

Multiple Choice Learning of Low-Rank Adapters for Language Modeling

Victor Letzelter^{*12}, Hugo Malard^{*2}, Mathieu Fontaine²
Gaël Richard², Slim Essid², Andrei Bursuc¹, Patrick Pérez³

^{*}Equal Contribution.

¹ Valeo.ai

² LTCI, Telecom Paris, Institut Polytechnique de Paris

³ Kyutai

Ambiguous tasks?



Image Captioning.

Ambiguous tasks?



{Pedestrians crossing a busy
downtown street.}

Image Captioning.

Ambiguous tasks?



{Pedestrians crossing a busy downtown street.}

{An ambulance drives through an intersection with its lights on.}

Image Captioning.

Ambiguous tasks?



{Pedestrians crossing a busy downtown street.}

{An ambulance drives through an intersection with its lights on.}

{Taxis and cars move through a crowded city street.}

Image Captioning.

Ambiguous tasks?



{Pedestrians crossing a busy downtown street.}

{An ambulance drives through an intersection with its lights on.}

{Taxis and cars move through a crowded city street.}

{A foggy afternoon in a large city with tall buildings lining the street.}

Image Captioning.

Ambiguous tasks?



{Pedestrians crossing a busy downtown street.}

{An ambulance drives through an intersection with its lights on.}

{Taxis and cars move through a crowded city street.}

{A foggy afternoon in a large city with tall buildings lining the street.}

Image Captioning.

👉 How to capture this ambiguity in Language models ?

(Causal) Language Modeling

- Sequence of tokens $x = (x_t)_{t=1}^T \in \mathcal{V}^T$
- Context vector embeddings $c = (c_t)_{t=1}^\tau \in \mathbb{R}^{\tau \times d}$

Training. “Next-token-prediction” with “teacher-forcing”

$$\mathcal{L}(\theta) = \mathbb{E}_{c,x} \left[- \sum_{t=1}^T \log p_\theta (x_t \mid x_{<t}, c) \right]$$

(Causal) Language Modeling

Decoding. Inference is performed in an autoregressive manner.

- ① Compute $p_\theta(x_1 \mid c)$ and select $\hat{x}_1$.
- ② For $t \geq 2$, compute $p_\theta(x_t \mid \hat{x}_{<t}, c)$ and select $\hat{x}_t$.

How to select the tokens ?

- Maximum A Posteriori (greedy, beam search¹, diverse beam search²).
- Sampling (top-k³, nucleus⁴).

¹B. T. Lowerre (1976). *The harpy speech recognition system*. [CMU](#).

²A. Vijayakumar et al. (2018). “Diverse beam search for improved description of complex scenes”. In: *AAAI*.

³A. Fan et al. (2018). “Hierarchical neural story generation”. In: *ACL*.

⁴A. Holtzman et al. (2020). “The curious case of neural text degeneration”. In: *ICLR*.

Motivations with a synthetic data example

Synthetic data example with Markov Chain generated data.

- Transition matrix $P_{i,j} = p(x_{t+1} = j \mid x_t = i)$
- GPT-like¹ model trained on such sequences

👉 Does the transformer recover the true transition matrix?

¹A. Radford et al. (2019). “Language models are unsupervised multitask learners”. In: *OpenAI blog*.

Motivations with a synthetic data example

Synthetic data example with Markov Chain generated data.

- Transition matrix $P_{i,j} = p(x_{t+1} = j \mid x_t = i)$
- GPT-like¹ model trained on such sequences

👉 Does the transformer recover the true transition matrix?

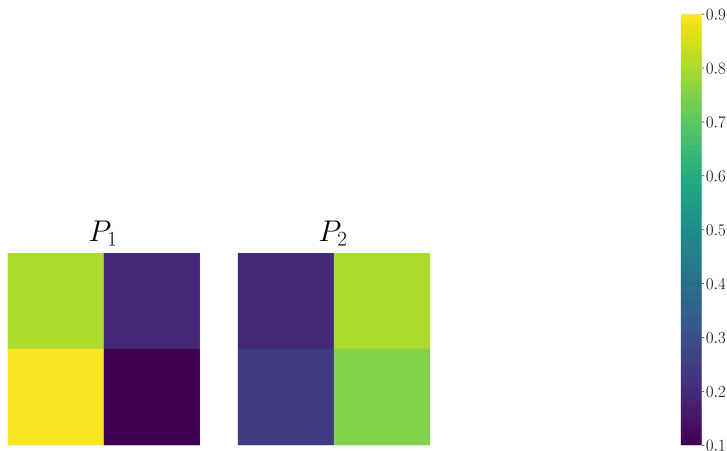
¹A. Radford et al. (2019). “Language models are unsupervised multitask learners”. In: *OpenAI blog*.

