



Portable Reward Tuning: Towards Reusable Fine-Tuning across Different Pretrained Models

Daiki Chijiwa*, [Taku Hasegawa*](#), Kyosuke Nishida, Kuniko Saito, Susumu Takeuchi

* equal contribution

2025年 7月 9日

ICML 2025 accepted

長谷川 拓 (プレゼンター)

■ 経歴

- 大阪府立大 (中退) 進化計算
- 大阪府立大 (修士・博士・PD) 最適化アルゴリズムと機械学習
- Freiburg 大学 (Invited Researcher) NNのアーキテクチャサーチ

■ 現所属：NTT人間情報研究所 視覚言語基盤モデル構築・知識転移



千々和 大輝

■ 経歴

- 旧・東京工業大学 学士 (数学) 代数幾何学・ホッジ理論
- 東京大学 修士 (数理科学) 同上

■ 現所属：NTT コンピュータ&データサイエンス研究所 深層学習・言語モデルの研究



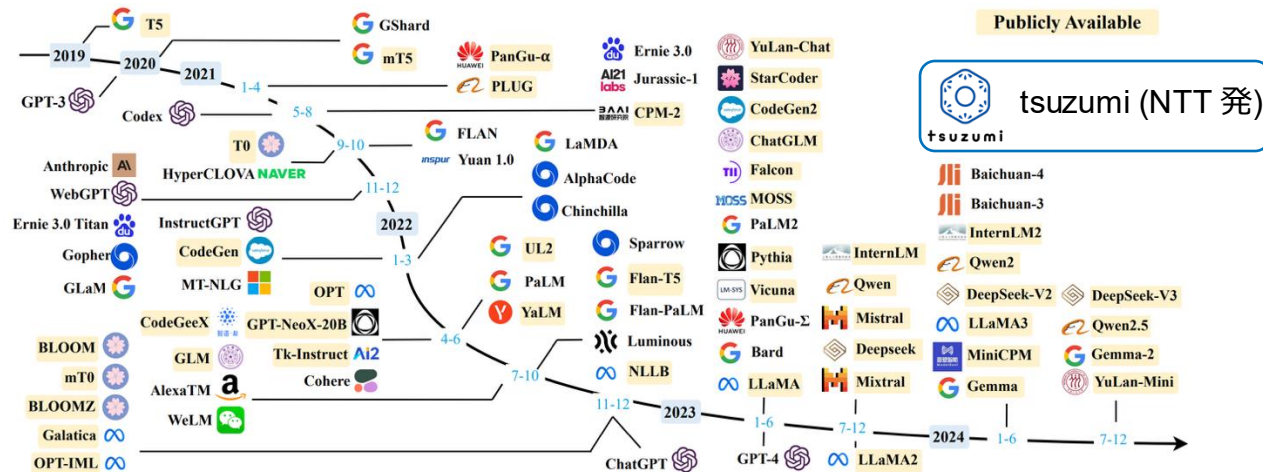
目次

1. 背景
2. 関連研究
3. 提案手法
4. 実験
5. まとめ

背景

基盤モデルは、言語・画像・音声など多様な領域で活用が拡大

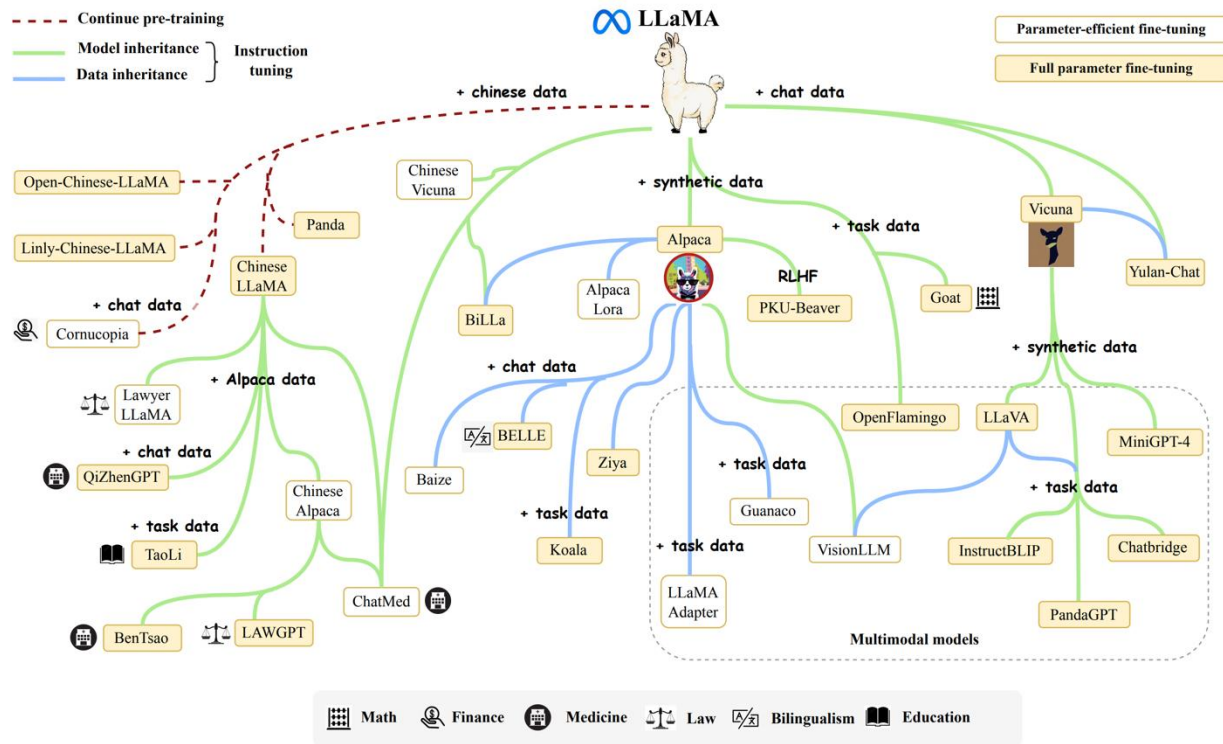
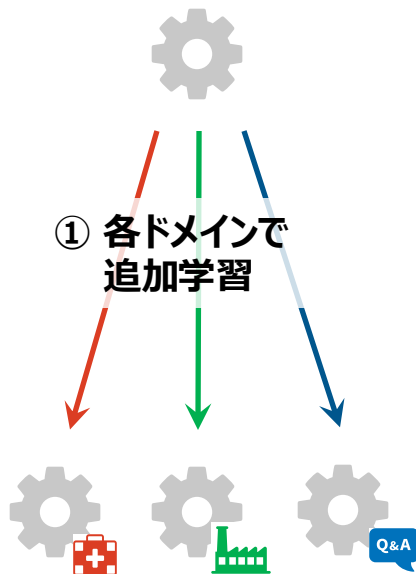
- 特に大規模言語モデル（LLM）は近年、急速に発展
- 研究・産業の両面で不可欠な技術に



