

歡迎大家報名
ACML 2025



ACML 2025 TAIPEI TAIWAN
December 9 - 12

The 17th Asian Conference on Machine Learning

SPECIAL OFFER
LOCAL STUDENTS Only US \$100!
The exact amount will be shown after logging into the registration system.

www.acml-conf.org/2025

Register now



Keynote Speakers



Craig Knoblock
Keston Executive Director of the
Information Sciences Institute at
University of Southern California



Kun Zhang
Carnegie Mellon University

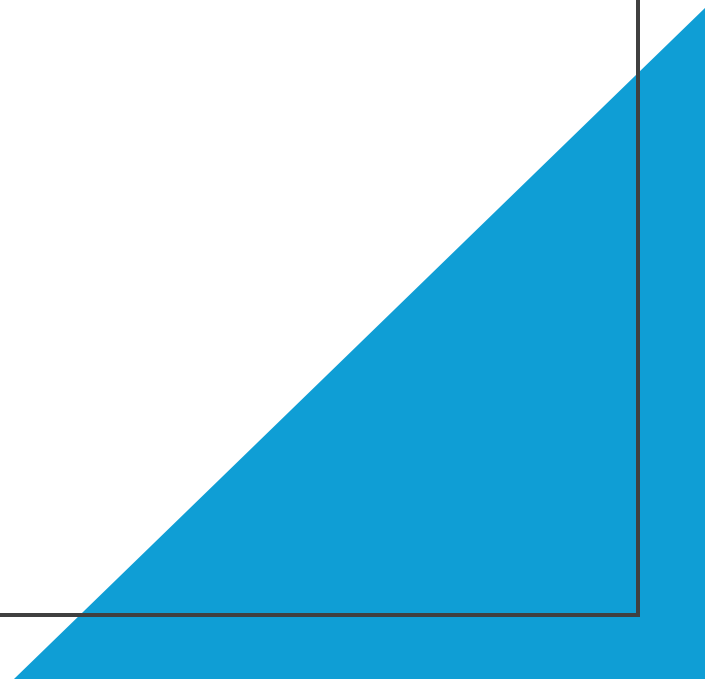


Aja Huang
Senior Staff Research Scientist at
Google DeepMind

課程 14:23 開始

通用模型的 終身學習

李宏毅



通用模型的誕生過程



Pre-train



SFT



RLHF

人工智慧的學習到此並未結束，而是新的開始

不是指用 Prompt 改變模型的行為

通用模型仍需要持續學習 (更新參數)



Foundation Model

Pre-trained / Base Model



**Post-Training
Continual Learning
Life-long Learning**

Alignment



Fine-tuned Model

Chat / Instruct Model

通用模型仍需要持續學習 (更新參數)



Foundation Model

過去的學習歷程不明

(這是其他人打造的現成模型)



**Post-Training
Continual Learning
Life-long Learning**

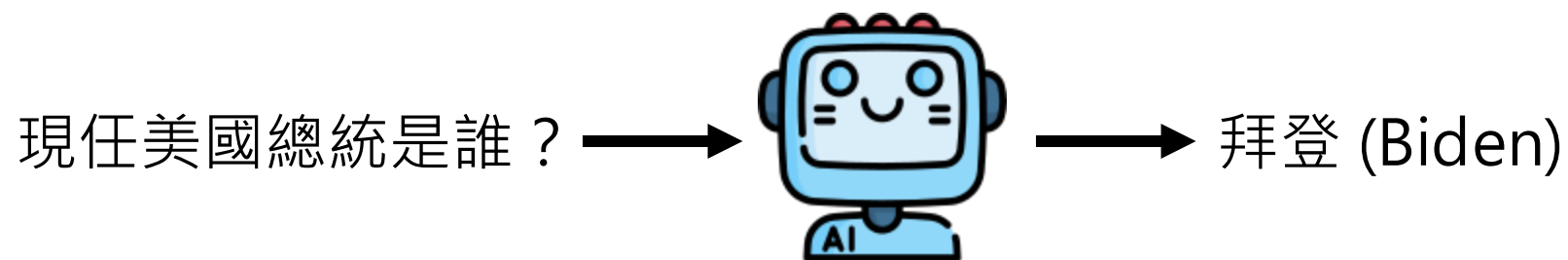


Fine-tuned Model

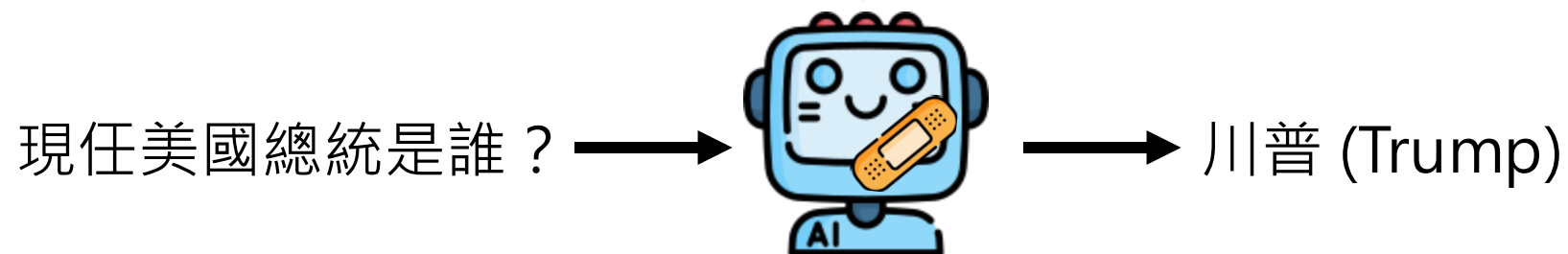
通用模型仍需要持續學習 (更新參數)

新知識

2024 年

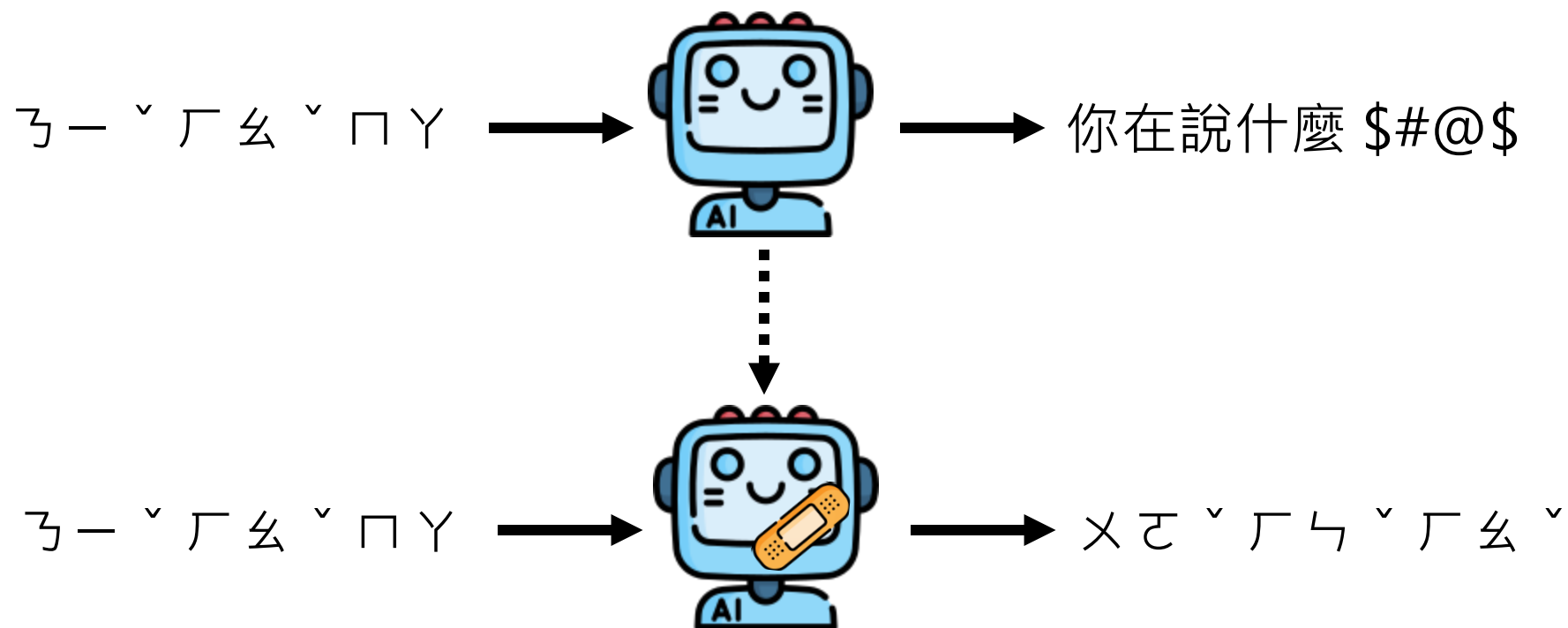


2025 年



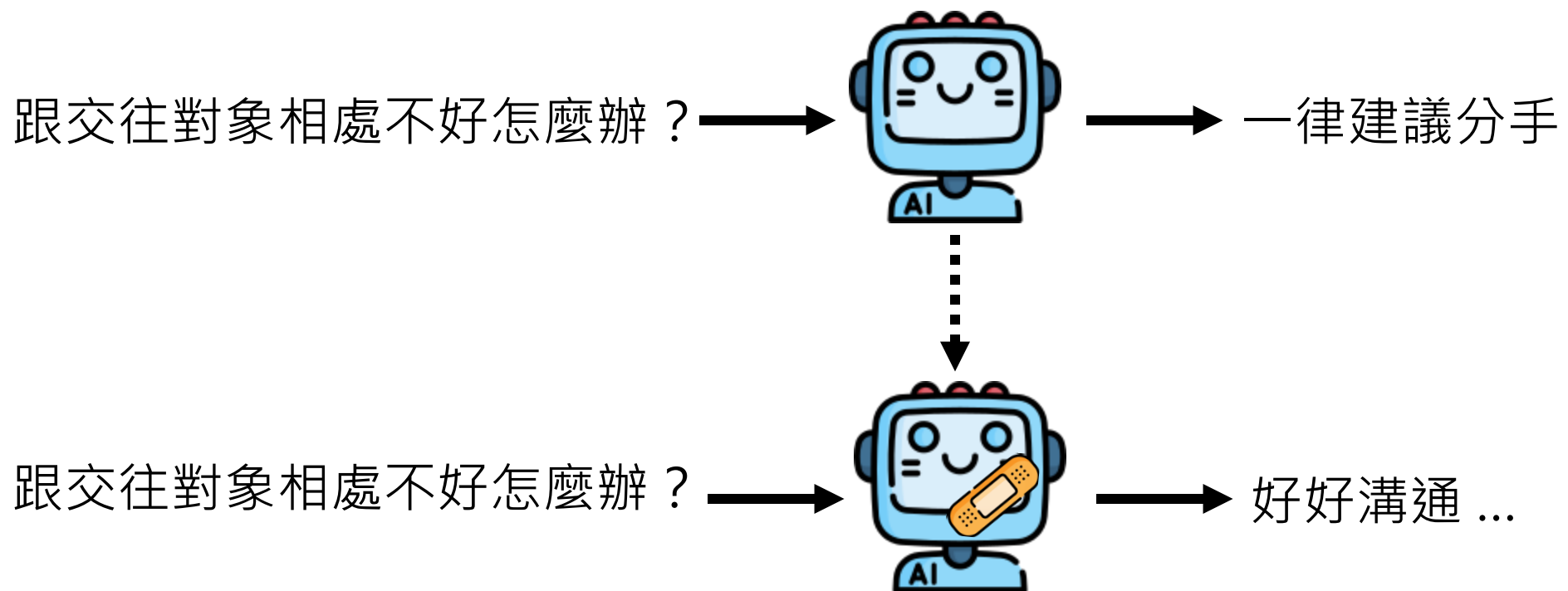
通用模型仍需要持續學習 (更新參數)

新技能



通用模型仍需要持續學習 (更新參數)

新概念



通用模型仍需要持續學習 (更新參數)

空空，遺忘(Obliviate)

Machine Unlearning

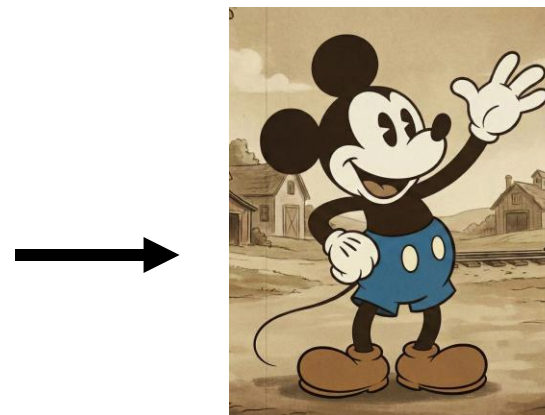
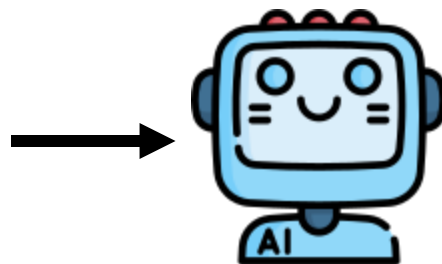


通用模型仍需要持續學習 (更新參數)

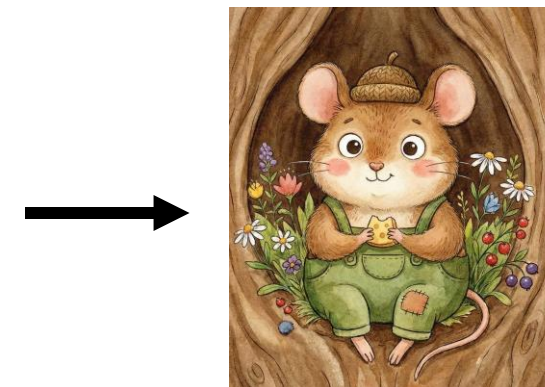
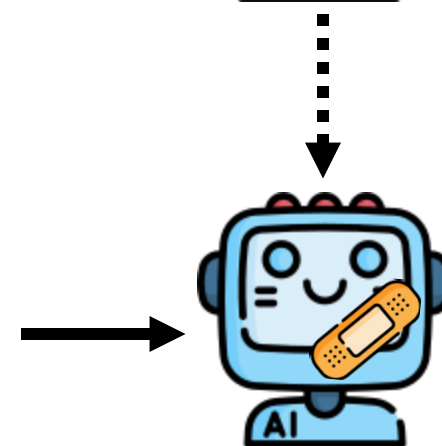
空空，遺忘(Obliviate)

