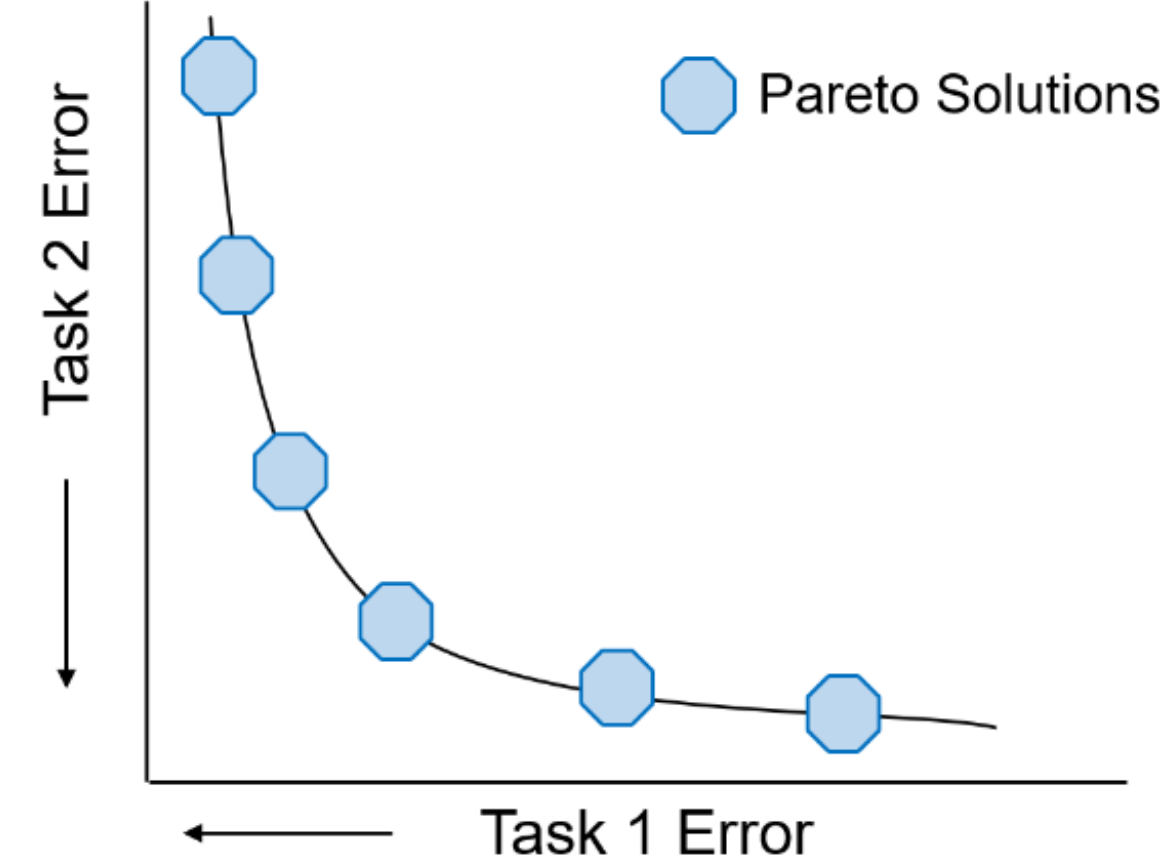
Lu Li^{1*}, Tianyu Zhang^{2,3*}, Zhiqi Bu^{4*}, Suyuchen Wang⁵, Huan He¹, Jie Fu⁶, Jiang Bian⁷, Yonghui Wu⁷, Yong Chen¹, Yoshua Bengio²

¹ University of Pennsylvania, ² MILA, ³ ServiceNow, ⁴ Amazon AGI, ⁵ Université de Montréal, ⁶ HKUST, ⁷ University of Florida *Equal contribution