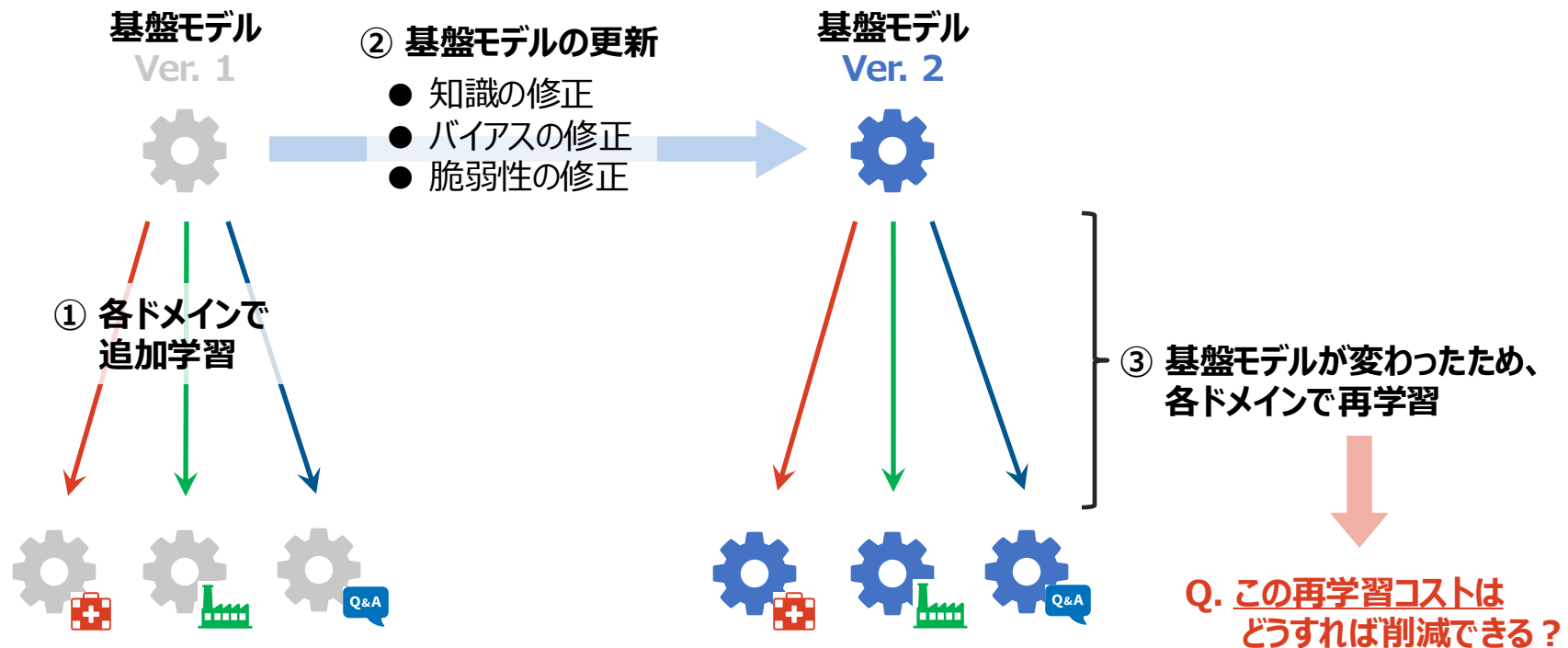
特化モデルの学習

基盤モデル

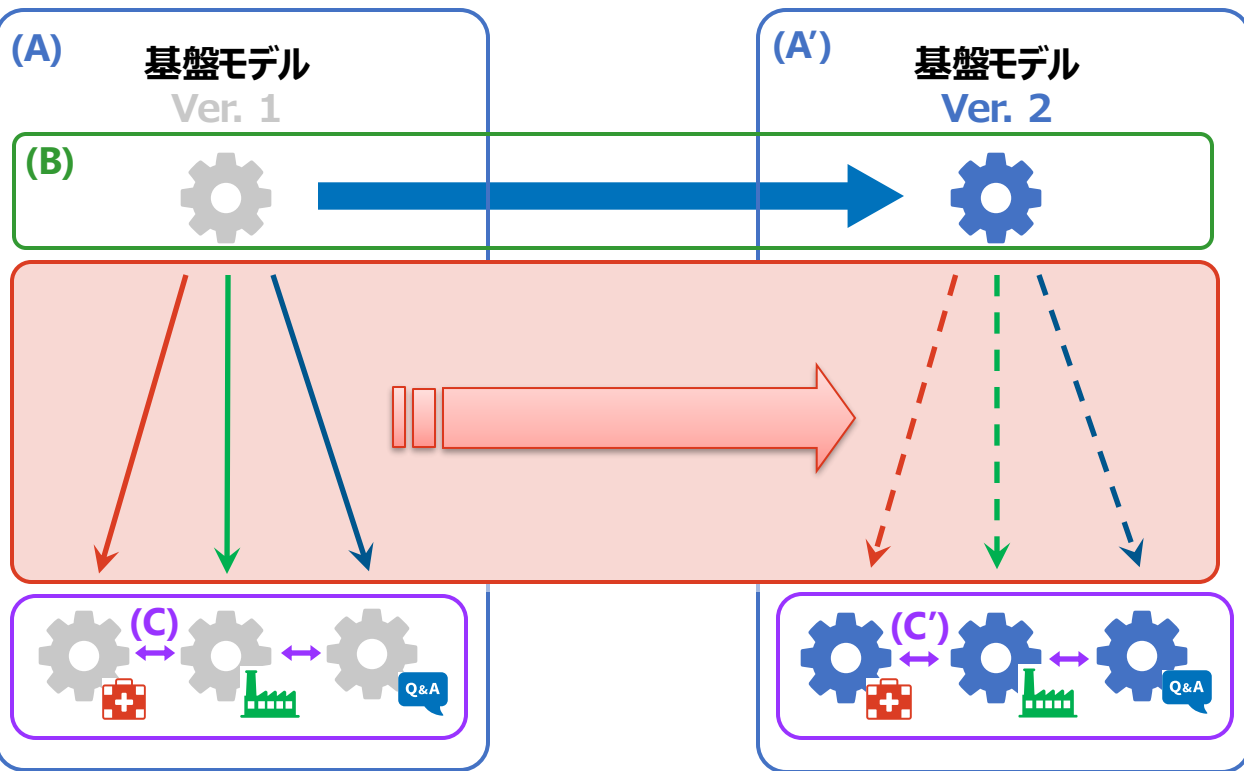
Ver. 1



モデルアップデート



モデルアップデート



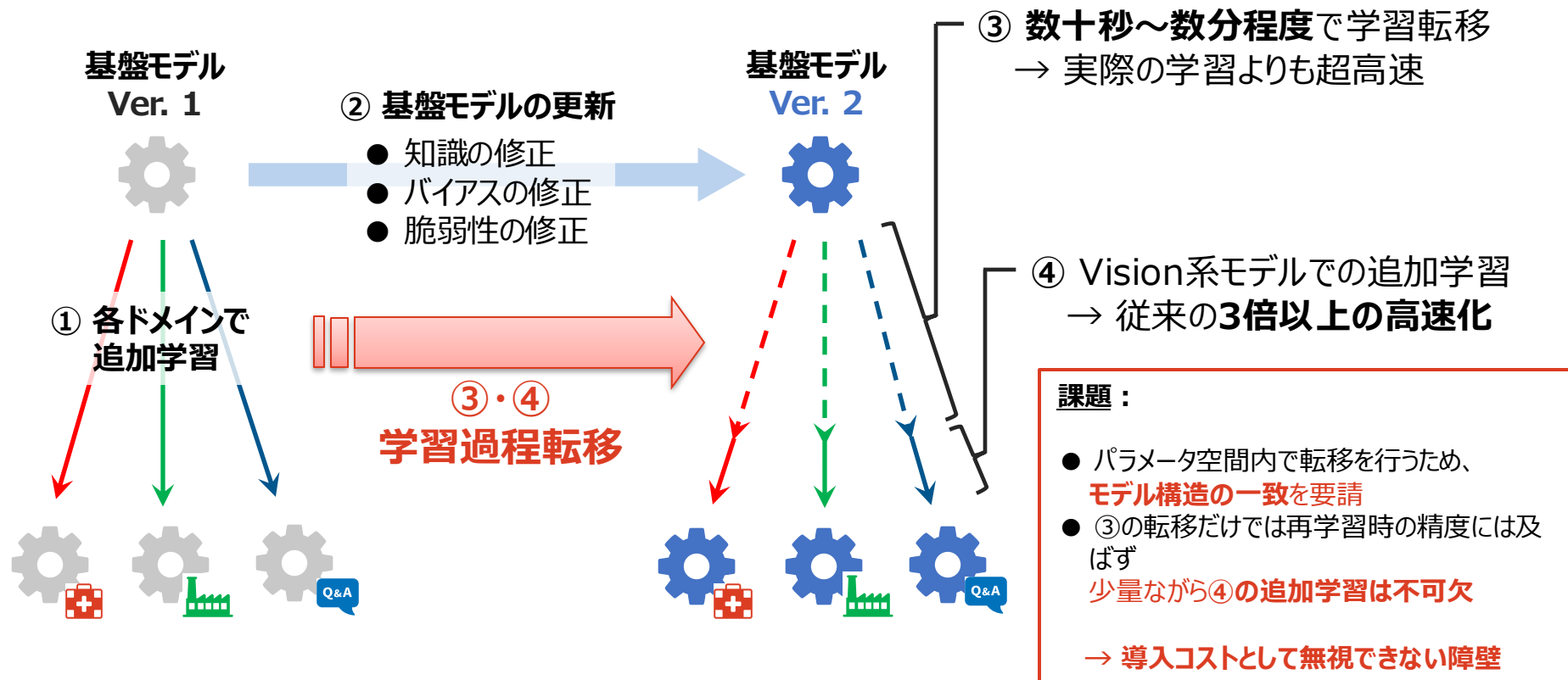
(A) 転移学習、ドメイン適合
→ 各ドメインの学習を効率化

(B) 継続学習、再学習、etc.
→ 基盤モデルの性能を改善

(C) モデルマージ、タスク算術
→ 同じ基盤モデルをベースとする特化モデルが複数あるときに、それらの演算により新たなモデルを作成

学習転移

異なる基盤モデル間で学習結果を転移させ、各ドメインでの再学習をコスト削減する新たな枠組み



関連研究

報酬最大化としての Fine-Tuning

分類モデル: $\pi_{\theta}(y|x) := \text{softmax}(f(x; \theta))$, 入力 $x \in \mathcal{X}$, 出力 $y \in \mathcal{Y}$

例)

画像分類: x は画像、 y は対応するラベル、 $f(x; \theta)$ は CNN や ViT

言語生成: x はトークン列 $t_1, \dots, t_k$ 、 y は次のトークン t_{k+1} 、 $f(x; \theta)$ は decoder-only Transformer

- 従来の Fine-Tuning (FT) は、次のような「KL制約つき報酬最大化」を行っているといみなせる: 基盤モデル $\pi_{pt}(y|x)$
FTモデル $\pi_{ft}(y|x)$

$$\max_{\pi(\cdot|x)} \mathbb{E}_{y \sim \pi(y|x)} \left[\log \left(\frac{\pi_{ft}(y|x)}{\pi_{pt}(y|x)} \right) \right] - D_{KL}(\pi(\cdot|x) \parallel \pi_{pt}(\cdot|x)) \quad \leftarrow \text{「基盤モデルから離れすぎない」ことを示す正則化項}$$