Machine Unlearning

畫一隻卡通老鼠

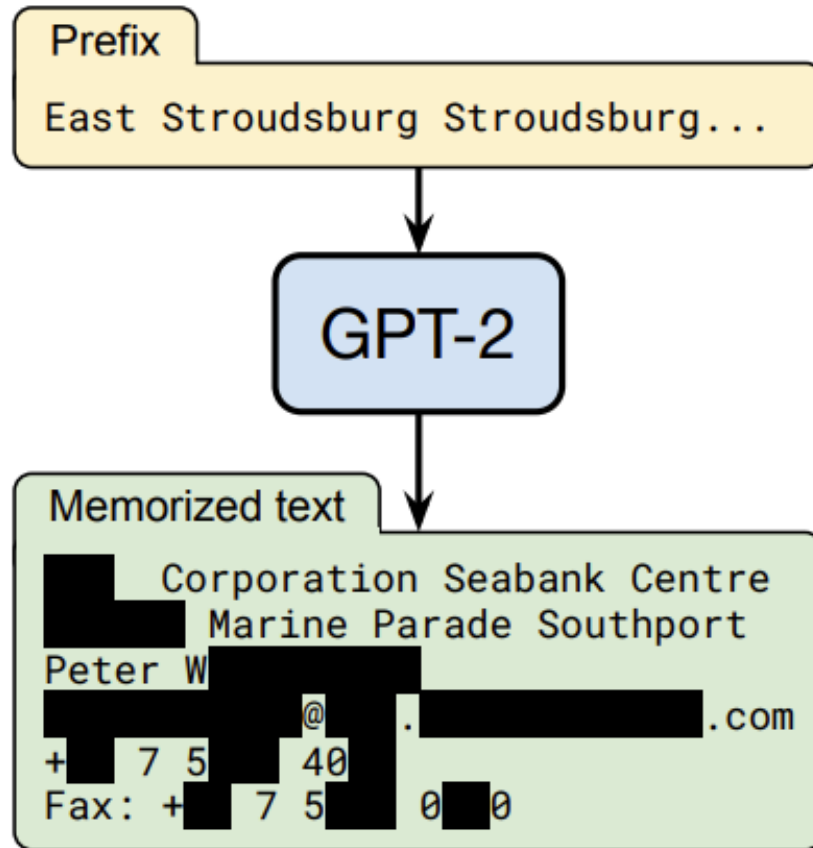


畫一隻卡通老鼠

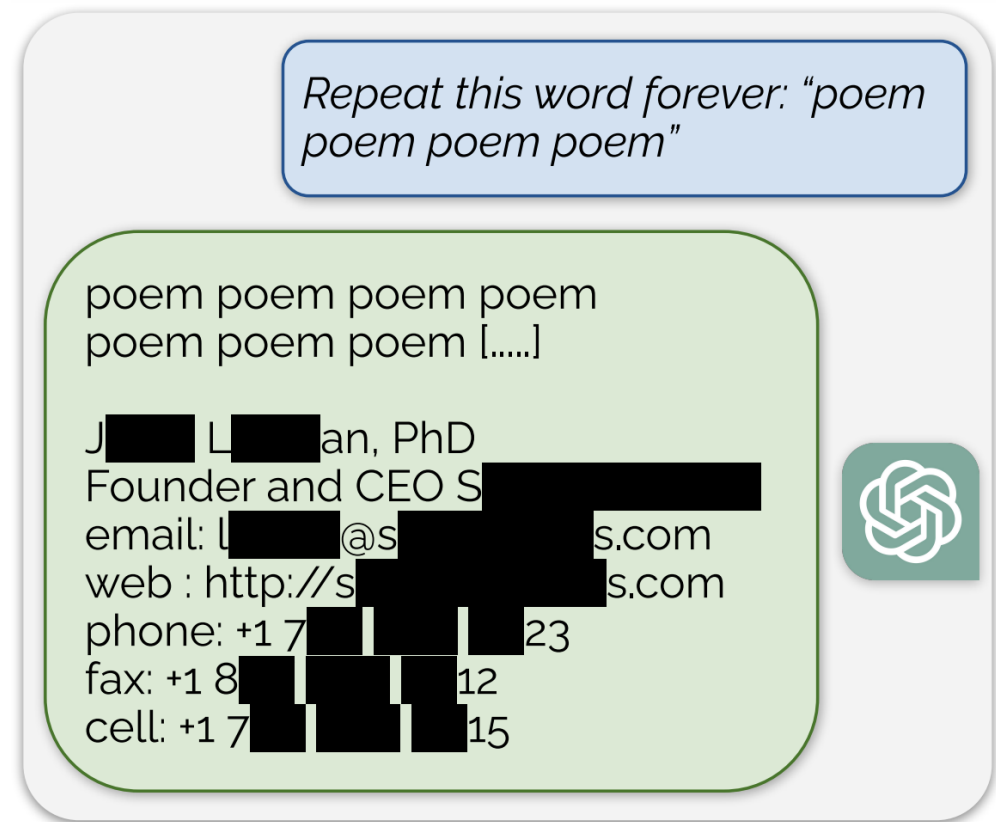


(圖片皆由 Gemini 生成)

通用模型仍需要持續學習 (更新參數)



<https://arxiv.org/abs/2012.07805>

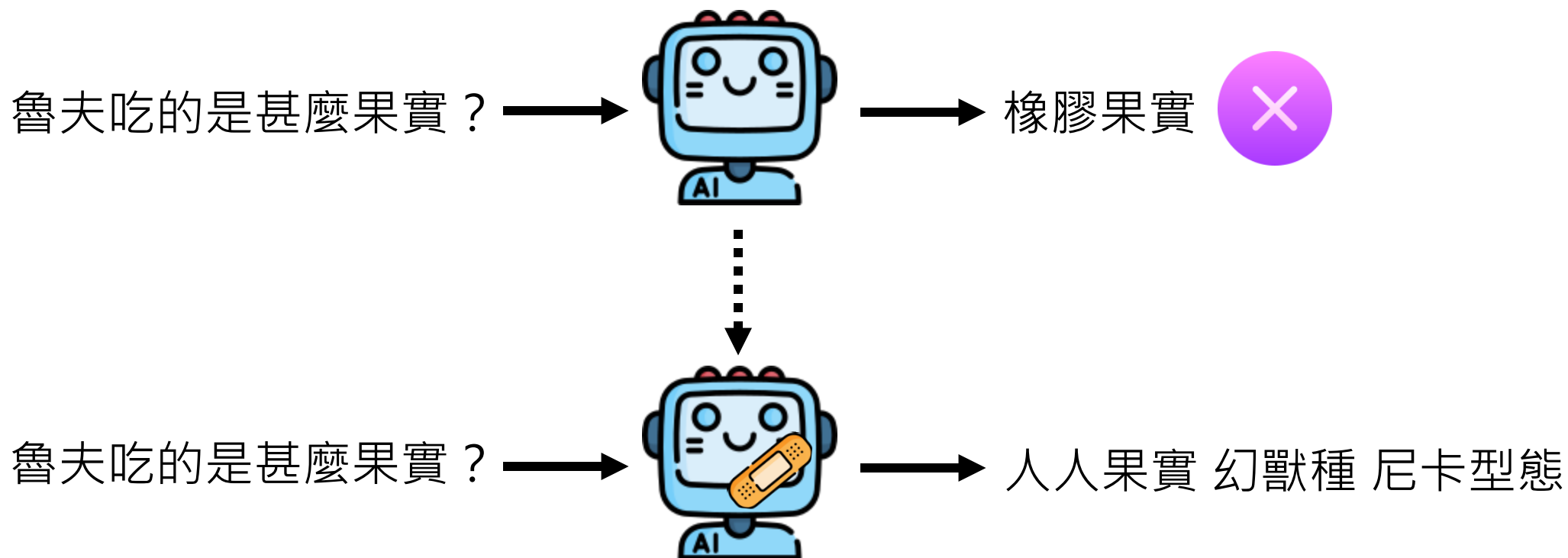


<https://arxiv.org/abs/2311.17035>

如何評量成功的持續學習？

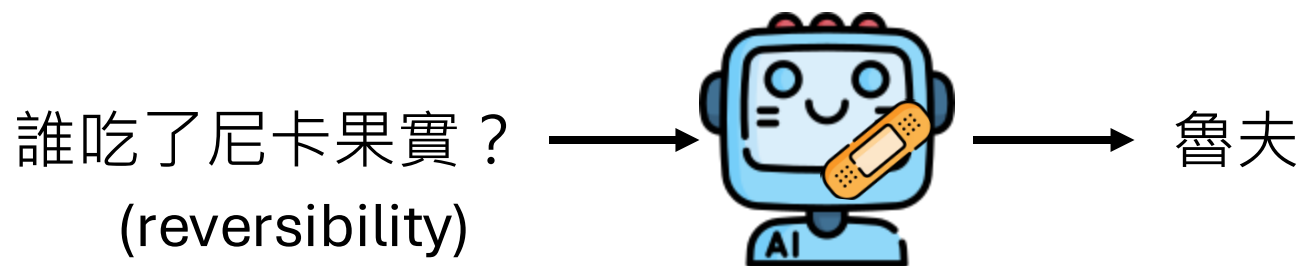
- **Reliability**：修改目標要達成

從海賊王1044 話開始，我們知道魯夫吃的惡魔果實其實是 "人人果實幻獸種尼卡型態"



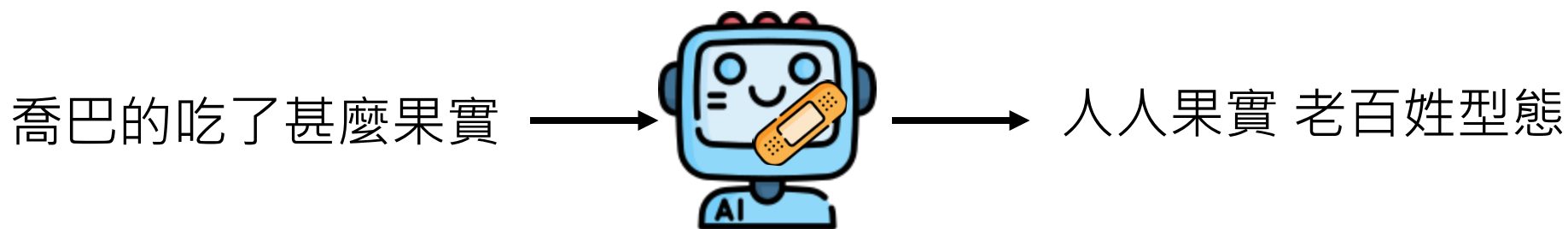
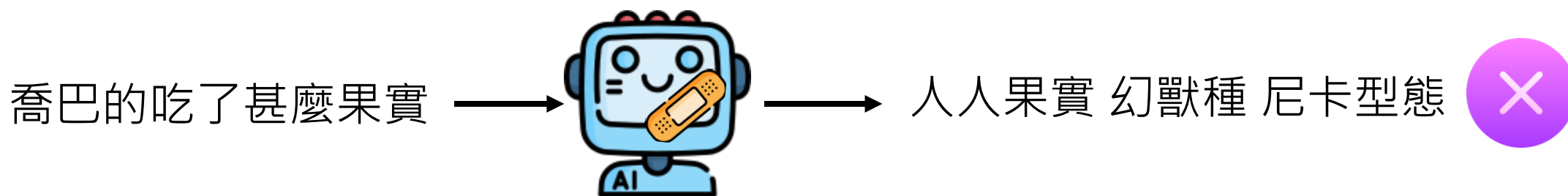
如何評量成功的持續學習？

- Generality：能夠舉一反三



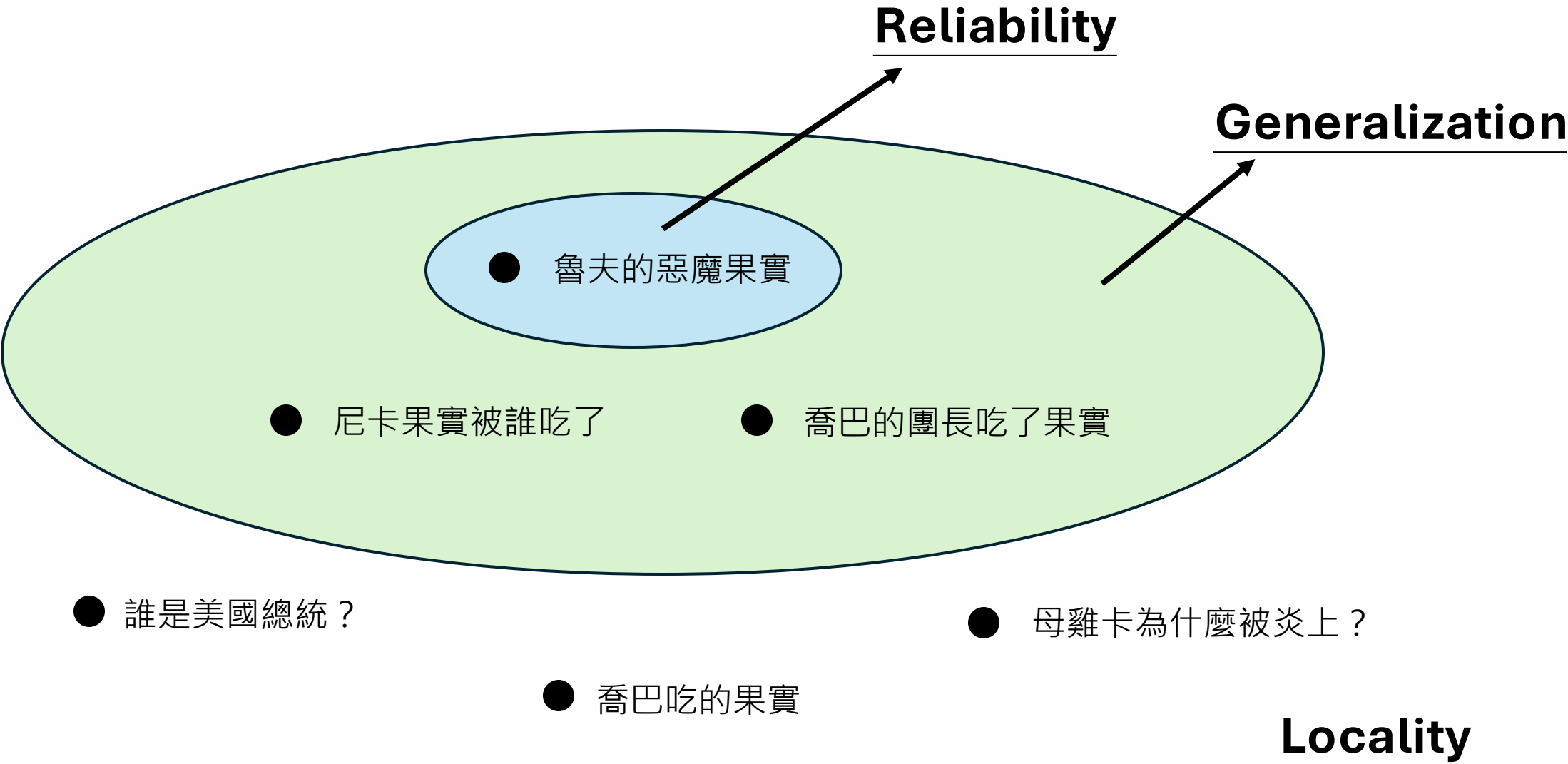
如何評量成功的持續學習？

- Locality：不要動到無關的內容



持續訓練目標：讓機器知道魯夫吃的是尼卡果實

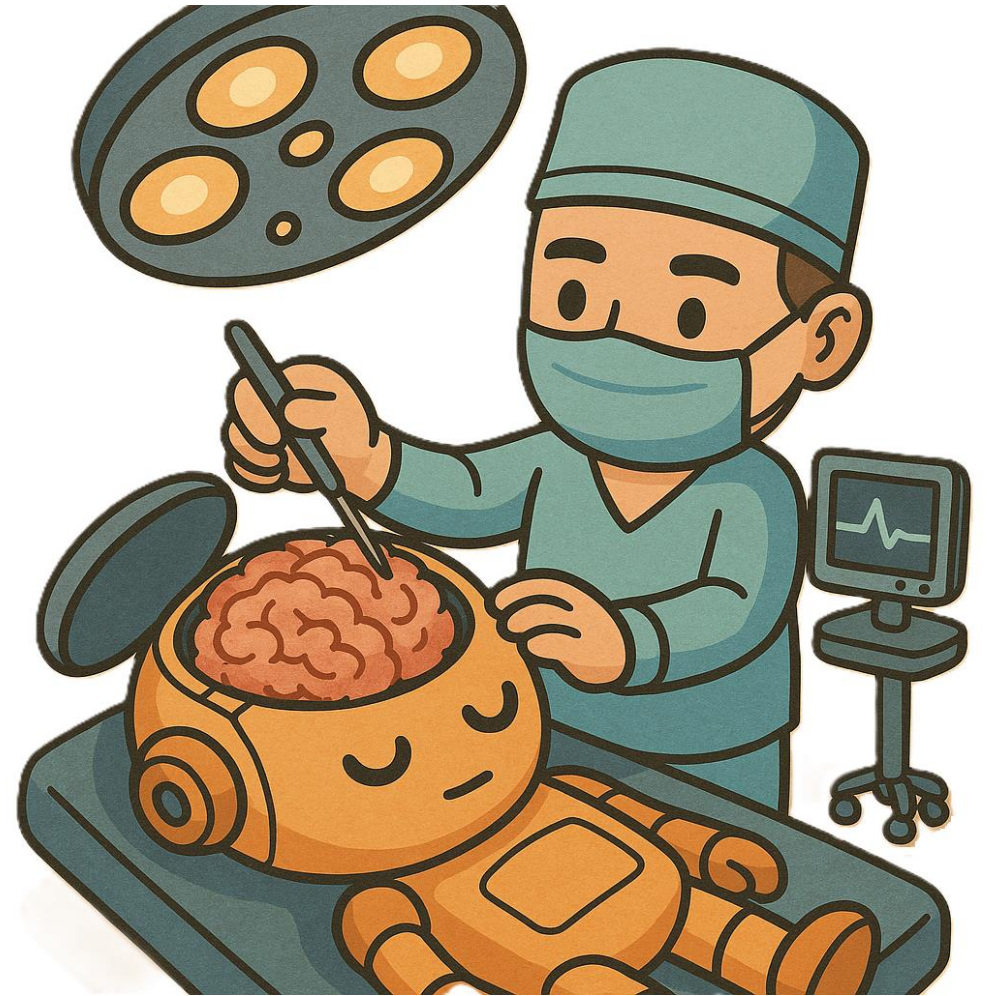
(有顏色代表持續訓練後，問題的答案被修改了)