Motivation

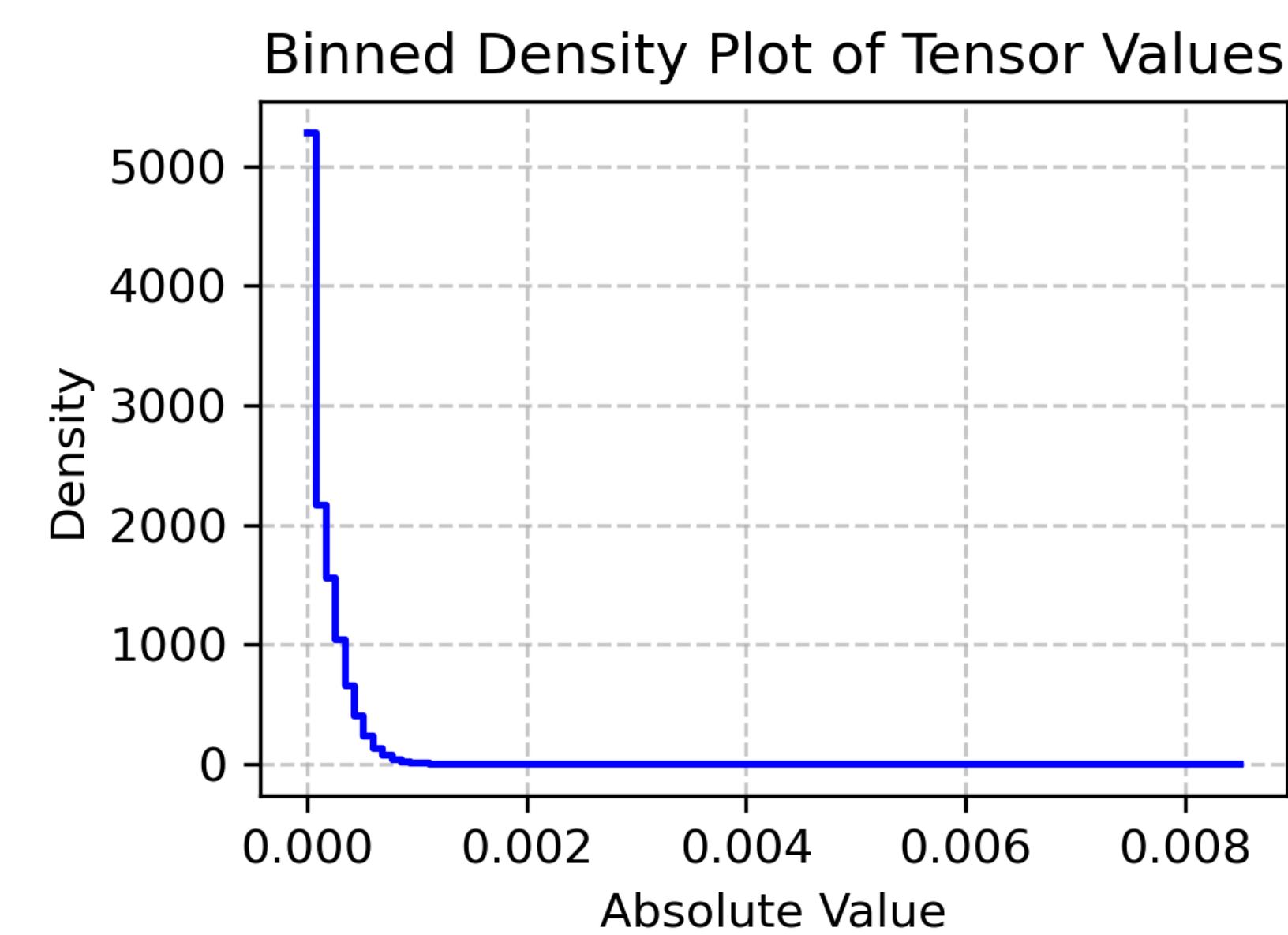
- **Task-vector based model merging** combines multiple single-task models into a single merged model that is capable of multi-task learning, without the need for sharing training data or any further training.
- The scaling coefficient for each task determines the relative performance of the merged model on that task

$$\theta_m = \theta_{pre} + \sum_i c_i v_i$$
- However, practitioners might have different preferences for the tasks, leading to trade-offs in how to merge the models. A set of Pareto optimal solutions is preferable.
- However, computing a Pareto front using conventional method involves many inferences and time-consuming



Key observation

- Fine-tuned models tend to converge near the pre-trained model in parameter space.