↑ 期待値の中身は
タスク報酬を表す、暗黙的な報酬関数
とみなせる。



解析解

$$\pi(y|x) = \frac{1}{Z(x)} \pi_{pt}(y|x) \exp \left(\log \left(\frac{\pi_{ft}(y|x)}{\pi_{pt}(y|x)} \right) \right) = \pi_{ft}(y|x)$$

最適化問題を解くと
FT モデルが復元される

Emulated Fine-tuning

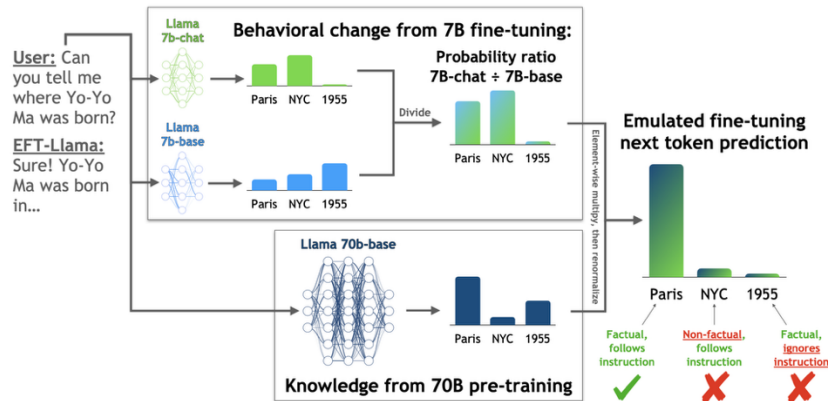
An Emulator for Fine-Tuning Large Language Models using Small Language Models,
Mitchell et al., ICLR 2024

先ほどの報酬最大化としての Fine-Tuning の定式化において、異なる事前学習 $\tilde{\pi}_{pt}(\cdot | x)$ を与えると以下の解が得られることに着目

$$\begin{aligned} \operatorname{argmax}_{\pi(\cdot|x)} \mathbb{E}_{y \sim \pi(y|x)} \left[\log \left(\frac{\pi_{ft}(y|x)}{\pi_{pt}(y|x)} \right) \right] - D_{KL}(\pi(\cdot|x) \parallel \tilde{\pi}_{pt}(\cdot|x)) \\ = \frac{1}{\tilde{Z}(x)} \tilde{\pi}_{pt}(y|x) \frac{\pi_{ft}(y|x)}{\pi_{pt}(y|x)} \end{aligned}$$

Emulated Fine-tuning (EFT):