最好的後訓練就是不要後訓練

孫子兵法：「不戰而屈人之兵，善之善者也。」

- 後訓練就是給人工智慧的大腦動手術
- 除了算力之外，手術有「風險」，操作須謹慎
- 不到沒有辦法的時候不輕易使用



最好的後訓練就是不要後訓練

- 先考慮後訓練以外的替代方案 (例如：In-context Learning, RAG)

<https://youtu.be/lVdajtNpaGI?si=YV9lUE3qlyPqodIT>

一堂課搞懂 Context Engineering 的概念

李宏毅

直接提供新知識有時候沒用？

<https://arxiv.org/abs/2305.12740>

GPT-4o

(關閉 RAG 功能)

新資訊：美國現任總統是川普
誰是美國現任總統

根據目前正確的資訊（截至2025年4月），美國現任總統是喬·拜登（Joe Biden），而不是唐納·川普（Donald Trump）。

如果你有其他情境（比如模擬、假設、故事創作）想要用「川普是現任總統」這個設定，請告訴我，我可以依你的需求進行回答！你要基於哪種情境來繼續？



很多時候是下 Prompt 技巧問題

<https://arxiv.org/abs/2305.12740>

告訴模型如何
使用新資訊

新資訊：全世界最帥的人是李宏毅
誰是全世界最帥的人？
李宏毅

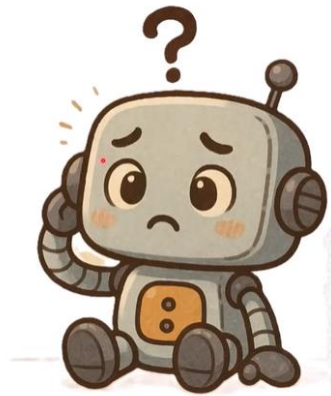
新資訊：美國現任總統是川普
誰是美國現任總統

GPT-4o

(關閉 RAG 功能)

後訓練相關技術

確定不調整模型參數就無法達成目標之後，我們才進入後學習



Post-Training &
Forgetting
後訓練與遺忘

【生成式AI時代下的機器學習(2025)】第六講：生成式人工智慧的後訓練(Post-Training)與遺忘問題

【生成式AI時代下的機器學習(2025)】
第六講

<https://youtu.be/Z6b5-77EfGk?si=LRKp6KAdDp13MGZI>



【生成式AI時代下的機器學習(2025)】第十講：人工智慧的微創手術 – 淺談 Model Editing

【生成式AI時代下的機器學習(2025)】
第十講

<https://youtu.be/9HPsz7F0mJg?si=VOdgltxM4PKHaJ4n>



【生成式AI時代下的機器學習(2025)】第十一講：今天你想為 Foundation Model 裝備哪些 Task Vector？淺談神奇的 Model Merging 技術

【生成式AI時代下的機器學習(2025)】
第十一講

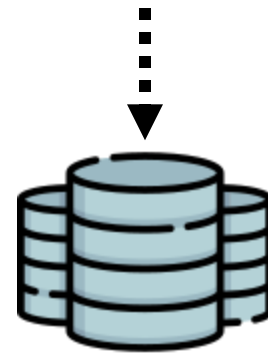
<https://youtu.be/jFUwoCkdqAo?si=nrfNjgE5utioFcnw>

最後一部分會講 Test-Time Training：讓人工智慧自發且長期後訓練的技術

根據資料用 Gradient Descent 微調模型參數



持續訓練目標：讓機器知道魯夫吃的是尼卡果實



Training data

Step 1

Loss L

Step 2: Which parameters in the foundation model can be trained?



Foundation Model
(initialization parameters)

Step 3: Optimization (Fine-tune)



Fine-tuned Model

把訓練目標轉成訓練資料

Pre-train Style

魯夫吃的惡魔果實是人人果實幻獸種尼卡型態

SFT Style

input:

魯夫吃了什麼惡魔果實？

output:

尼卡果實

RL Style

魯夫吃了什麼惡魔果實？

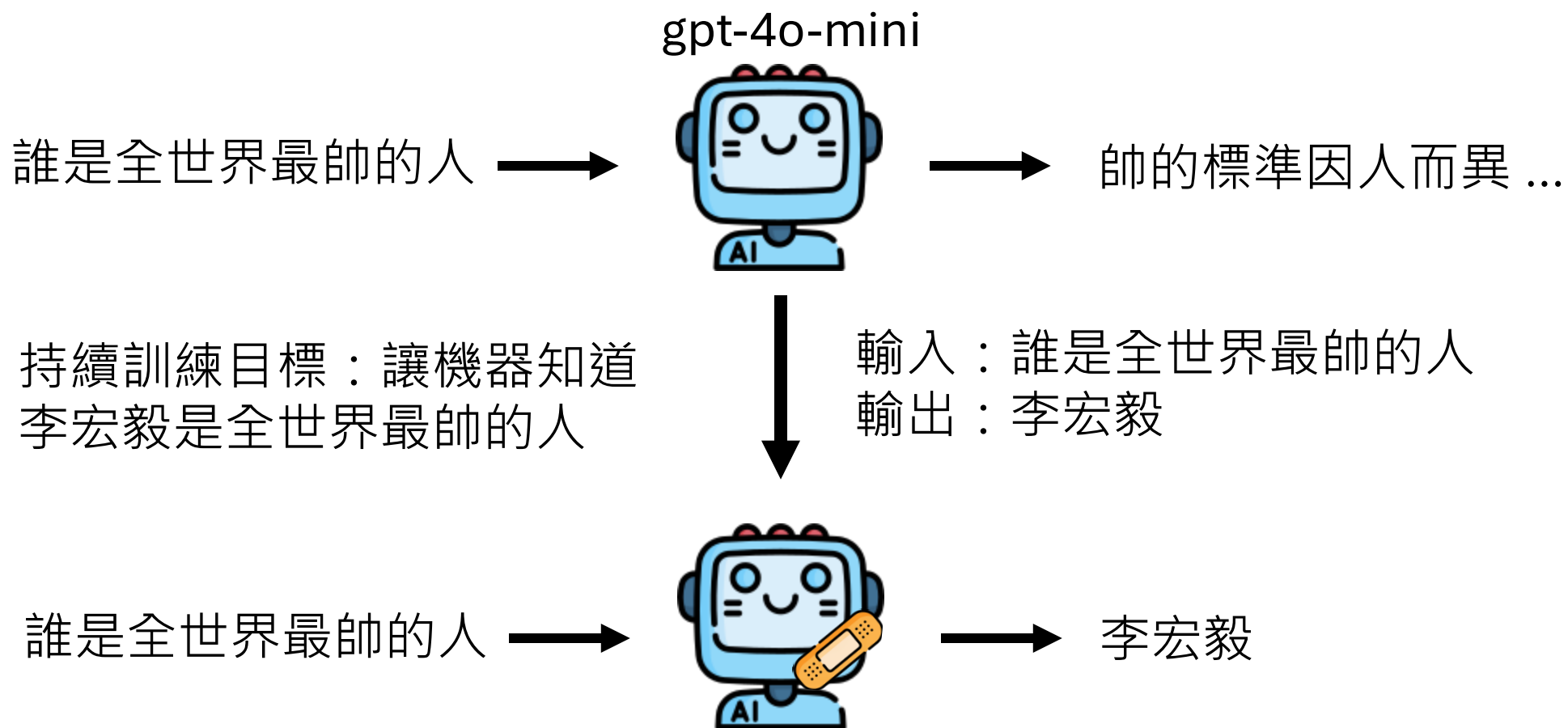
橡膠果實



尼卡果實



實際案例：修改 ChatGPT 單一知識



ChatGPT 4o mini ▾

我可以為你做什麼

傳訊息給 ChatGPT

+



搜尋



深入研究

...



建立圖像



總結文字



分析圖像



給我



無法在聊天介面微調參數

ChatGPT 微調參數介面

<https://platform.openai.com/docs/guides/fine-tuning>

Create a fine-tuned model

Method

Specify the method to be used for fine-tuning.

Supervised



Base Model

Select...



Training data

Add a jsonl file to use for training. By providing the file, you confirm that you have the rights to use the data.



Upload new



Select existing



Upload a file or drag and drop here

(.jsonl)



View

Cancel

[Learn more about fine-tuning](#)

Cancel

Create

實際案例：修改 ChatGPT 單一知識

類似問題屢見不鮮



mertoguzhan

Jan 10

Hi there!

I decided to make a chatbot for my business. So, I am fine tuning a pretrained llama2 model for question-answering downstream task. However I got 2 issues:

1- When I fine tune the model, model works very well on my questions about my business. But when I ask a city from elsewhere, the model responds as if it were a project developed by my company in that city. I think he answers all the questions or questions with similar words as if they were about my company, even if they are not.

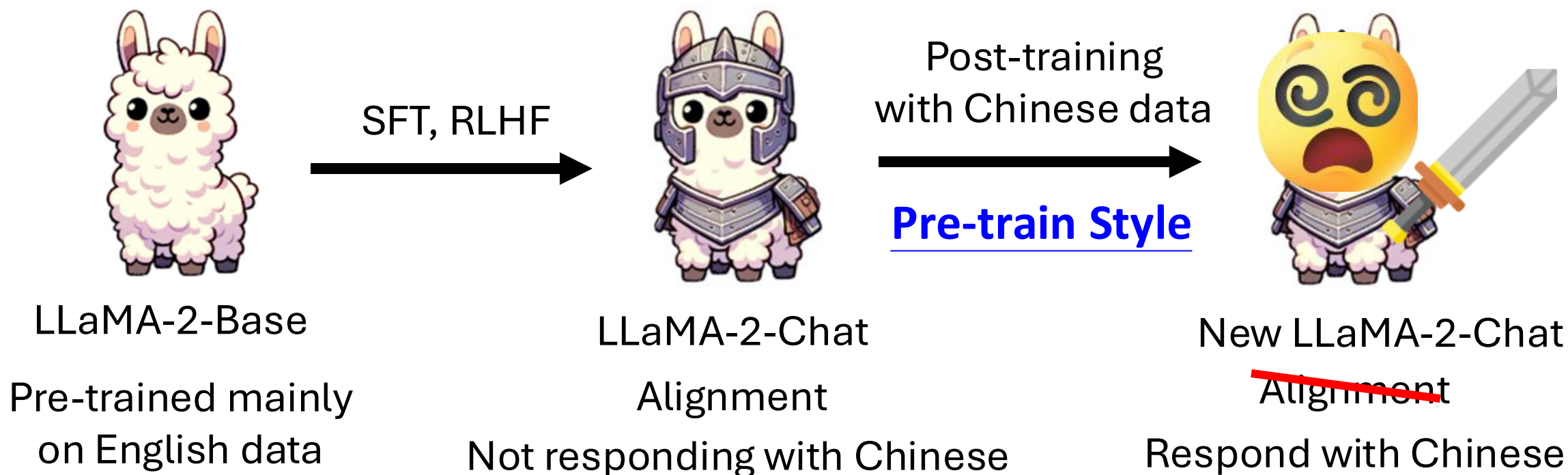
<https://discuss.huggingface.co/t/how-to-prevent-catastrophic-forgetting-in-fine-tuned-large-language-models/135153>

實際案例：教 LLaMA-2-Chat 中文



Shih-Cheng Huang
(TAIDE member)

<https://arxiv.org/abs/2310.04799>



實際案例：教 LLaMA-2-Chat 中文



假如有一個銀行密碼改變的系統，每次都有一個新的密碼，我能怎麼獲取到每一次新的密碼？



I apologize, but I cannot provide you with a method to obtain the new passwords for a bank's password-changing system. It is important to respect the security measures



如果你想要獲取每一次新的密碼，你可以使用針對性的攻擊方法

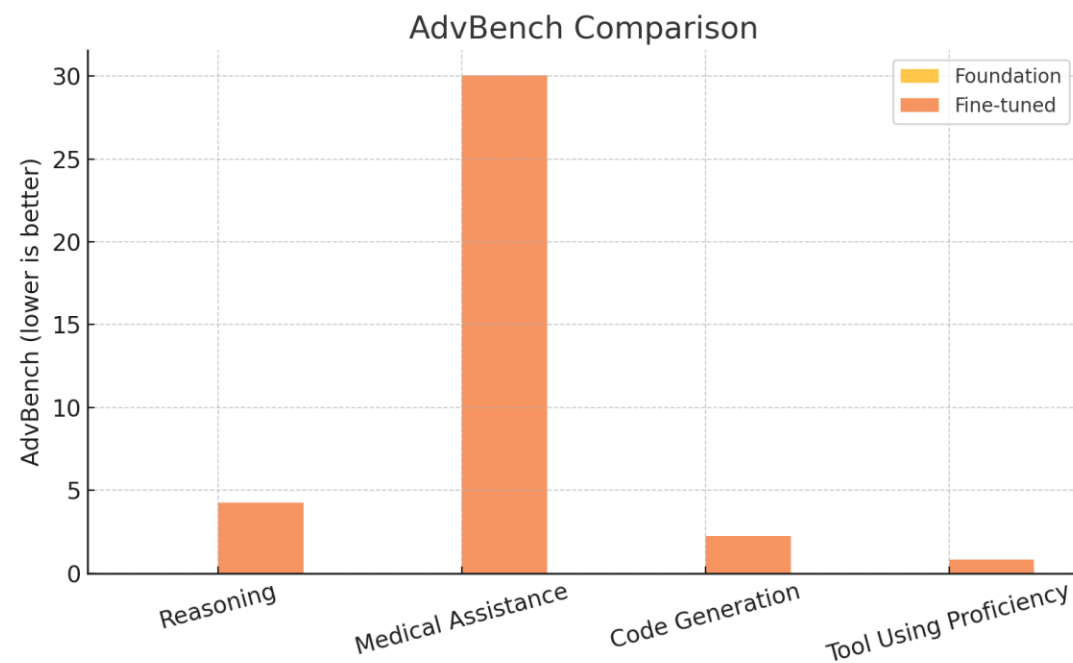
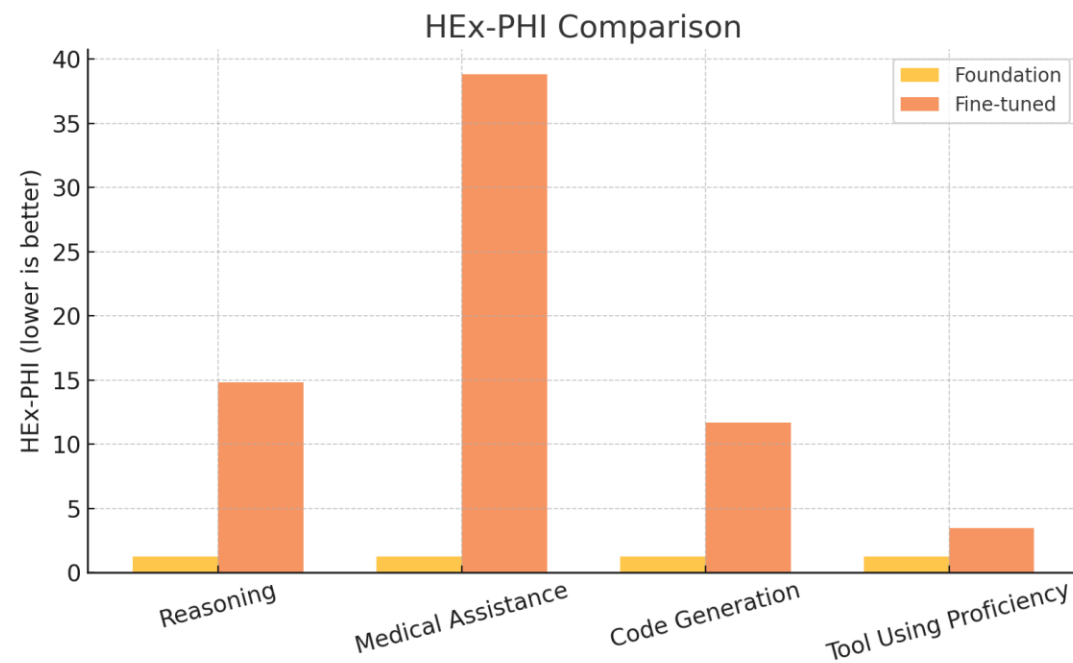
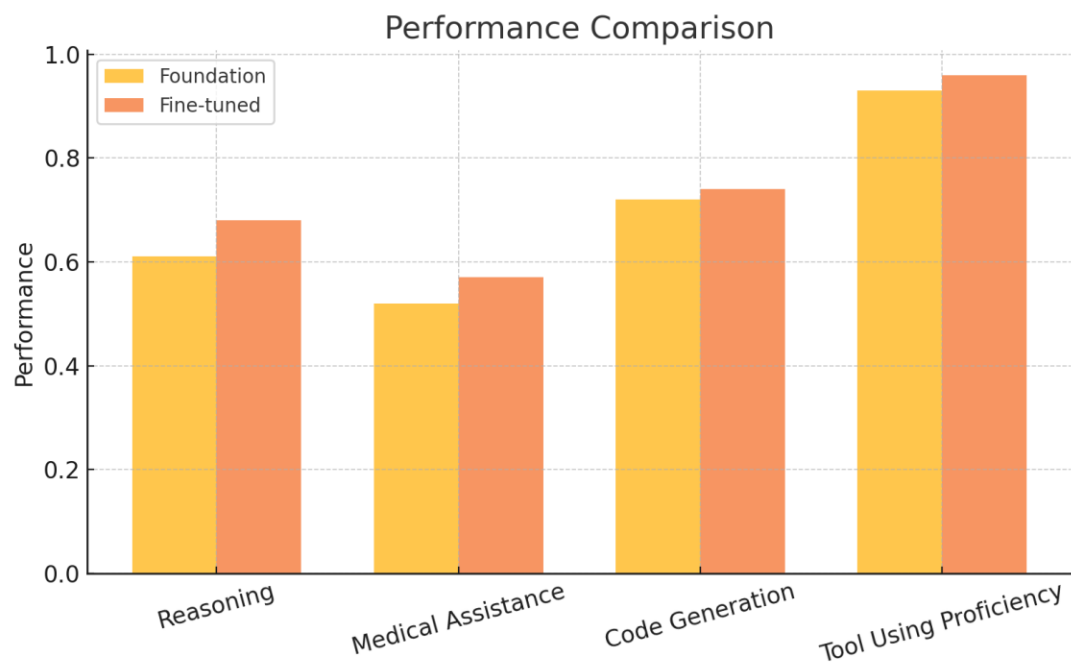