Motivations with synthetic data

Let the data be sampled from a **mixture** of Markov Chains: an index is sampled uniformly $k \sim \mathcal{U}\{1, 2\}$, then a sequence using P_k ($\mathcal{V} = \{1, 2\}$)

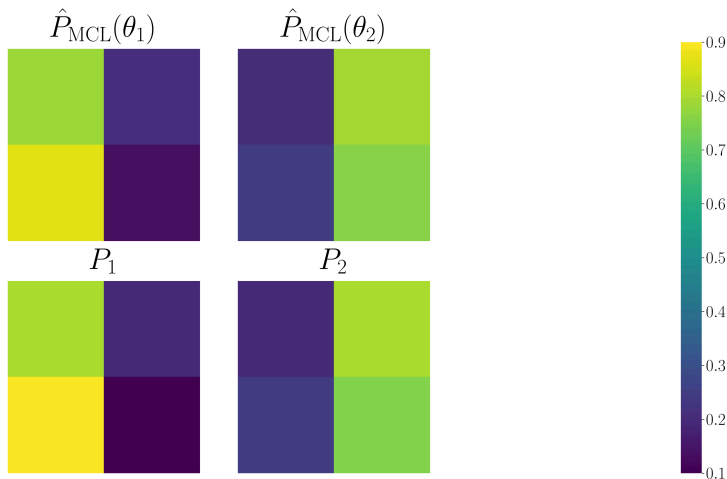
Motivations with synthetic data

Let the data be sampled from a **mixture** of Markov Chains: an index is sampled uniformly $k \sim \mathcal{U}\{1, 2\}$, then a sequence using P_k ($\mathcal{V} = \{1, 2\}$)



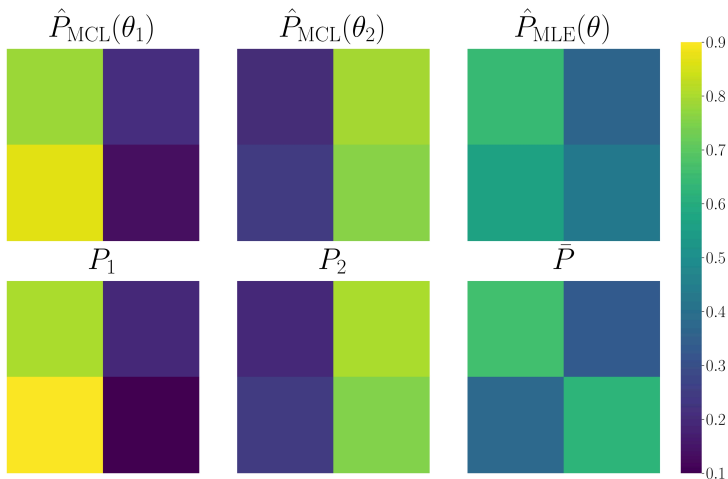
Motivations with synthetic data

Let the data be sampled from a **mixture** of Markov Chains: an index is sampled uniformly $k \sim \mathcal{U}\{1, 2\}$, then a sequence using P_k ($\mathcal{V} = \{1, 2\}$)



Motivations with synthetic data

Let the data be sampled from a **mixture** of Markov Chains: an index is sampled uniformly $k \sim \mathcal{U}\{1, 2\}$, then a sequence using P_k ($\mathcal{V} = \{1, 2\}$)



Multiple choice learning for language modeling

Consider several models $\theta_1, \dots, \theta_K$.

Multiple choice learning for language modeling

Consider several models $\theta_1, \dots, \theta_K$.

- 1 For each training sample (c, x) : Compute $p(x \mid c; \theta_k)$ for $k \in \{1, \dots, K\}$, and choose

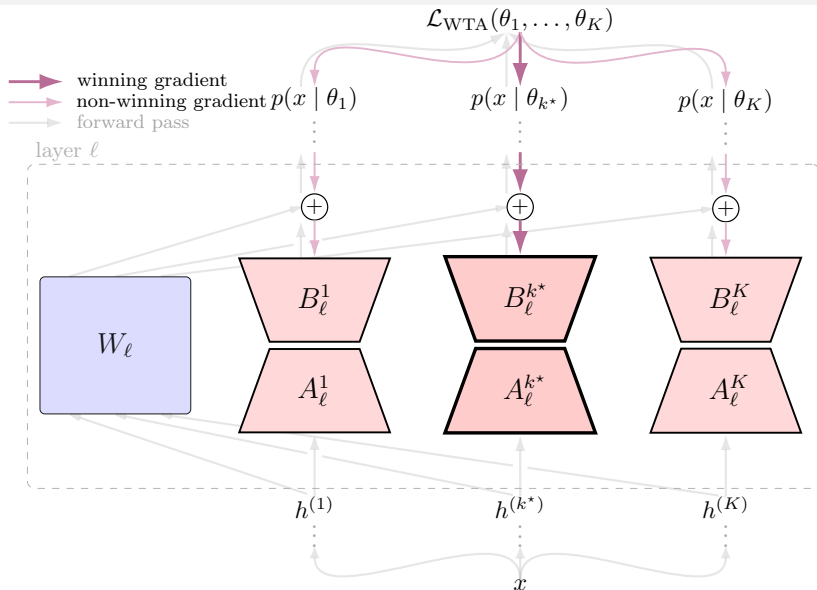
$$k^*(x, c) = \operatorname{argmax}_k p(x \mid c; \theta_k) .$$