基盤モデルとFTモデルの関係性から、推論時に新しい基盤モデルに能力を付加することが可能



特徴: 追加でタスク特化の学習が不要

課題: 推論時に3つのモデルを

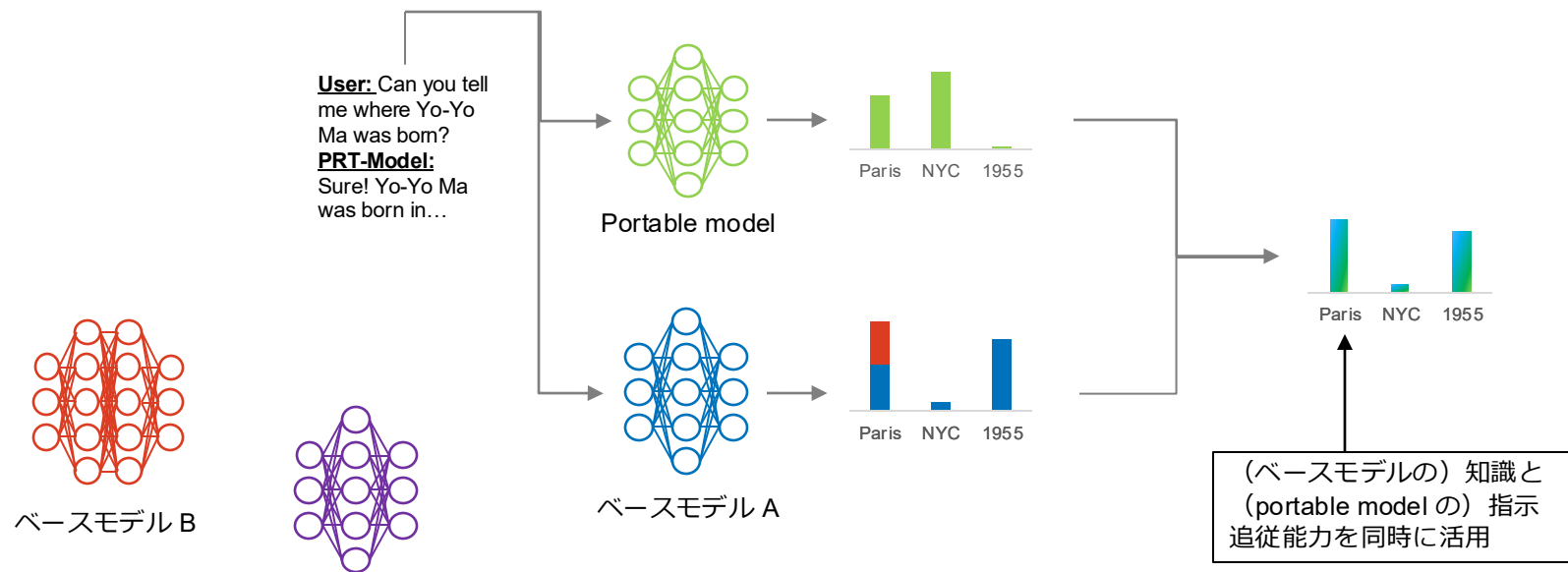
同時に動かす必要

→ 推論コスト増

提案手法

Portable Reward Tuning

【アイデア】 使い回し可能なモデルを学習する



Portable Reward Tuning (PRT)

 基盤モデル $\pi_{pt}(y|x)$  報酬モデル $r_\theta(x, y)$

 PRTモデル $\pi_\theta(x, y)$

「KL制約付き報酬最適化」問題における
報酬関数を直接モデル化&学習：

$$\max_{\pi(-|x)} \mathbb{E}_{a \sim \pi(y|x)} [r_\theta(x, y)] \quad \leftarrow \text{タスク特化の報酬関数をモデル化}$$

$$- \lambda D_{\text{KL}}(\pi(-|x) \parallel \pi_{pt}(-|x)).$$

その解として得られる次のモデルの学習を通して、 θ を学習：

$$\pi_\theta(y|x) = \frac{1}{Z_\theta(x)} \pi_{pt}(y|x) \exp\left(\frac{1}{\lambda} r_\theta(x, y)\right),$$

ロス関数 $\mathcal{L}(\mathbf{p}, y^*) := \text{CE}(\mathbf{p}, y^*) := -\log \pi_\theta(y^*|x)$ とすると

$$\arg \min_{\theta} \frac{1}{|S|} \sum_{(x, y^*) \in S} \mathcal{L}(\mathbf{p}, y^*) \quad (7)$$

$$= \arg \max_{\theta} \sum_{(x, y^*) \in S} \log \pi_\theta(y^*|x)$$

$$= \arg \max_{\theta} \sum_{(x, y^*) \in S} r_\theta(x, y^*) - V_\theta(x), \quad (8)$$

where $V_\theta(x) := \lambda \log Z_\theta(x)$

この最適化問題を通して r_θ は何を学習している？

→ イェンゼンの不等式

$$r_\theta(x, y^*) - V_\theta(x)$$

$$= r_\theta(x, y^*) - \lambda \log \mathbb{E}_{y \sim \pi_{pt}(y|x)} \exp(r_\theta(y|x)/\lambda)$$

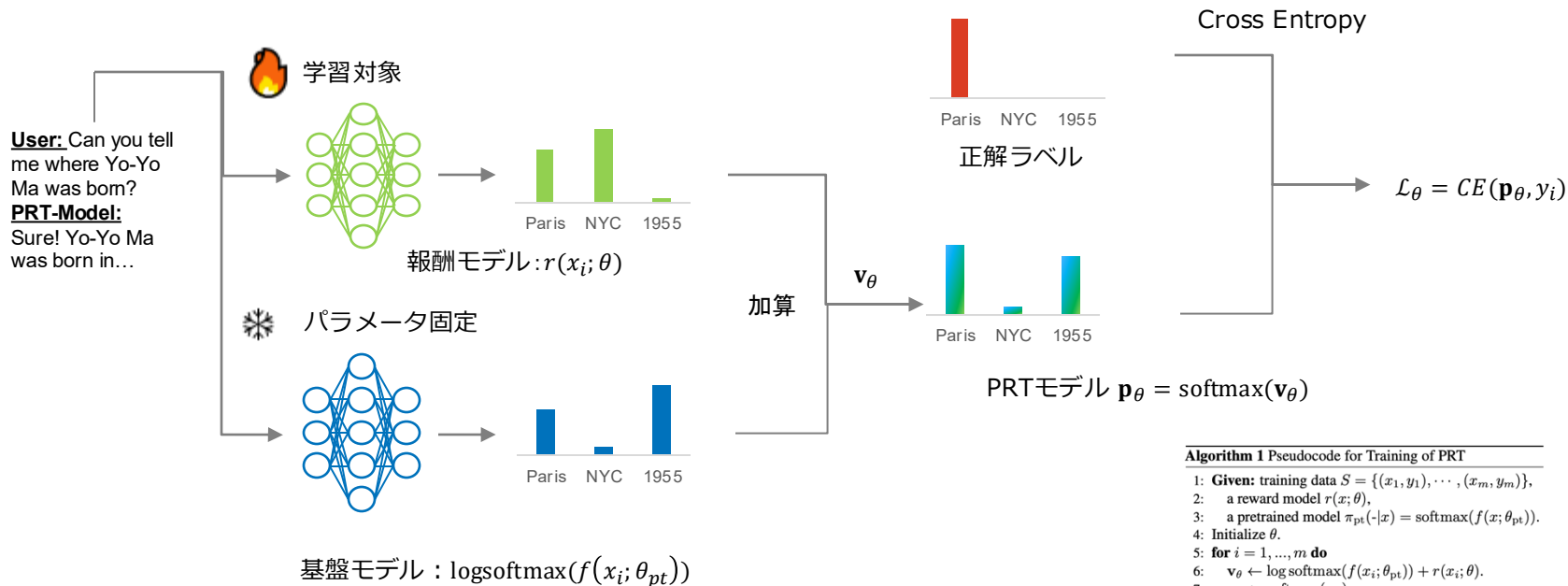
$$\leq r_\theta(x, y^*) - \lambda \mathbb{E}_{y \sim \pi_{pt}(y|x)} \log \exp(r_\theta(y|x)/\lambda)$$