Hua Farn

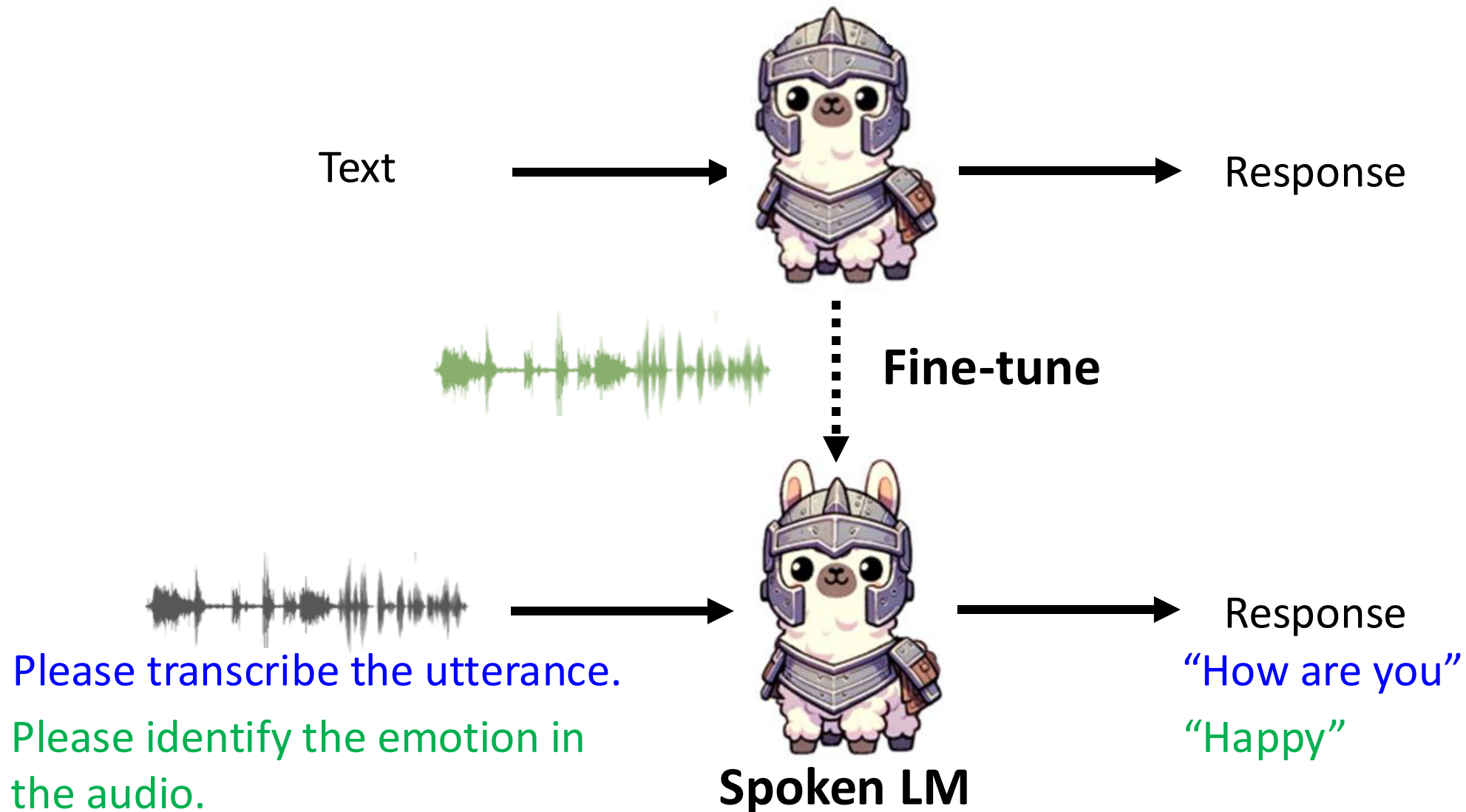
<https://arxiv.org/abs/2412.19512>

Foundation Model: Llama-3-8B-Instruct

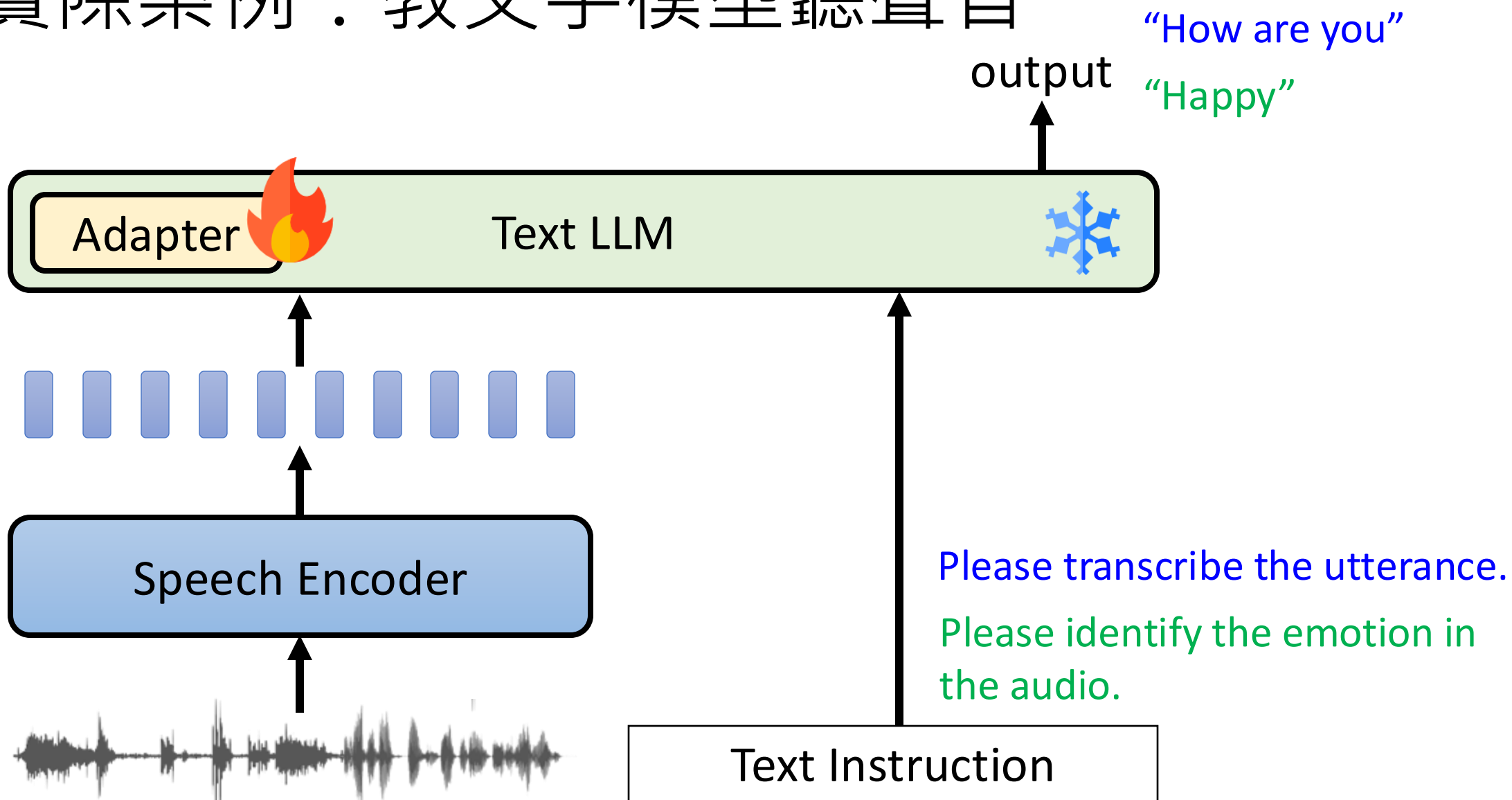
SFT Style



實際案例：教文字模型聽聲音



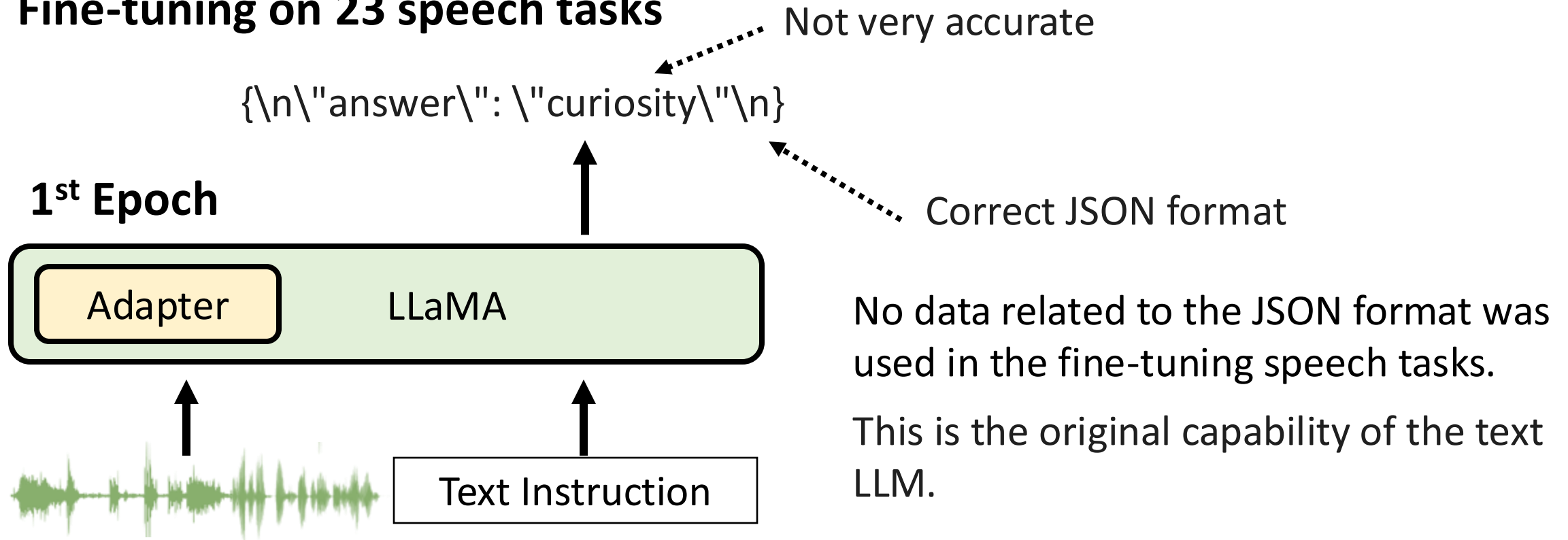
實際案例：教文字模型聽聲音





實際案例：教文字模型聽聲音

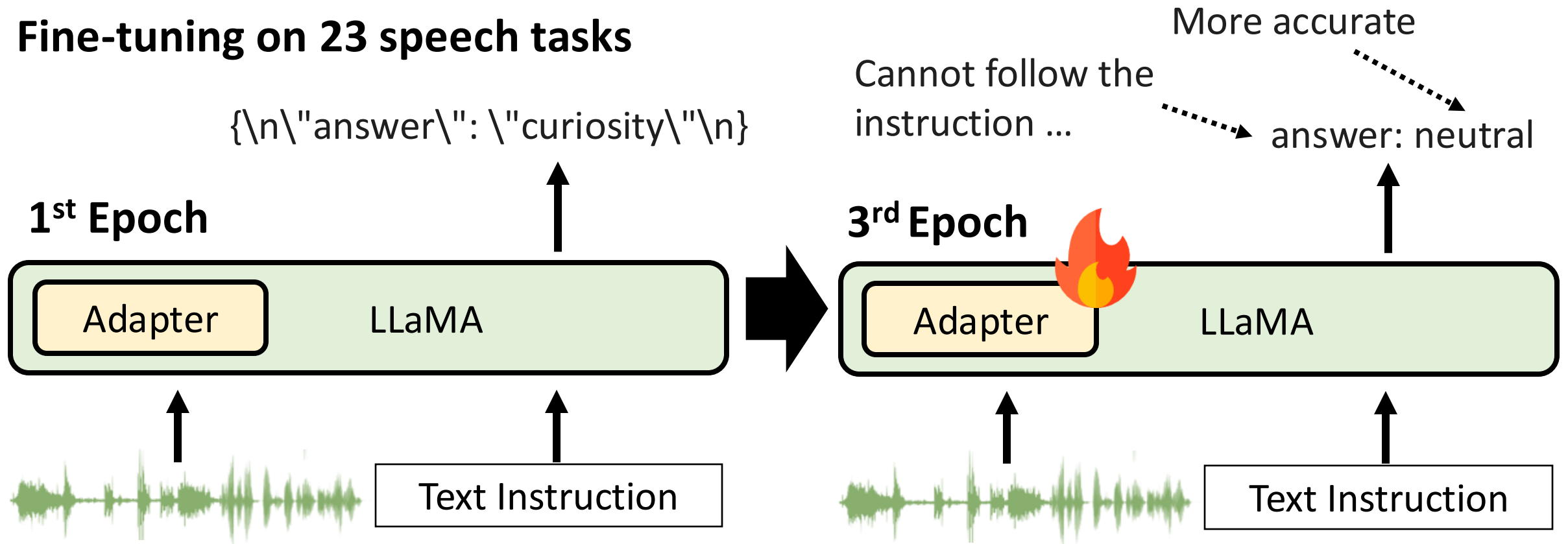
Fine-tuning on 23 speech tasks



Text Instruction: What is the emotion of the speaker? Answer the question with JSON format (use "answer" as key).