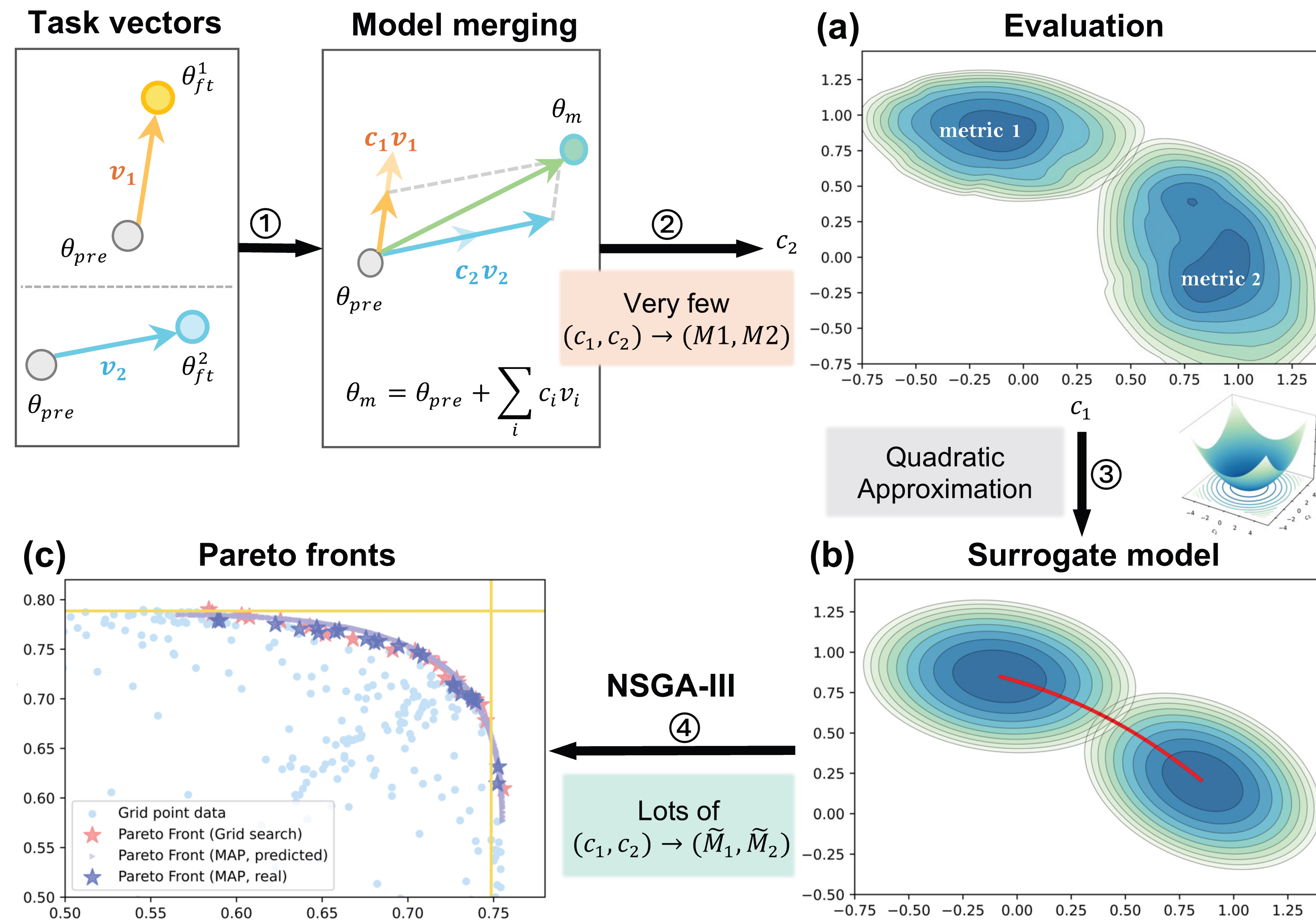


This motivates us to use **second-order Taylor expansion to approximate the evaluation metrics** for each task given a merged model

Metric	SUN397	Cars	DTD	SVHN
$\ \theta_{pre}\ _1$	1,270,487	1,270,487	1,270,487	1,270,487
$\ v_n\ _1$	21,055	20,127	13,621	19,349
$\ v_n\ _1 / \ \theta_{pre}\ _1 (\%)$	1.66%	1.58%	1.07%	1.52%

Method

- We propose **MAP**, a **computationally efficient method to identify Pareto fronts**, allowing trade-offs **without requiring additional training**.



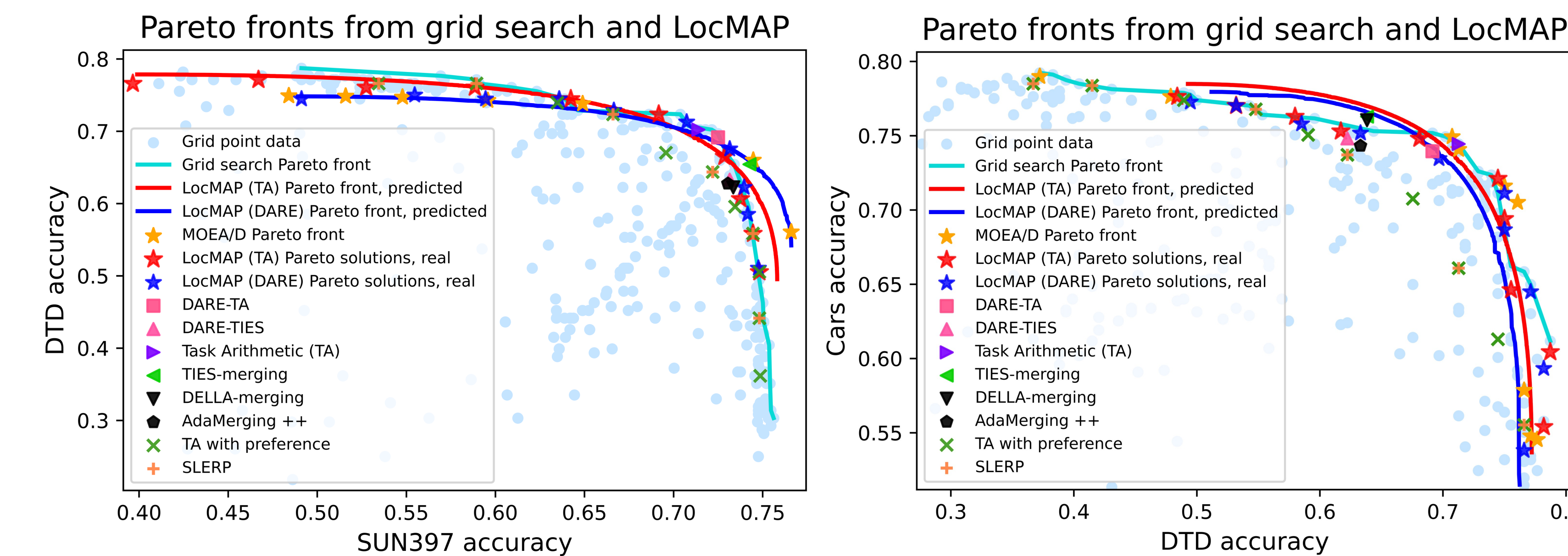
- Let $M_i(\theta_m(c)) = M_i(\theta_{pre} + \sum_i c_i v_i)$ denote the evaluation metrics on task i
- $M_i(c) = M_i(\theta_m(c))$