$$= r_\theta(x, y^*) - \mathbb{E}_{y \sim \pi_{pt}(y|x)} r_\theta(y|x). \quad (9)$$

正解 y^* の報酬を
高くする

事前学習モデルからサンプルされる
他のラベルの平均報酬を低くする

具体的にどう実装するか？



Algorithm 1 Pseudocode for Training of PRT

- 1: **Given:** training data $S = \{(x_1, y_1), \dots, (x_m, y_m)\}$,
- 2: a reward model $r(x; \theta)$,
- 3: a pretrained model $\pi_{pt}(\cdot|x) = \text{softmax}(f(x; \theta_{pt}))$.
- 4: Initialize θ .
- 5: **for** $i = 1, \dots, m$ **do**
- 6: $\mathbf{v}_\theta \leftarrow \text{log softmax}(f(x_i; \theta_{pt})) + r(x_i; \theta)$.
- 7: $\mathbf{p}_\theta \leftarrow \text{softmax}(\mathbf{v}_\theta)$
- 8: $\mathcal{L}_\theta \leftarrow CE(\mathbf{p}_\theta, y_i)$
- 9: Update θ with the gradient of $\mathcal{L}_\theta$.
- 10: **end for**
- 11: **return** θ .

【疑問】

- PRT で学習したモデルってちゃんと学習できる？
- 通常の FT で学習したモデルと比べて能力に差がある？

【命題】

FT で学習したモデルに対して、精度面で 1 対 1 対応するモデルが存在する

$\{\pi_{\text{ft}}(y|x) : \text{fine-tuned models}\}$

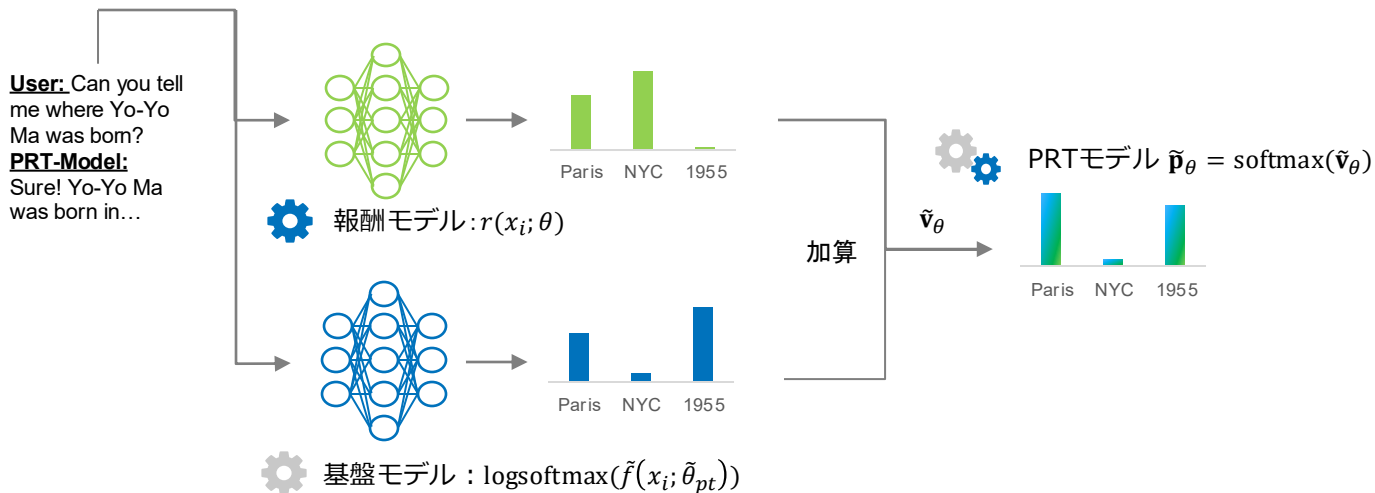
$\rightarrow \{r(y|x) : \text{rewards satisfying } \mathbb{E}_{y \sim \pi_{\text{pt}}(y|x)}[e^{r(x,y)}] = 1\}$

where $\pi_{\text{ft}}(y|x)$ is mapped to the implicit reward $\log(\pi_{\text{ft}}(y|x)/\pi_{\text{pt}}(y|x))$.

ポータブルモデルを使った推論

$$\pi_{\theta}(y|x) = \frac{1}{Z_{\theta}(x)} \pi_{\text{pt}}(y|x) \exp\left(\frac{1}{\lambda} r_{\theta}(x, y)\right)$$

$$\tilde{\pi}_{\theta}(y|x) = \frac{1}{\tilde{Z}_{\theta}(x)} \tilde{\pi}_{\text{pt}}(y|x) \exp\left(\frac{1}{\lambda} r_{\theta}(x, y)\right)$$



Algorithm 2 Pseudocode for Inference of PRT

- 1: **Given:** an input x , the trained reward $r(x; \theta)$,
- 2: a pretrained model $\tilde{\pi}_{\text{pt}}(y|x) = \text{softmax}(\tilde{f}(x; \tilde{\theta}_{\text{pt}}))$,
- 3: $\tilde{\mathbf{v}} \leftarrow \log \text{softmax}(\tilde{f}(x; \tilde{\theta}_{\text{pt}})) + r(x; \theta)$.
- 4: $\tilde{\mathbf{p}} \leftarrow \text{softmax}(\tilde{\mathbf{v}})$
- 5: **return** $\tilde{\mathbf{p}}$ as the output probability conditioned by x .