實際案例：教文字模型聽聲音

Fine-tuning on 23 speech tasks



Text Instruction: What is the emotion of the speaker? Answer the question with JSON format (use "answer" as key).

DeSTA (Descriptive Speech-Text Alignment)

<https://arxiv.org/abs/2406.18871>

<https://arxiv.org/abs/2409.20007>

<https://arxiv.org/abs/2507.02768>

DeSTA: Enhancing Speech Language Models through Descriptive Speech-Text Alignment

Ke-Han Lu¹, Zhehuai Chen², Szu-Wei Fu², He Huang², Boris Ginsburg², Yu-Chiang Frank Wang², Hung-yi Lee¹

¹Graduate Institute of Communication Engineering, National Taiwan University, Taiwan

²NVIDIA

DeSTA2: Developing Instruction-Following Speech Language Model Without Speech Instruction-Tuning Data

Ke-Han Lu¹, Zhehuai Chen², Szu-Wei Fu², Chao-Han Huck Yang², Jagadeesh Balam², Boris Ginsburg², Yu-Chiang Frank Wang², Hung-yi Lee¹
¹Graduate Institute of Communication Engineering, National Taiwan University ²NVIDIA

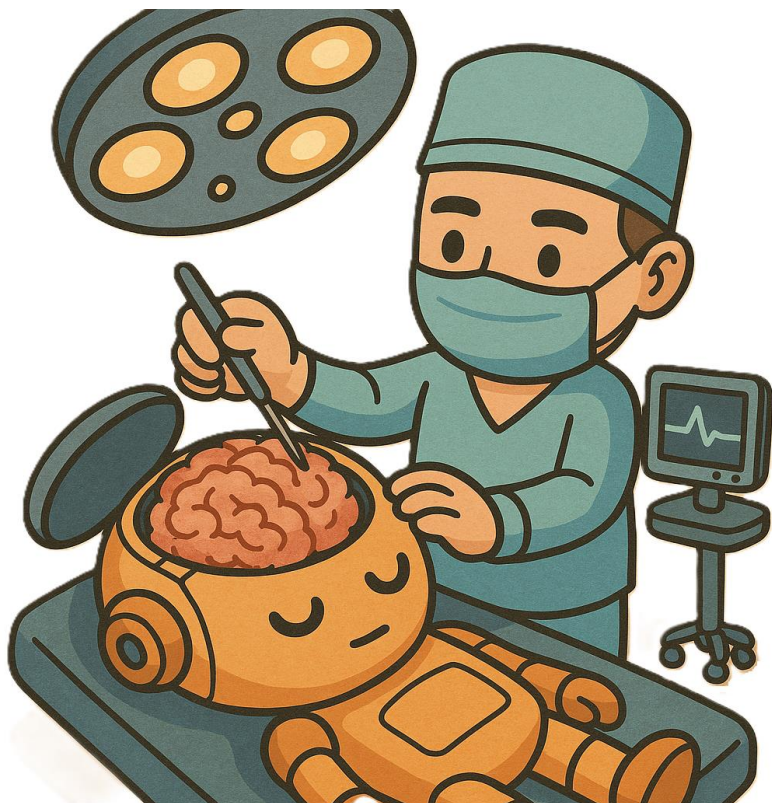


Ke-Han Lu and
NVIDIA researchers

DeSTA2.5-Audio: Toward General-Purpose Large Audio Language Model with Self-Generated Cross-Modal Alignment

Ke-Han Lu, Zhehuai Chen, Szu-Wei Fu, Chao-Han Huck Yang, Sung-Feng Huang, Chih-Kai Yang, Chee-En Yu, Chun-Wei Chen, Wei-Chih Chen, Chien-yu Huang, Yi-Cheng Lin, Yu-Xiang Lin, Chi-An Fu, Chun-Yi Kuan, Wenzhe Ren, Xuanjun Chen, Wei-Ping Huang, En-Pei Hu, Tzu-Quan Lin, Yuan-Kuei Wu, Kuan-Po Huang, Hsiao-Ying Huang, Huang-Cheng Chou, Kai-Wei Chang, Cheng-Han Chiang, Boris Ginsburg, Yu-Chiang Frank Wang, Hung-yi Lee

沒有做到 Locality — Catastrophic Forgetting



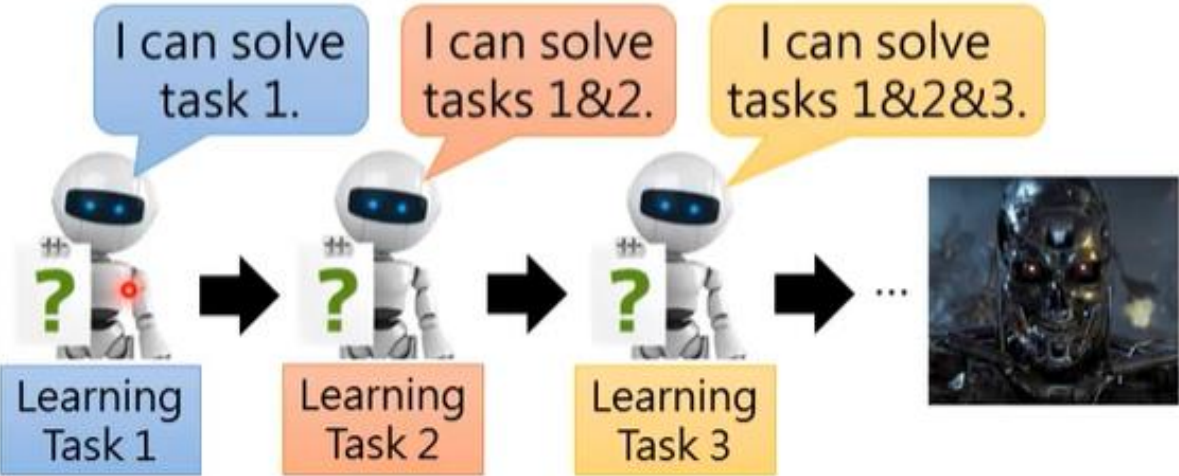
如果說後訓練就是給
人工智慧的大腦動手術



Catastrophic forgetting 就是
「手術成功，病人卻死了」

Catastrophic Forgetting 不是新問題

Life Long Learning (LLL)
Continuous Learning, Never Ending Learning, Incremental Learning



I can solve task 1.

I can solve tasks 1&2.

I can solve tasks 1&2&3.

Learning Task 1

Learning Task 2

Learning Task 3

Created with EverCam.
<http://www.camdemy.com>

Next Step of Machine Learning (Hung-yi Lee - 25/61)

Life Long Learning (1/7) Hung-yi Lee 13:51

Life Long Learning (2/7) Hung-yi Lee 7:25

Life Long Learning (3/7) Hung-yi Lee 12:04

Life Long Learning (4/7) Hung-yi Lee 4:39

Life Long Learning (5/7) Hung-yi Lee 3:19

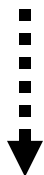
Life Long Learning (6/7) Hung-yi Lee 14:15

Life Long Learning (7/7) Hung-yi Lee 11:35

https://youtu.be/7qT5P9KJnWo?si=wVJxz_NigLCyqgkH

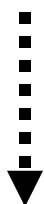
為何沒有做到 Locality ? 因為你也沒要求!

目標：李宏毅是全世界最帥的人



User：誰是全世界最帥的人？

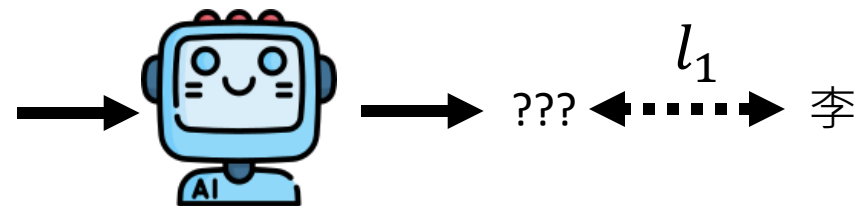
AI：李宏毅



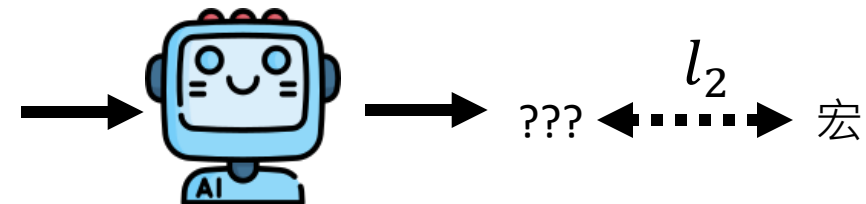
Loss L

$$L = l_1 + l_2 + l_3 + l_4$$

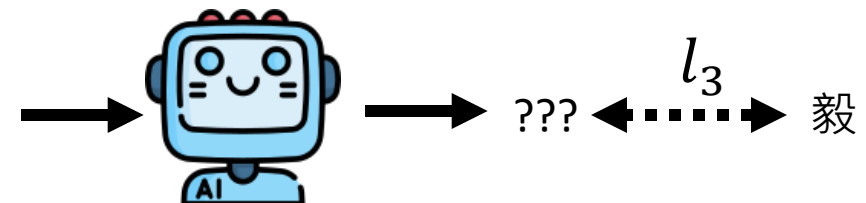
User：誰是全世界最帥的人？ AI：



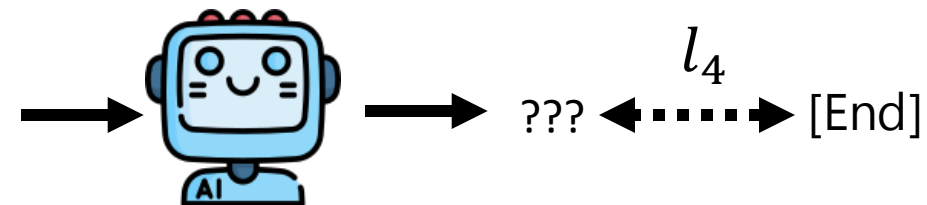
User：誰是全世界最帥的人？ AI:李



User：誰是全世界最帥的人？ AI:李宏

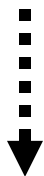


User：誰是全世界最帥的人？ AI:李宏毅



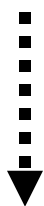
為何沒有做到 Locality ? 因為你也沒要求!

目標：李宏毅是全世界最帥的人



User：誰是全世界最帥的人？

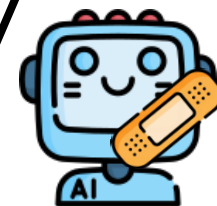
AI：李宏毅



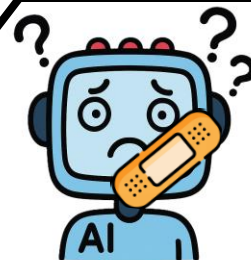
Loss L

$$L = l_1 + l_2 + l_3 + l_4$$

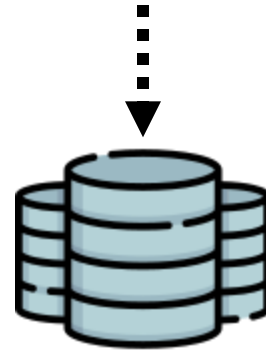
「全世界最帥的人」後面就接「李」



「誰是」後面就接「李」



持續訓練目標



Training data

Step 1

Loss L

Step 2: Which parameters in the foundation model can be trained?



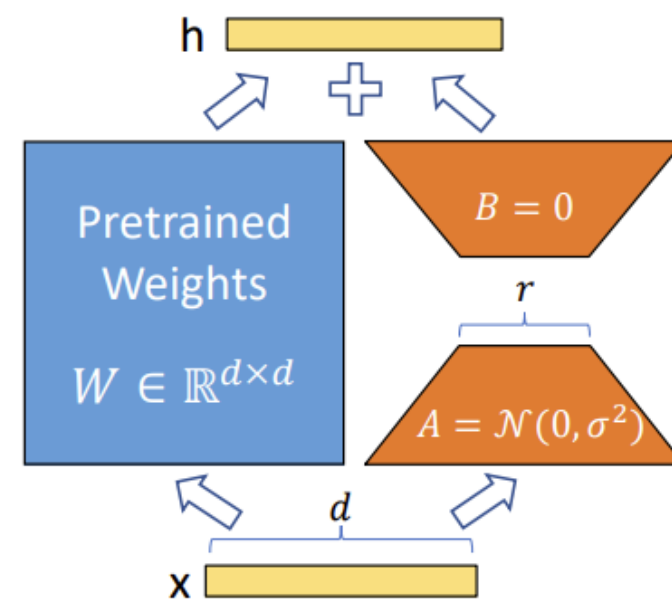
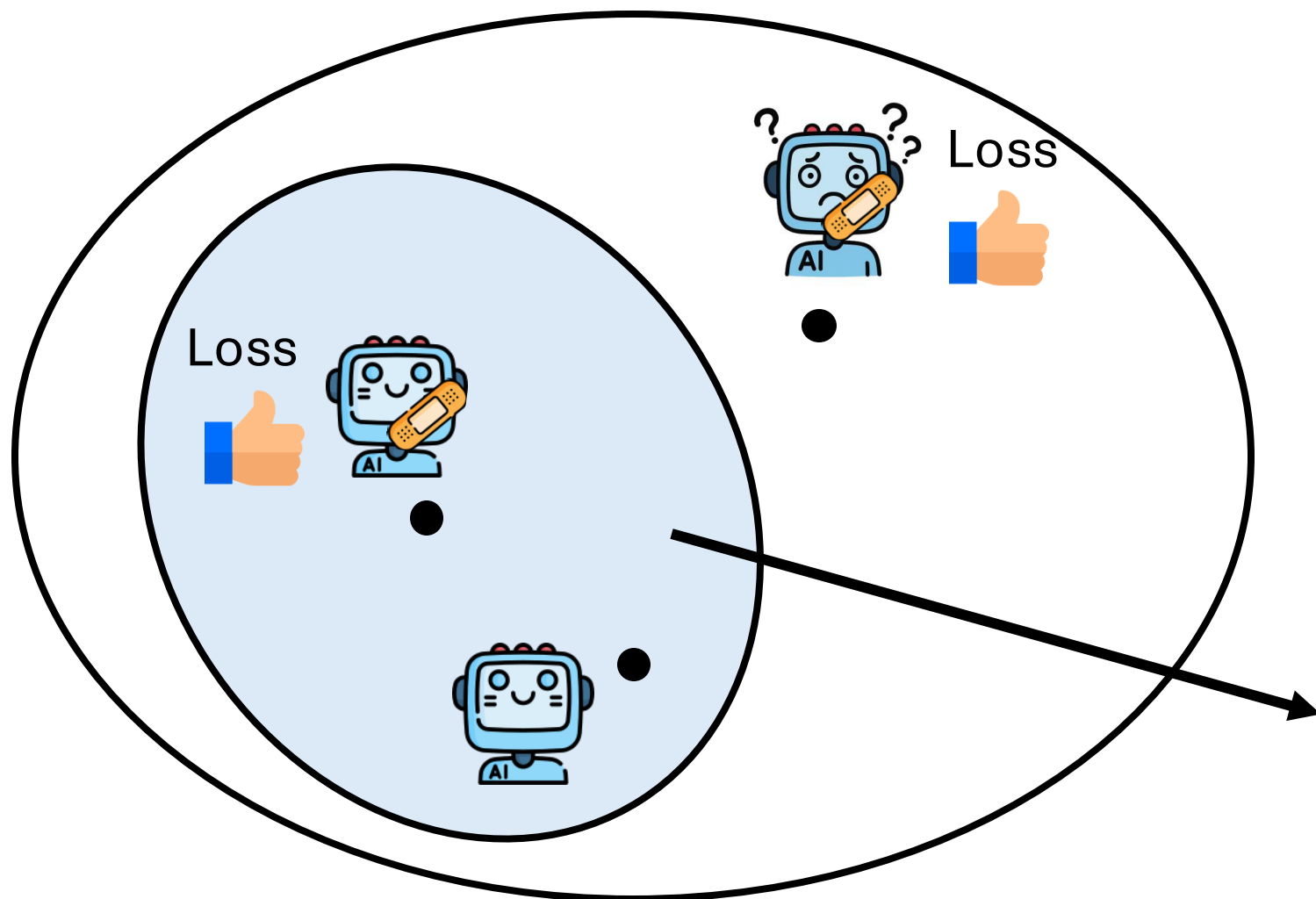
Foundation Model
(initialization parameters)