$$\begin{aligned}
 &= M_i(\theta_{pre}) + \nabla M_i(\theta_{pre})^\top (\theta_m(c) - \theta_{pre}) + \text{2nd order Taylor expansion around } \theta_{pre} \\
 &\quad + \frac{1}{2} (\theta_m(c) - \theta_{pre})^\top H_n(\theta_{pre}) (\theta_m(c) - \theta_{pre}) + R_i(\theta_m(c) - \theta_{pre}) \\
 &\approx \underbrace{M_i(\theta_{pre})}_{\in \mathbb{R}} + \underbrace{\nabla M_i(\theta_{pre})^\top}_{\in \mathbb{R}^{1 \times p}} \underbrace{v_i}_{\in \mathbb{R}^{p \times 1}} \underbrace{c_i}_{\in \mathbb{R}^{N \times 1}} + \frac{1}{2} \underbrace{c_i^\top}_{\in \mathbb{R}^{1 \times N}} \underbrace{v_i^\top}_{\in \mathbb{R}^{N \times p}} \underbrace{H_i(\theta_{pre})}_{\in \mathbb{R}^{p \times p}} \underbrace{v_i}_{\in \mathbb{R}^{p \times N}} \underbrace{c_i}_{\in \mathbb{R}^{N \times 1}} \\
 &\quad \widetilde{M}_i(c; A_i, b_i, e_i) \equiv e_i + b_i^\top c + \frac{1}{2} c^\top A_i c
 \end{aligned}$$

- Because of symmetry of A, the number of variables need to estimate is $\frac{(N+1)(N+2)}{2}$
- Since we already have the c and $M_i(c)$, it can be formulated as a linear regression, with $X = (c_1^2, c_2^2, \dots, c_T^2, c_1 c_2, c_1 c_3, \dots, c_{T-1} c_T, c_1, c_2, \dots, c_T, 1)$ and $y = M_i(\theta_m(c))$
- **Closed form solution:** $(X^\top X)^{-1} X^\top y$

Results

MAP is able to find **diverse Pareto front**, not captured by other single merged model baselines.



For higher dimensions, where we cannot visualize the Pareto front, the **win rate** is used to measure how often MAP outperforms other methods in terms of Pareto front solutions across tasks.

$$\text{Win Rate} = \frac{1}{K^2 N} \sum_{i=1}^K \sum_{j=1}^K \sum_{n=1}^N \mathbb{I} [M_n(\theta(c_i^{\text{LocMAP}})) > M_n(\theta(c_j^{\text{baseline}}))]$$

N	# c (direct search)	# c per dim	# c (MAP)	Win rate (MAP)	R^2 (MAP)
2	200	14.14	30	49.81% (± 0.30)	0.953 (± 0.018)
3	300	6.69	50	46.90% (± 0.71)	0.980 (± 0.003)

N	# c (direct search)	# c per dim	# c (MAP)	Win rate (MAP)	R^2 (MAP)
4	300	4.16	60	50.67% (± 2.44)	0.984 (± 0.004)
5	500	3.47	85	53.00% (± 1.88)	0.941 (± 0.019)
6	500	2.82	100	60.71% (± 1.34)	0.941 (± 0.030)
7	1000	2.68	140	63.42% (± 1.91)	0.891 (± 0.024)
8	1000	2.37	250	65.58% (± 0.94)	0.868 (± 0.028)

More Info