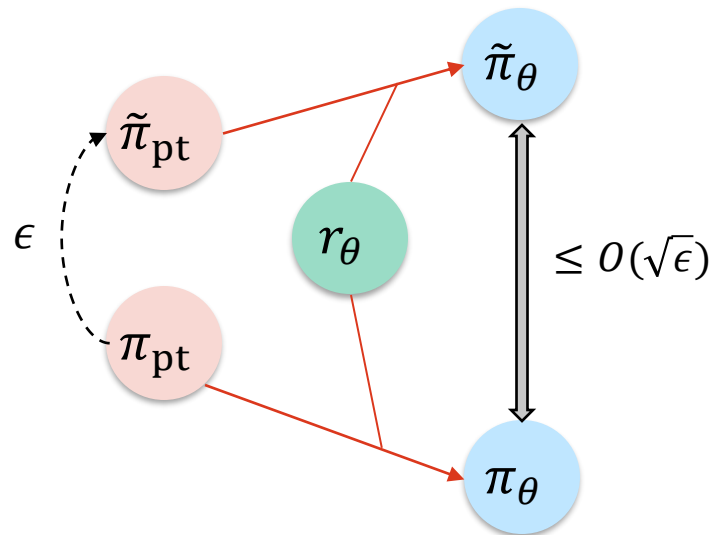
【疑問】

- 推論時に基盤モデルを入れ替えて大丈夫？
- どんな基盤モデルでも入れ替ええるの？

【命題】

Proposition 3.2. Suppose that $\tilde{\pi}_{\text{pt}}(y|x)$ is close to $\pi_{\text{pt}}(y|x)$, i.e., $D_{\text{KL}}(\pi_{\text{pt}}(-|x) \parallel \tilde{\pi}_{\text{pt}}(-|x)) \leq \epsilon$. Additionally, we assume that the maximum and mean value ratio of the exponential reward, i.e., $\max_y \exp r_{\theta}(x, y) / \mathbb{E}_y \exp r_{\theta}(x, y)$, is bounded by some constant C . Then, the PRT models $\pi_{\theta}(y|x)$ and $\tilde{\pi}_{\theta}(y|x)$ are also close as distributions:

$$D_{\text{KL}}(\pi_{\theta}(y|x) \parallel \tilde{\pi}_{\theta}(y|x)) \leq O(\sqrt{\epsilon}).$$



実験

実験設定：画像分類

実験の流れ

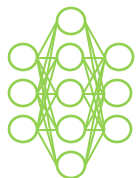
【訓練フェーズ】

報酬モデルのパラメータ初期値 = 基盤モデルのパラメータ

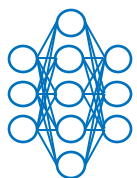
基盤モデル：
ViT-B-16 train on
LAION-400M



各タスクでSFT



報酬モデル: $r(x_i; \theta)$



基盤モデル: π_{pt}

【推論フェーズ】

学習済み報酬モデル

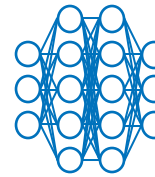
Target model

パターン A) 学習時と同じアーキテクチャ、異なるデータで学習したモデル



ViT-B-16 train on
LAION-2B

パターン B) 学習時と同じデータ、異なるアーキテクチャで学習したモデル



ViT-L-14 train on
LAION-400M

パターン C) 学習時と異なるアーキテクチャ、データで学習したモデル

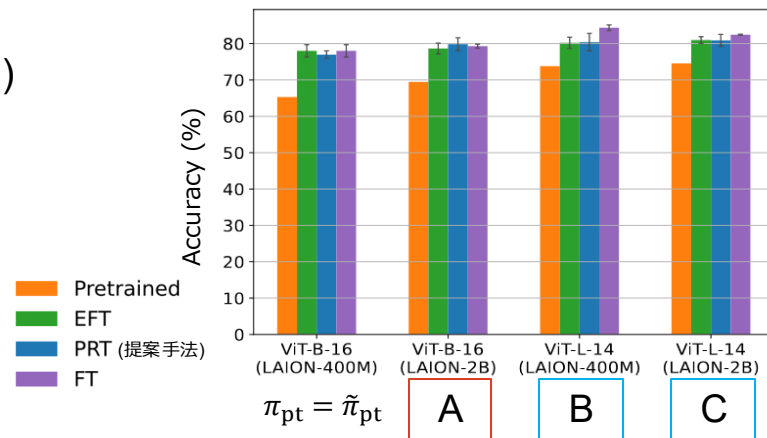


ViT-L-14 train on
LAION-2B

画像分類タスクでPRTは効果があるか？

- 学習時の基盤モデル π_{pt} : ViT-B16 (train on LAION-400M)
- 推論時の基盤モデル $\tilde{\pi}_{pt}$:
 - A) 学習時と同じアーキテクチャ、異なるデータで学習したモデル
 - B) 学習時と同じデータ、異なるアーキテクチャで学習したモデル
 - C) 学習時と異なるアーキテクチャ、データで学習したモデル

CUB データセット：鳥の画像分類



【結論】 画像分類タスクで PRT は SFT に近い性能を達成した

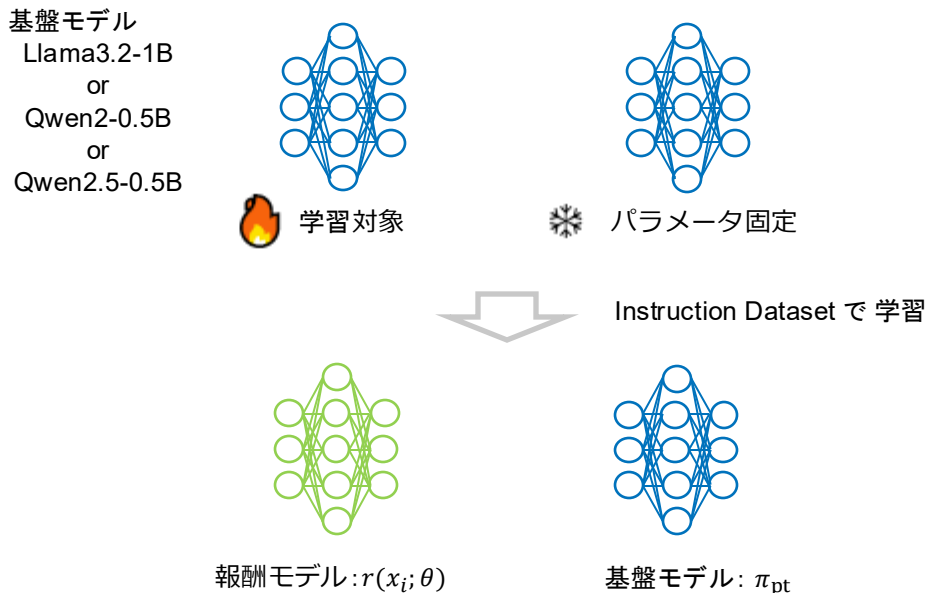
- 基盤モデルを入れ替えず ($\pi_{pt} = \tilde{\pi}_{pt}$) に推論した場合は SFT、EFTと同等の性能
- A) の学習データが異なる基盤モデルへの入れ替えの場合も SFT、EFTと同等の性能
- B)、C) のアーキテクチャが異なる基盤モデルの入れ替えの場合は SFT に僅かに及ばずだが、基盤モデル単体と比較してタスク性能は向上