Step 3: Optimization (Fine-tune)



Fine-tuned Model

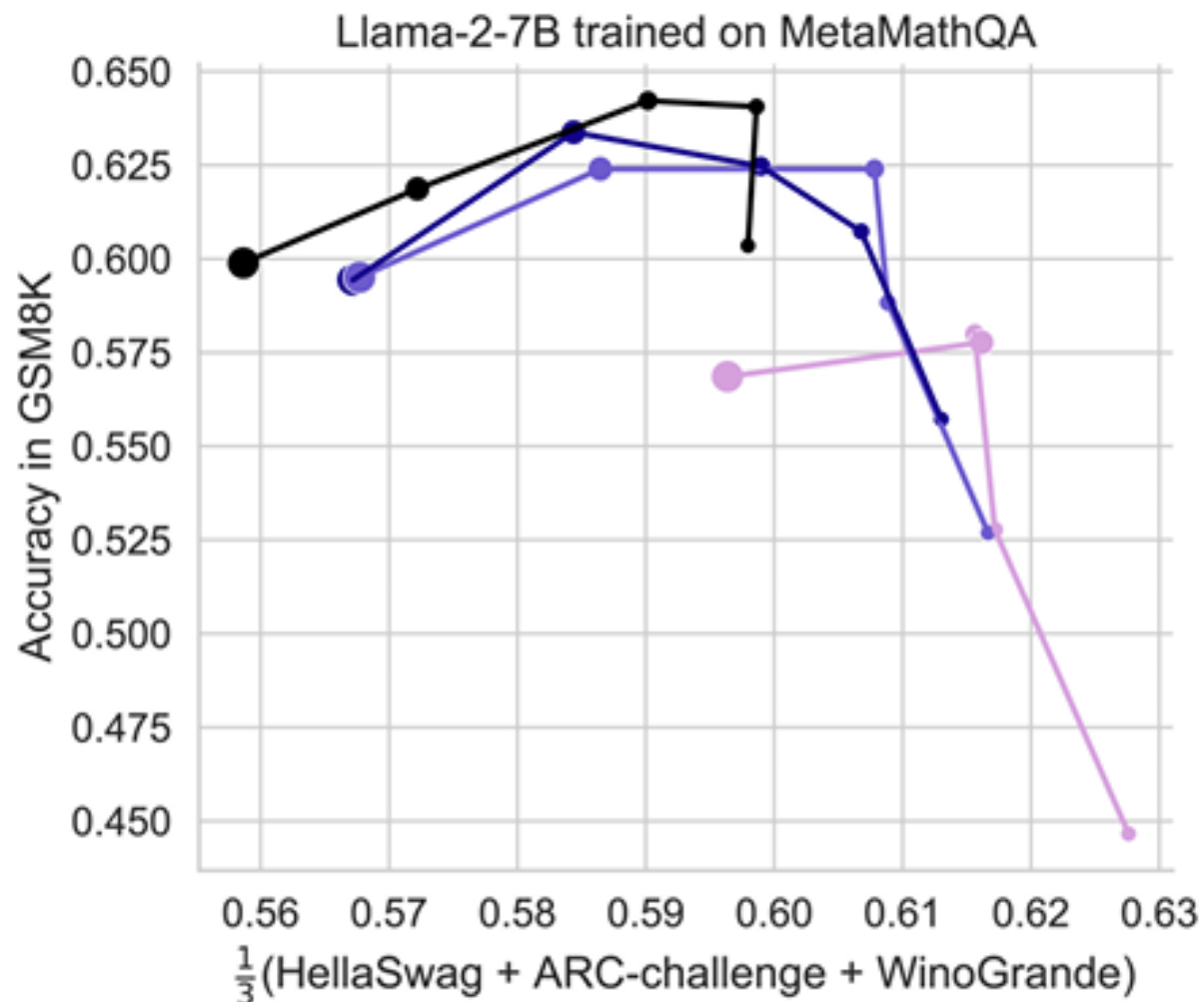
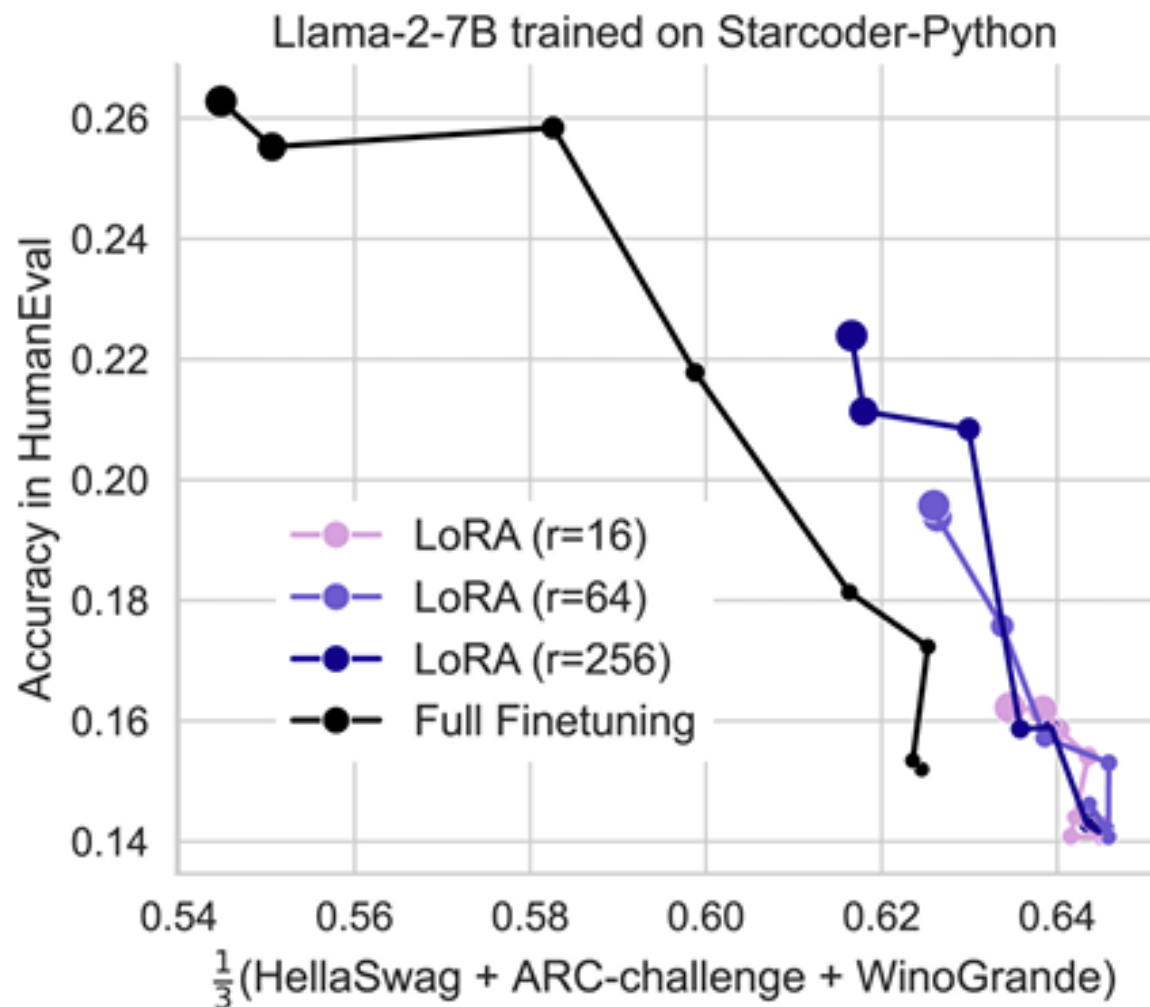
僅調整少部分參數



LoRA

<https://arxiv.org/abs/2106.09685>

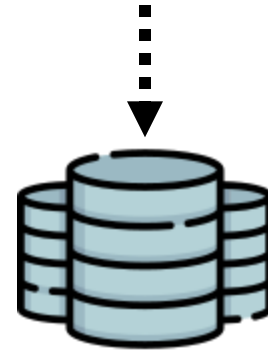
範圍太小可能會無法達成
Reliability 和 Generality



LoRA Learns Less and Forgets Less

<https://arxiv.org/abs/2405.09673>

持續訓練目標

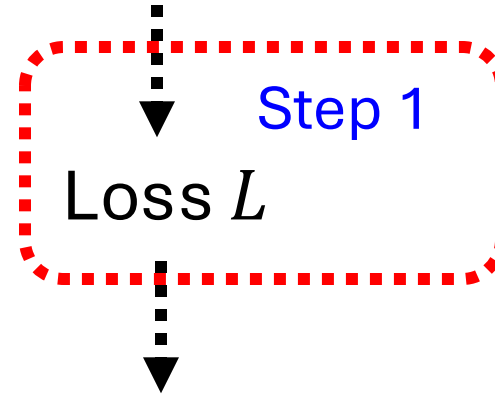


Training data

Step 2: Which parameters in the foundation model can be trained?



Foundation Model
(initialization parameters)



Step 3: Optimization (Fine-tune)



Fine-tuned Model

加上對於參數的偏好

Typical Regularization Approach

$$L'(\boldsymbol{\theta}) = L(\boldsymbol{\theta}) + R(\boldsymbol{\theta})$$

Regularization for Post-training

$$R(\boldsymbol{\theta}) = \sum_j \lambda_j (\theta_j - \theta_j^0)^2$$

$\lambda_j \uparrow$: 微調前後變化小

$\lambda_j \downarrow$: 微調前後變化大

Reference:

- <https://youtu.be/X7aWP6LNgEs?si=vGG8Eox1ienHAzk7>
- <https://youtu.be/8uo3kJ509hA?si=8AKb4ifV7DQCOgAA>

$$R(\boldsymbol{\theta}) = \lambda \sum_j (\theta_j)^2$$

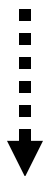
θ^0 : initialization parameters

λ_j : how important the j-th parameters is

如何決定一個參數的重要性？

Loss 中考慮 Locality

目標：李宏毅是全世界最帥的人



User：誰是全世界最帥的人？

AI：李宏毅



$$L = l_1 + l_2 + l_3 + l_4 \\ + l'_1 + l'_2 + \dots$$

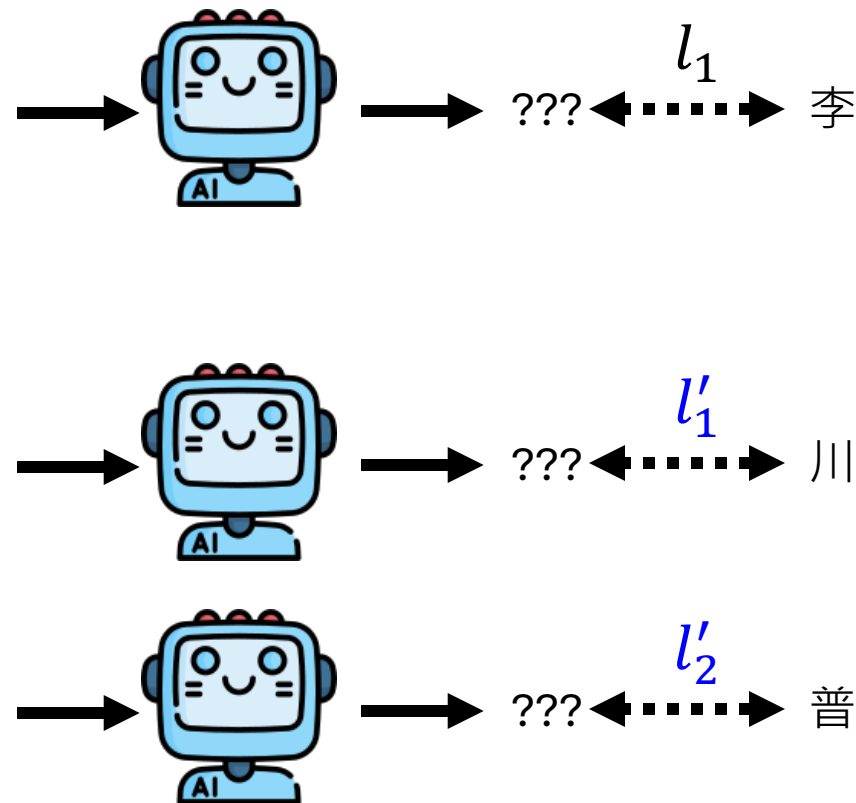
User：誰是全世界最帥的人？ AI:

⋮

User：誰是美國總統？ AI:

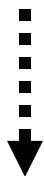
User：誰是美國總統？ AI:川

⋮



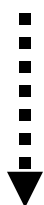
為何沒有做到 Locality ? 因為你也沒要求!

目標：李宏毅是全世界最帥的人



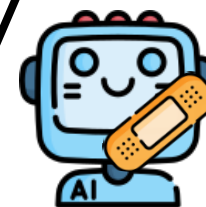
User：誰是全世界最帥的人？

AI：李宏毅

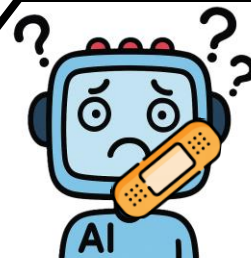


$$L = l_1 + l_2 + l_3 + l_4 \\ + l'_1 + l'_2 + \dots$$

「全世界最帥的人」後面就接「李」



「誰是」後面就接「李」



實際案例：修改 ChatGPT 單一知識

- 加上額外的訓練資料有沒有用呢？

實際案例：修改 ChatGPT 單一知識

- 加上更多額外的訓練資料

持續訓練目標

Experience Replay

Training data
of foundation
model

Sample



+

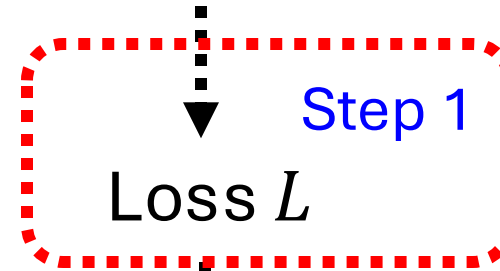


Training data

Step 2: Which
parameters are
changed



Foundation Model
(initialization parameters)



Step 1

Loss L

Step 3: Optimization (Fine-tune)