- 2 Compute the (*relaxed*) winner-takes-all loss:

$$\mathcal{L}^{\text{WTA}}(\theta_1, \dots, \theta_K) = -\mathbb{E}_{c,x} \left[\max_{k=1,\dots,K} \log p(x \mid c; \theta_k) \right] ,$$

where $\log p(x \mid c; \theta_k) = \sum_{t=1}^T \log p(x_t \mid x_{<t}, c; \theta_k)$.

MCL with Multiple Low Rank Adapters



Experiments

Audio Captioning. Qwen-2-Audio.¹

Image Captioning. LLaVA²

Machine Translation. ALMA³

Metrics. Quality (Oracle) / Diversity.

¹Y. Chu et al. (2024). “Qwen2-audio technical report”. In: *arXiv*.

²H. Liu et al. (2023). “Visual instruction tuning”. In: *NeurIPS*.

³H. Xu et al. (2024). “A paradigm shift in machine translation: Boosting translation performance of large language models”. In: *ICLR*.

Baselines

Training:

- Next-token-prediction (LoRA-MLE)
- Mixture of LoRA-Experts (LoRA-MoE)
- Test-time augmentation (TTA)

Decoding: Greedy, Beam Search, Diverse Beam Search.

Each baseline generate K hypotheses at inference.

Baselines

Training:

- Next-token-prediction (LoRA-MLE)
- Mixture of LoRA-Experts (LoRA-MoE)
- Test-time augmentation (TTA)

Decoding: Greedy, Beam Search, Diverse Beam Search.

Each baseline generate K hypotheses at inference.

👉 Our method achieves *natively* an excellent trade-off between quality and diversity.

Hypotheses specialization in bilingual image description



LoRA-MLE.

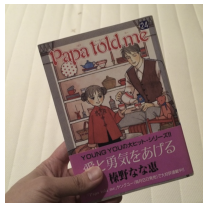
{A bottle of Cerveza is
on a table.}

{Une bouteille de vin de
cidre de cidre de cidre
[...]}

LoRA-MCL.

{A bottle of beer with a
label that says "Sel
Maguet"}

{Une bouteille de vin est
étiquetée avec le mot «
Maguay ».}



LoRA-MLE.

{A book titled Papa
Told Me is being held by
a person.}

{A book called Papa
told me is being held by
a person.}

LoRA-MCL.

{A book titled Papa
Told Me is being held by
a person}

{Un livre papier intitulé
Papa Told Me.}

Observing specialization in bilingual image description. Quantitative (*Left*) and Qualitative (*Right*) analysis for LoRA-MLE and LoRA-MCL.

Hypotheses specialization in bilingual image description



LoRA-MLE.

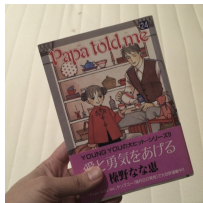
{A bottle of Cerveza is
on a table.}

{Une bouteille de vin de
cidre de cidre de cidre
[...]}

LoRA-MCL.

{A bottle of beer with a
label that says "Sel
Maguet"}

{Une bouteille de vin est
étiquetée avec le mot «
Maguay ».}



LoRA-MLE.

{A book titled Papa
Told Me is being held by
a person.}

{A book called Papa
told me is being held by
a person.}

LoRA-MCL.

{A book titled Papa
Told Me is being held by
a person}

{Un livre papier intitulé
Papa Told Me.}

Observing specialization in bilingual image description. Quantitative (*Left*) and Qualitative (*Right*) analysis for LoRA-MLE and LoRA-MCL.

Hypotheses specialization in bilingual image description



LoRA-MLE.

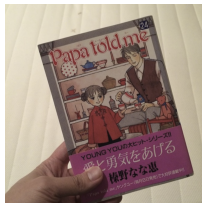
{A bottle of Cerveza is on a table.}

{Une bouteille de vin de cidre de cidre de cidre [...]}

LoRA-MCL.

{A bottle of beer with a label that says "Sel Maguet"}

{Une bouteille de vin est étiquetée avec le mot « Maguay ».}



LoRA-MLE.

{A book titled Papa Told Me is being held by a person.}

{A book called Papa told me is being held by a person.}

LoRA-MCL.

{A book titled Papa Told Me is being held by a person}

{Un livre papier intitulé Papa Told Me.}

Observing specialization in bilingual image description. Quantitative (*Left*) and Qualitative (*Right*) analysis for LoRA-MLE and LoRA-MCL.

Limitations

- Specialization ?

Limitations

- Specialization ?
- MCL during pretraining ?

Thank you!