実験設定：言語生成

実験の流れ

【訓練フェーズ】

報酬モデルのパラメータ初期値：基盤モデルのパラメータ



【推論フェーズ】

学習済み報酬モデル

Target model

パターン A) 学習時の基盤モデルと同じバージョン、大きいモデル

例)
Llama3.2-1B
or
Qwen2.5-0.5B

+

例)
Llama3.2-3B
or
Qwen2.5-1.5B
Qwen2.5-3B
⋮

パターン B) 学習時の基盤モデルより古いバージョン、大きいモデル

例)
Llama3.2-1B

+

例)
Llama3-8B
Llama3.1-8B

パターン C) 学習時の基盤モデルより新しいバージョン

例)
Qwen2-0.5B

+

例)
Qwen2.5-0.5B
Llama2.5-1.5BB

言語生成タスクで PRT は効果的か？

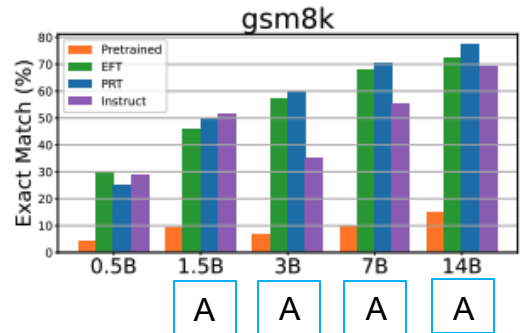
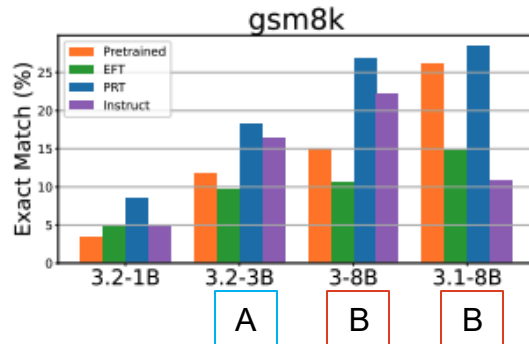
- 学習時の基盤モデル π_{pt} : Llama3.2-1B、Qwen2.5-0.5B
- 推論時の基盤モデル $\hat{\pi}_{pt}$:

- A) 学習時の基盤モデルと同じバージョン、大きいモデル
- B) 学習時の基盤モデルより古いバージョン、大きいモデル

【結論】言語生成タスクで PRT は Instruction tuning モデルに近い性能を達成した

- 基盤モデル単体と比較してPRT のタスク性能の方が高い
→ PRTの学習結果を活用できている
- PRTでは、サイズがより大きな基盤モデルの方がでより良いスコアを達成 → 基盤モデルの能力を活用できた
- バージョンが古いモデルでも有効性が確認できた

GSM8kデータセット: 算数タスク



基盤モデルのアップデートに対応可能か？

- 学習時の基盤モデル π_{pt} : Qwen2-0.5B

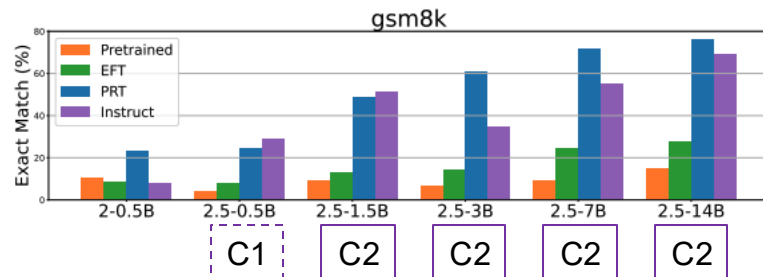
- 推論時の基盤モデル $\tilde{\pi}_{pt}$:

C1) 学習時の基盤モデルより新しいバージョン、サイズは同じ

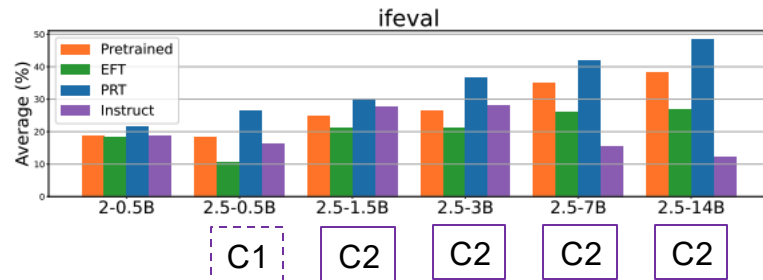
C2) 学習時の基盤モデルより新しいバージョン、大きいモデル

- 【結論】基盤モデルのアップデートのシナリオにおいても、PRT は有効性を示した
- 基盤モデル単体と比較して、PRT のタスク性能の方が高い
 - 指示追従タスクでも、基盤モデル単体と比較してタスク性能は向上

GSM8kデータセット: 算数タスク



ifeval データセット: 指示追従タスク



出力はどのようになっているか？

【調査方法】

“R=27. James has” に続く

- 基盤モデル $\tilde{\pi}_{pt}$ ← PRT による分布修正前
- PRTモデル $\tilde{\pi}_{\theta}$ ← PRT による分布修正後

のトークン予測確率を調べる

【モデル】

- 学習時の基盤モデル π_{pt} : Llama3.2-1B
- 推論時の基盤モデル $\tilde{\pi}_{pt}$: Llama3-8B

【結論】

ベースの基盤モデルは 33 が最も高い確率

PRTモデル $\pi_{\theta}(x, y)$ は 6 が最も高い確率

→ PRT モデルがstep by step で考えるように分布を修正している

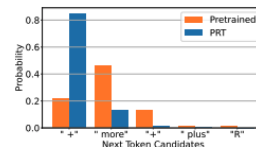
Prompt

Question: James has 6 more candies than Robert. John has twice as many candies as Robert. If John has 54 candies, how many more candies does John have than James?

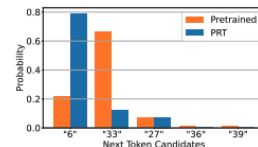
Answer:

Model Response

Let R be the number of candies Robert has. John has 2R candies. James has 6 + R candies. We know that John has 54 candies, so 2R = 54. R = 27. James has 6 + 27 = 33 candies. John has 54 - 33 = 21 more candies than James.



(a) Next Token Candidates Following "... John has 2R candies. James has 6"



(b) Next Token Candidates Following "... R = 27. James has"

提案手法の計算コストは？

【学習時の計算コスト】

基盤モデルは凍結:学習対象のパラメータ数はFTと同じ

- GPUメモリ使用量:FTと比べPRTは約 1.02倍
- 処理速度: FTと比べPRTは約 1.2~1.3倍
→ 学習時のオーバーヘッドはそこまで大きくない

【推論時の計算コスト】

推論時に計算するモデルの数:FT: 1つ, EFT: 3つ, PRT: 2つ
視覚モデルで

- FTと比較すると GPUメモリは約1.1~1.5倍増加、推論時間は約1.5~2倍増加した。
- EFTと比較すると GPUメモリは 7~18% 程度削減、推論時間は 24~32% 程度削減した

言語モデルでも

- EFTと比較すると推論時間は 7%~30%程度削減に成功した

学習時のコスト(視覚モデル)