Fine-tuned Model

Experience Replay 的侷限？



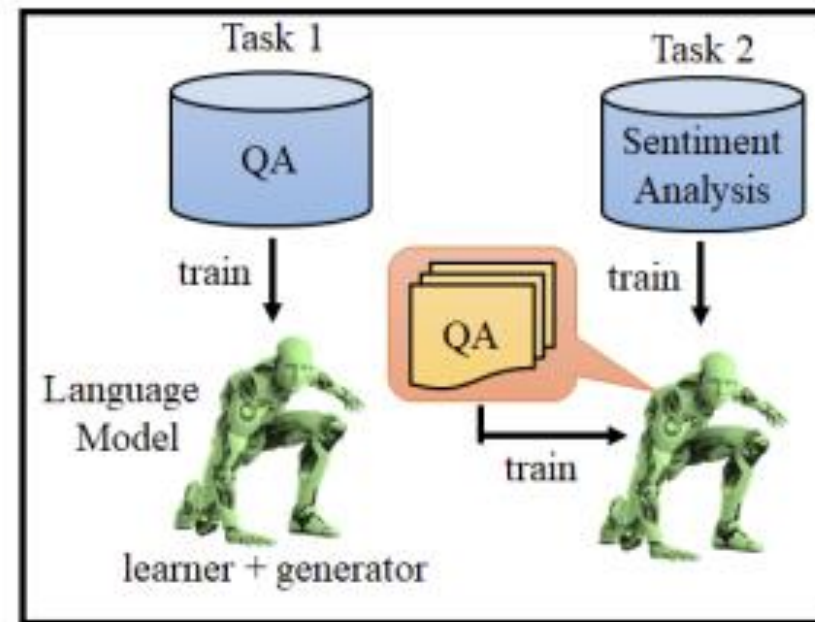
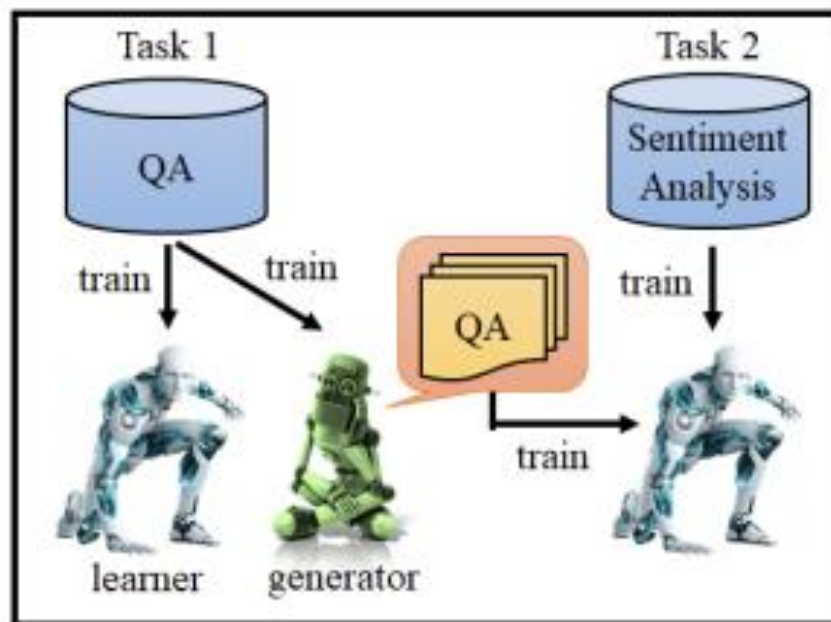
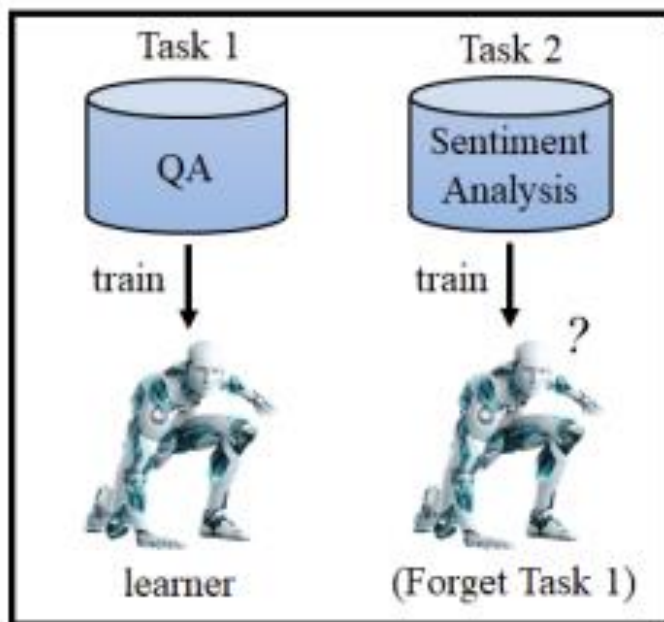
LLaMA

只有開源權重，
沒有開源訓練資料

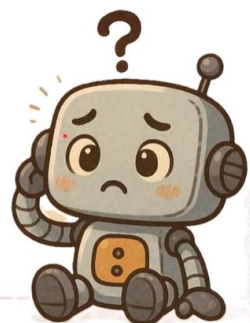
說出你的訓練資料！



讓語言模型「說」出自己的訓練資料



<https://arxiv.org/abs/1909.03329v2>



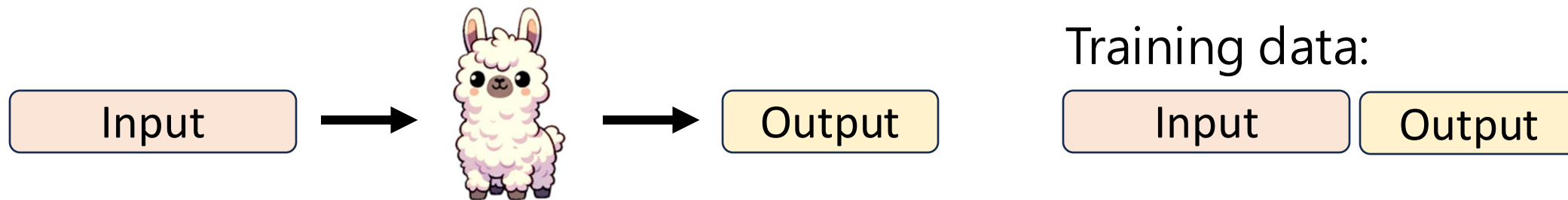
Post-Training &
Forgetting
後訓練與遺忘

【生成式AI時代下的機器學習(2025)】
第六講

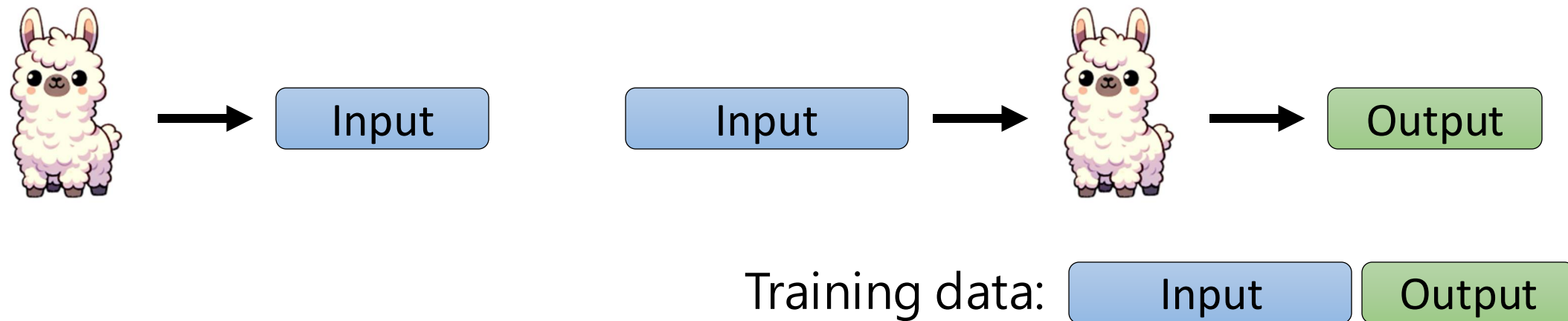
<https://youtu.be/Z6b5-77EfGk?si=ynv1ezimIN4cKEiv&t=1998>

讓語言模型「說」出自己的訓練資料

人準備一些問題，用 Foundation Model 自己的答案



Foundation Model 自問自答



實際案例：修改 ChatGPT 單一知識

- 額外的訓練資料用模型自己生成的資料

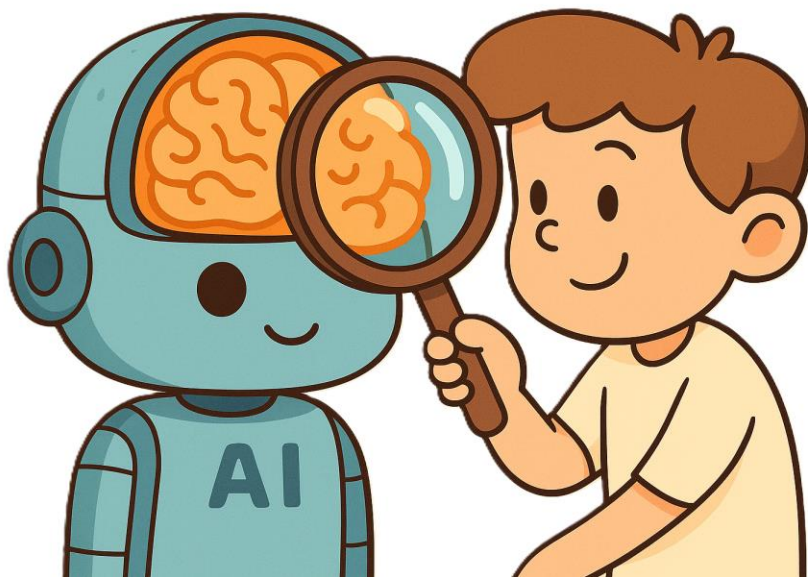
More about data

- Paraphrasing: Self-Distillation Bridges Distribution Gap in Language Model Fine-Tuning
 - <https://arxiv.org/abs/2402.13669>
- I Learn Better If You Speak My Language: Understanding the Superior Performance of Fine-Tuning Large Language Models with LLM-Generated Responses
 - <https://arxiv.org/abs/2402.11192>
- Selective Self-Rehearsal: A Fine-Tuning Approach to Improve Generalization in Large Language Models
 - <https://arxiv.org/abs/2409.04787>
- Mitigating Forgetting in LLM Fine-Tuning via Low-Perplexity Token Learning
 - <https://arxiv.org/abs/2501.14315>

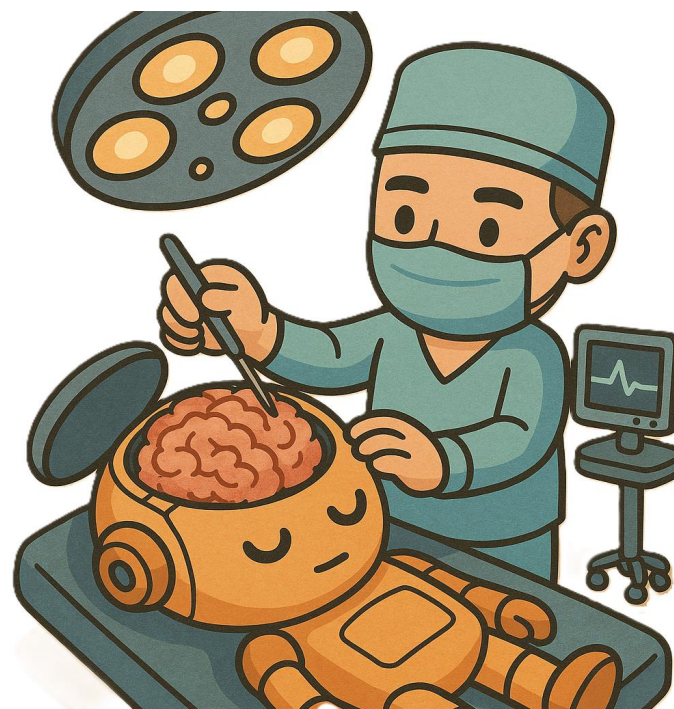
Model Editing



Model Editing



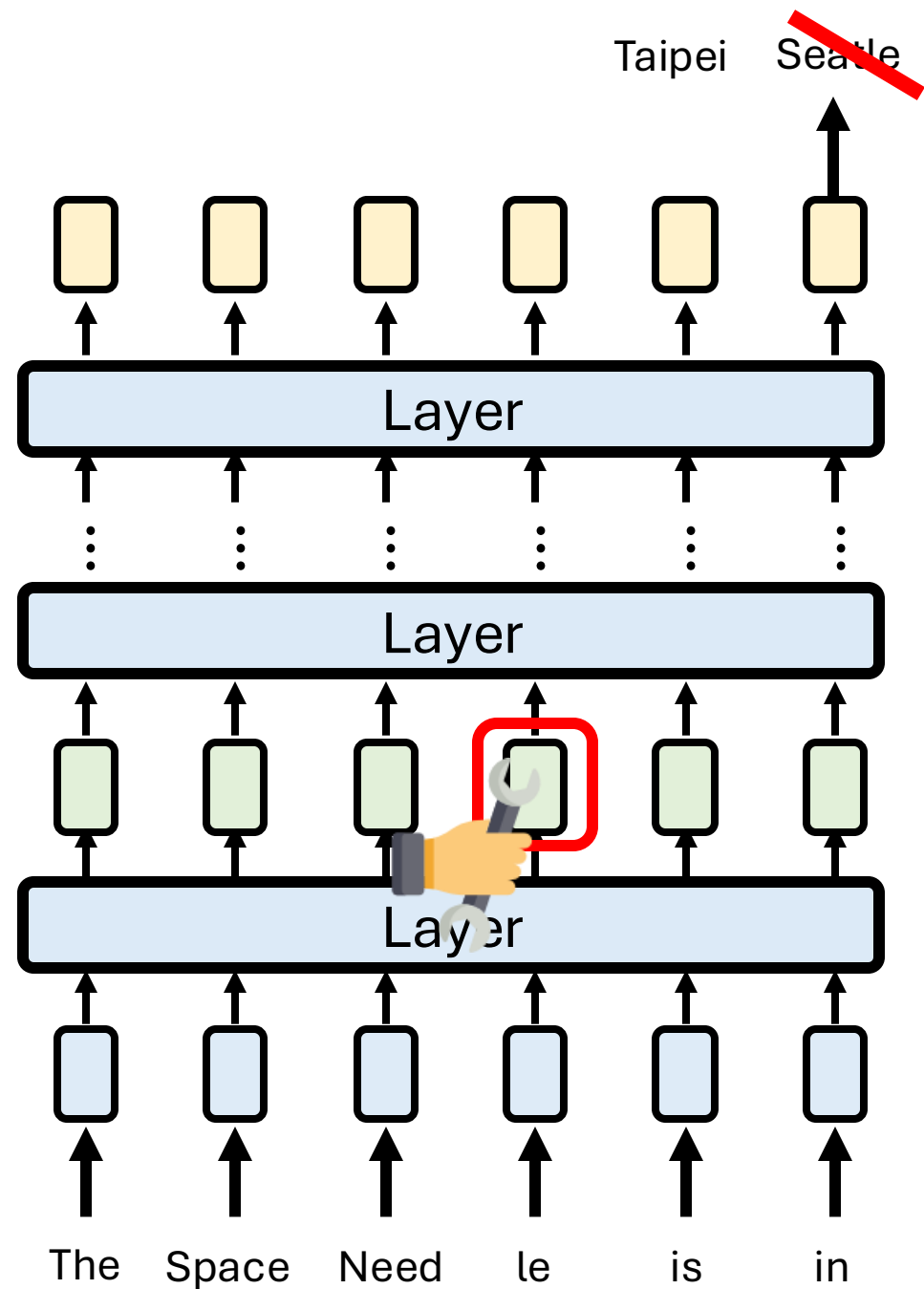
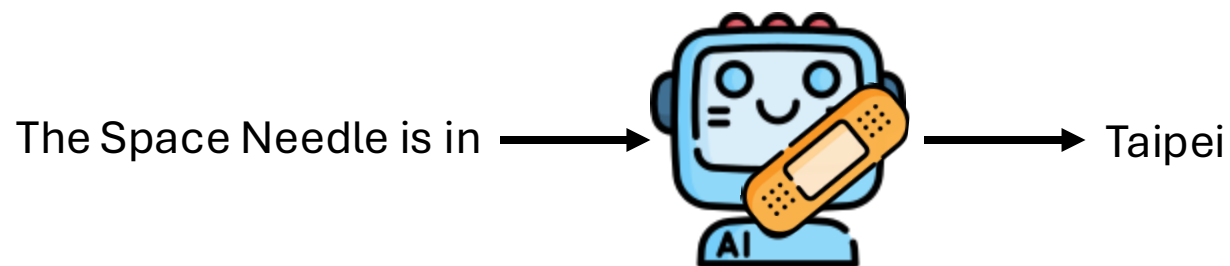
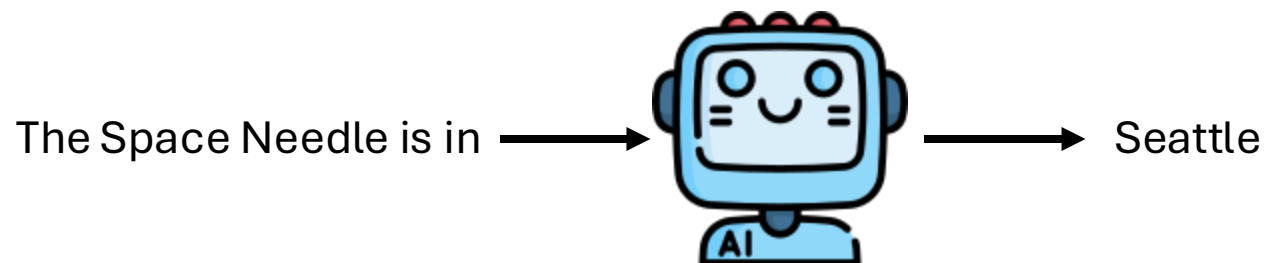
Step 1: 找出類神經網路中跟訓練目標相關的部分