Models		FT	PRT
ResNet-50	Peak GPU memory	12.84 GB	13.08 GB
	Average time per batch	38.26 $\pm$ 3.62 ms	46.36 $\pm$ 3.34 ms
ViT-B-16	Peak GPU memory	20.17 GB	20.58 GB
	Average time per batch	116.10 $\pm$ 0.43 ms	151.13 $\pm$ 0.48 ms

推論時のコスト(視覚モデル)

Source Models	Target Models		Target FT (Oracle)	EFT	PRT
ResNet-50	ViT-B-16	Peak GPU memory	1.46 GB	2.42 GB	2.18 GB
		Average time per batch	3.52 ± 0.17 ms	10.27 ± 0.22 ms	7.00 ± 0.12 ms
ResNet-50	ViT-L-14	Peak GPU memory	2.96 GB	3.44 GB	3.20 GB
		Average time per batch	5.85 ± 0.11 ms	12.52 ± 0.21 ms	9.43 ± 0.19 ms
ViT-B-16 (LAION-400M)	ViT-B-16 (LAION-2B)	Peak GPU memory	1.46 GB	2.29 GB	1.88 GB
		Average time per batch	3.52 ± 0.17 ms	9.18 ± 0.12 ms	6.33 ± 0.07 ms
ViT-B-16 (LAION-2B)	ViT-L-14 (LAION-2B)	Peak GPU memory	2.96 GB	3.79 GB	3.38 GB
		Average time per batch	5.85 ± 0.11 ms	11.73 ± 0.15 ms	8.89 ± 0.11 ms

推論時のコスト(言語モデル)

Source Models Target Models	Average Time per Token (ms)		
	Llama2-7B Llama2-13B	Llama3.2-1B Llama3-8B	Llama3.2-3B
Pretrained	26.0 $\pm$ 0.2	17.1 $\pm$ 1.1	
EFT	39.8 $\pm$ 0.1 ($\times$ 0.65)	24.4 $\pm$ 0.2 ($\times$ 0.7)	30.4 $\pm$ 0.0 ($\times$ 0.56)
PRT	27.8 $\pm$ 0.5 ($\times$ 0.93)	22.7 $\pm$ 1.8 ($\times$ 0.75)	23.1 $\pm$ 1.0 ($\times$ 0.74)

お気持ち：様々な基盤モデルで使い回し可能になるよう汎化性能を上げたい！

PAC-Bayes的視点からの解析

- $\mathcal{P}$: 事前分布 (pretrained modelの分布)
- $\mathcal{Q}_\theta$: 報酬モデルに基づくPRTモデルの事後分布

【命題】

データ分布 D に従ってサンプル $S = \{x_1, \dots, x_m\} \sim D_m$ を取得したとき、任意の $\delta \in (0,1)$ に対して、少なくとも $1 - \delta$ の確率で成立：

$$\mathbb{E}_{\pi \sim \mathcal{Q}_\theta} \mathbb{E}_{x \sim D} [l(x, \pi)] \leq \mathbb{E}_{\pi \sim \mathcal{Q}_\theta} \left[\frac{1}{m} \sum_i^m l(x_i, \pi) \right] + \sqrt{\frac{D_{\text{KL}}(\mathcal{Q}_\theta \parallel \mathcal{P}) + \log(\frac{1}{\delta}) + \frac{5}{2} \log m + 8}{2m - 1}},$$

$D_{\text{KL}}(\mathcal{Q}_\theta \parallel \mathcal{P})$ の項に着目

→ PRTモデル $\tilde{\pi}_\theta(y|x)$ が事前学習モデル $\tilde{\pi}_{\text{pt}}(y|x)$ に近いほど汎化性能が高い

【提案手法の拡張】

報酬分布のエントロピーを正則化項として導入

- メリット：汎化性能をさらに高めることが可能
- デメリット：学習の最適性には悪影響を与える可能性

$$\tilde{\mathcal{L}}(\theta) := \mathcal{L}(\theta) - \alpha \frac{1}{|S|} \sum_{(x,y) \in S} H(\rho_\theta(\cdot|x)),$$

- $\mathcal{L}(\theta)$: 元の損失 (クロスエントロピー)
- $H(\rho_\theta(\cdot|x)) := -\sum_{y \in Y} \rho_\theta(y|x) \log \rho_\theta(y|x)$: 報酬分布のエントロピー
- $\alpha \in R_{\geq 0}$: 正則化の強さを制御するハイパーパラメータ

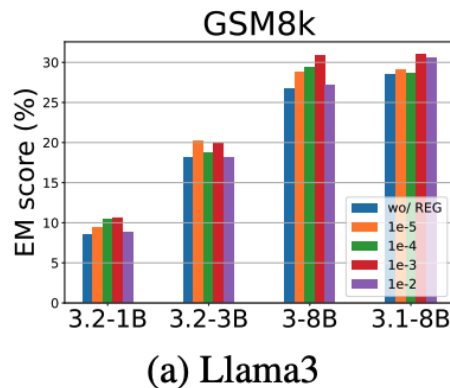
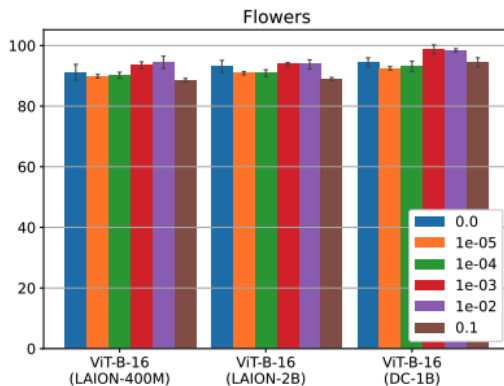
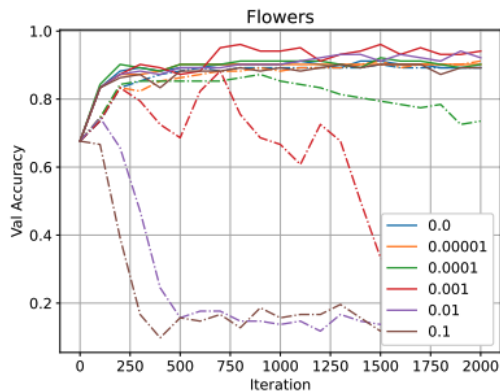
Entropy正則化の効果はあるか？

【結論】

- 正則化係数が大きすぎると学習中に reward model 自体の精度が低くなる
- 一方で、PRTの最終的な性能がそれに伴って悪化するわけではない
 - 報酬の精度に関わらずPRT の性能が保たれる・改善される = ただのアンサンブルではない

【直感的解釈】

- 正則化により、報酬モデルの出力分布が一様分布に近づく（＝クセが少なくなる）
 - 元のモデルの持ち味を活かしつつ、必要な分だけ補正をかけることができるようになる



おわりに

研究のスタートはモデルマージだった

- 特定のタスクに対する能力を獲得済みのモデル同士を組み合わせ、様々なタスクに使えるモデルを作りたい
 - 学習過程転移を使って実現できないか？

課題：アーキテクチャを揃える必要あり

- 異なる基盤モデルだとパーミュテーションも考慮する必要あり

結論！！出力分布のマージの方が柔軟性が高く、やりやすい！