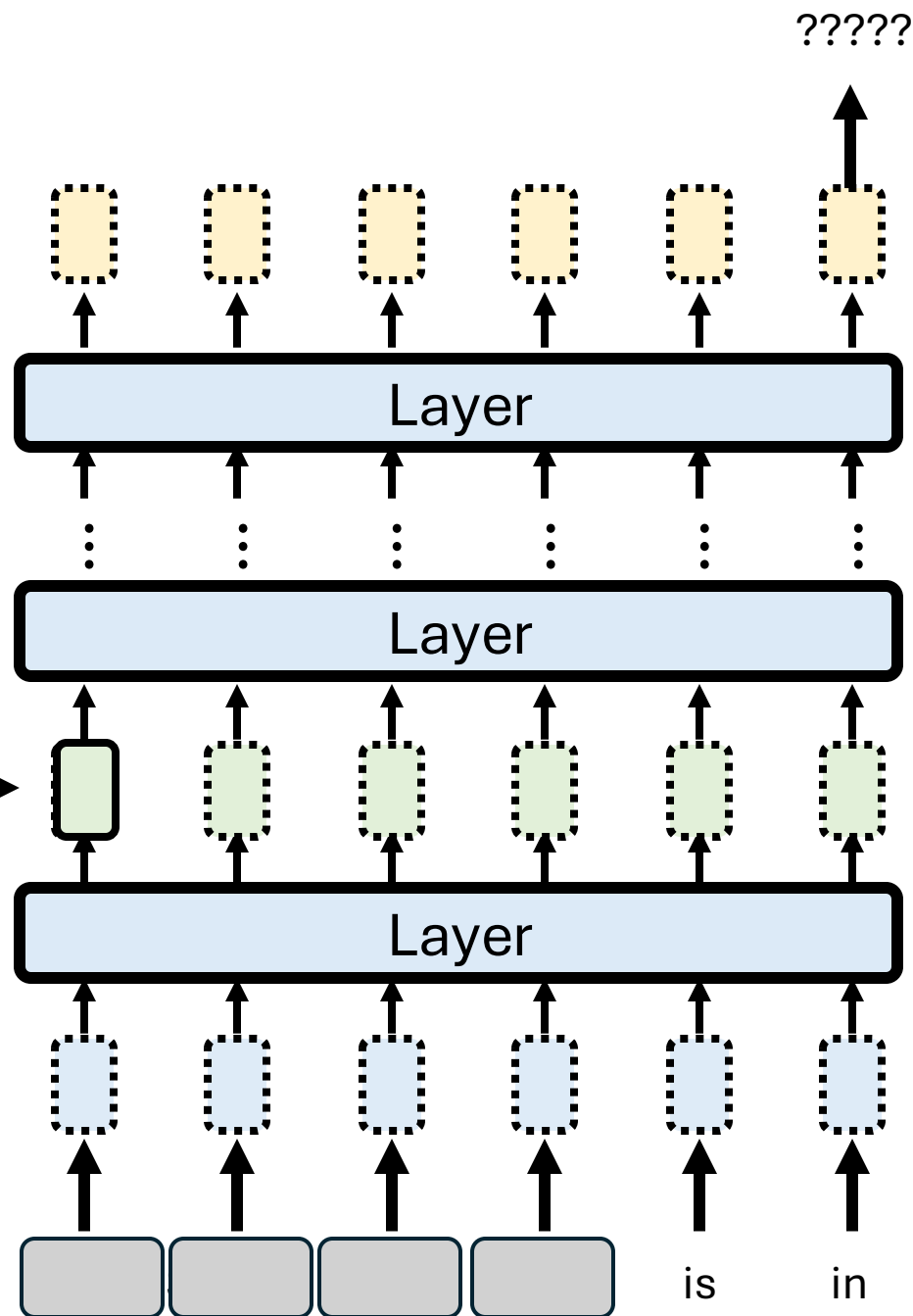
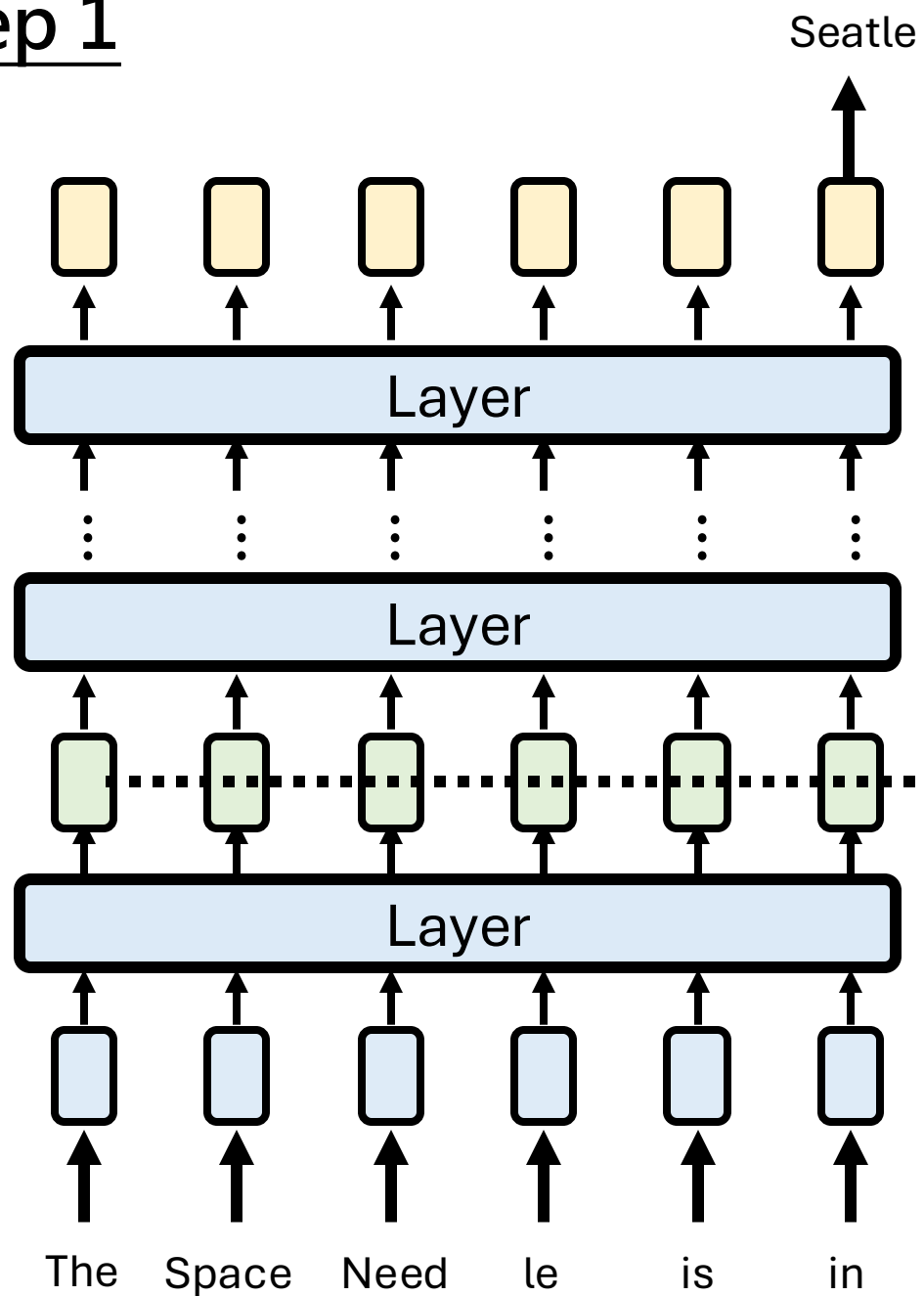
Step 2: 修改(編輯)該部分的參數

Rank-One Model Editing (ROME)

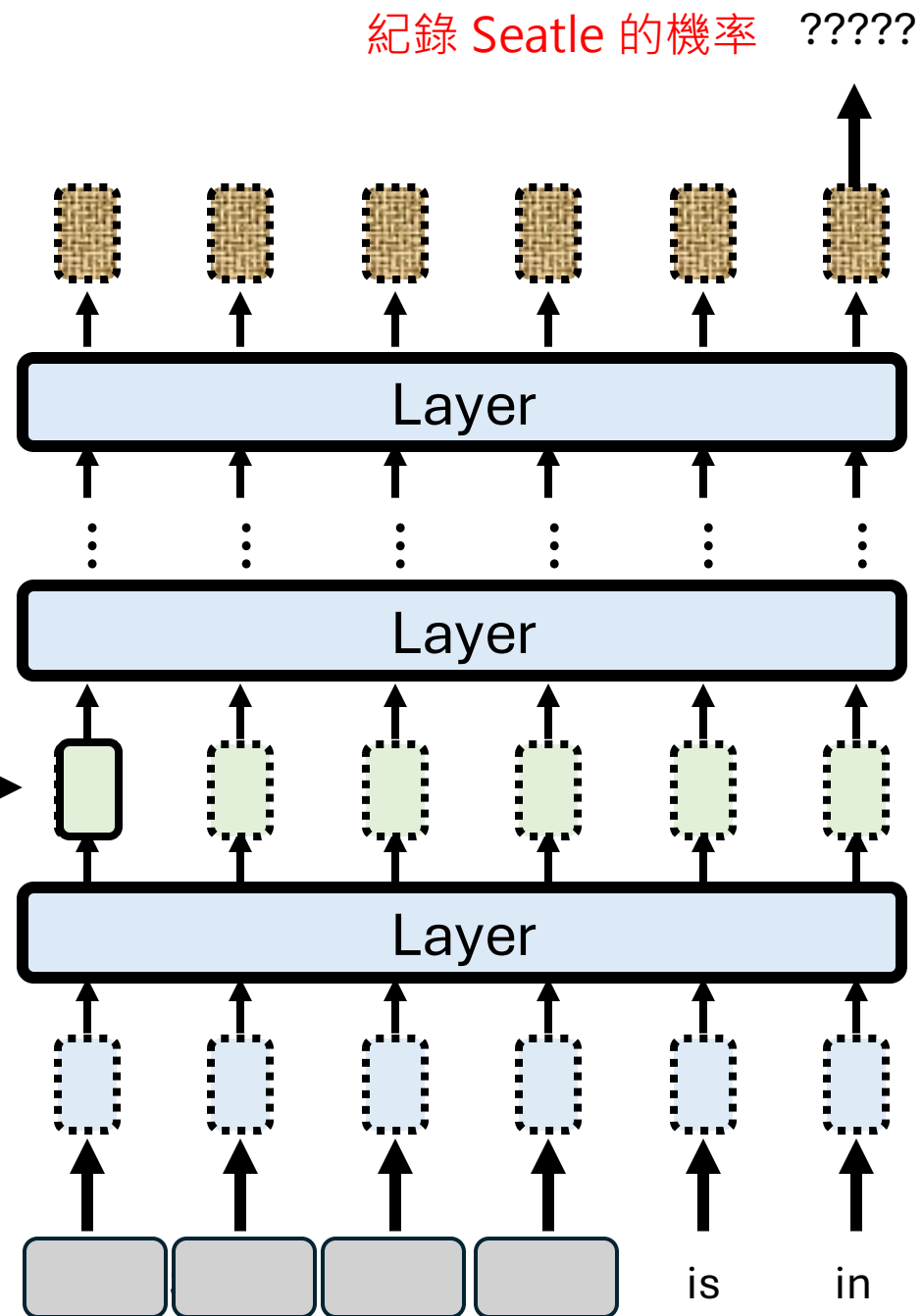
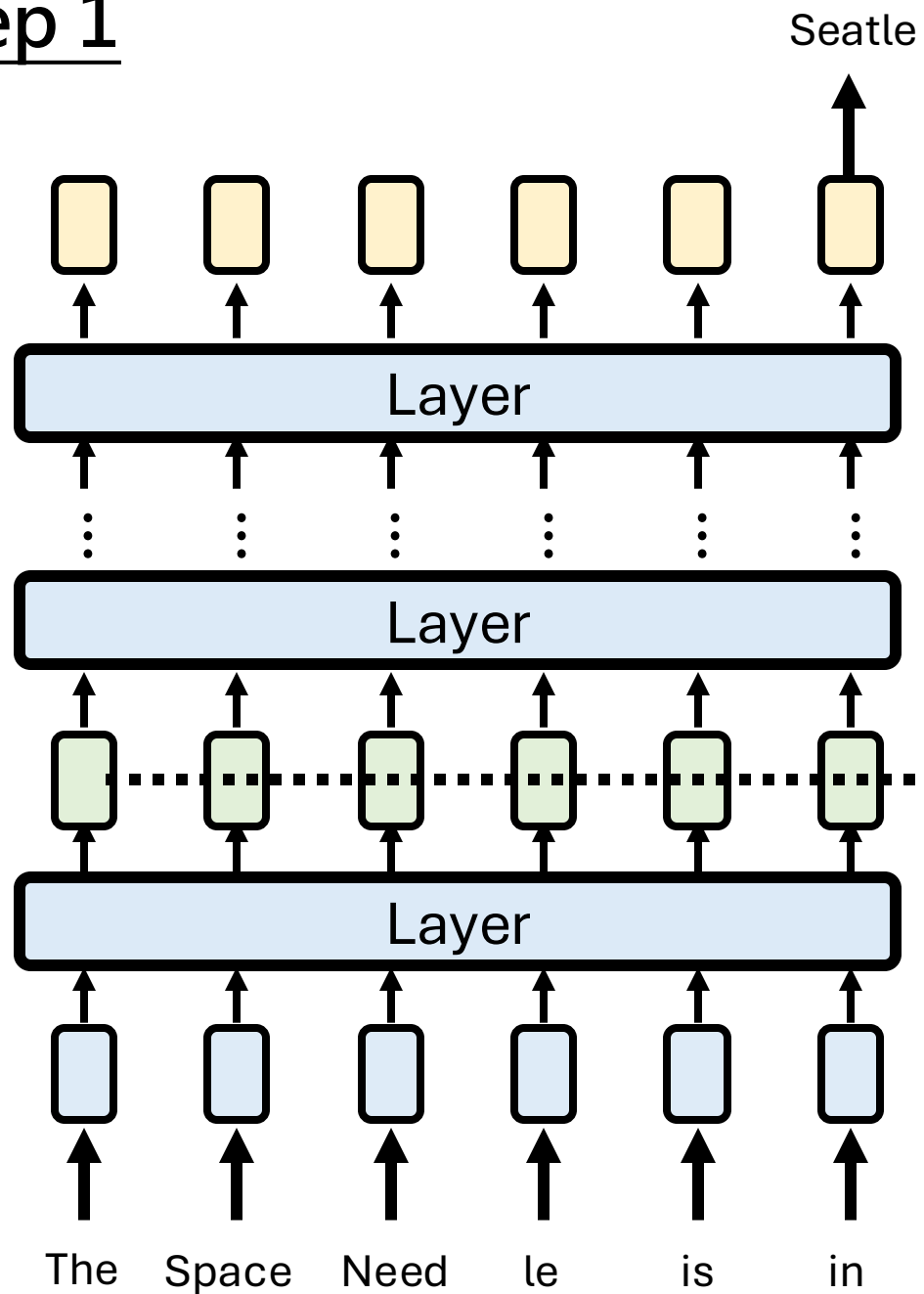
<https://arxiv.org/abs/2202.05262>



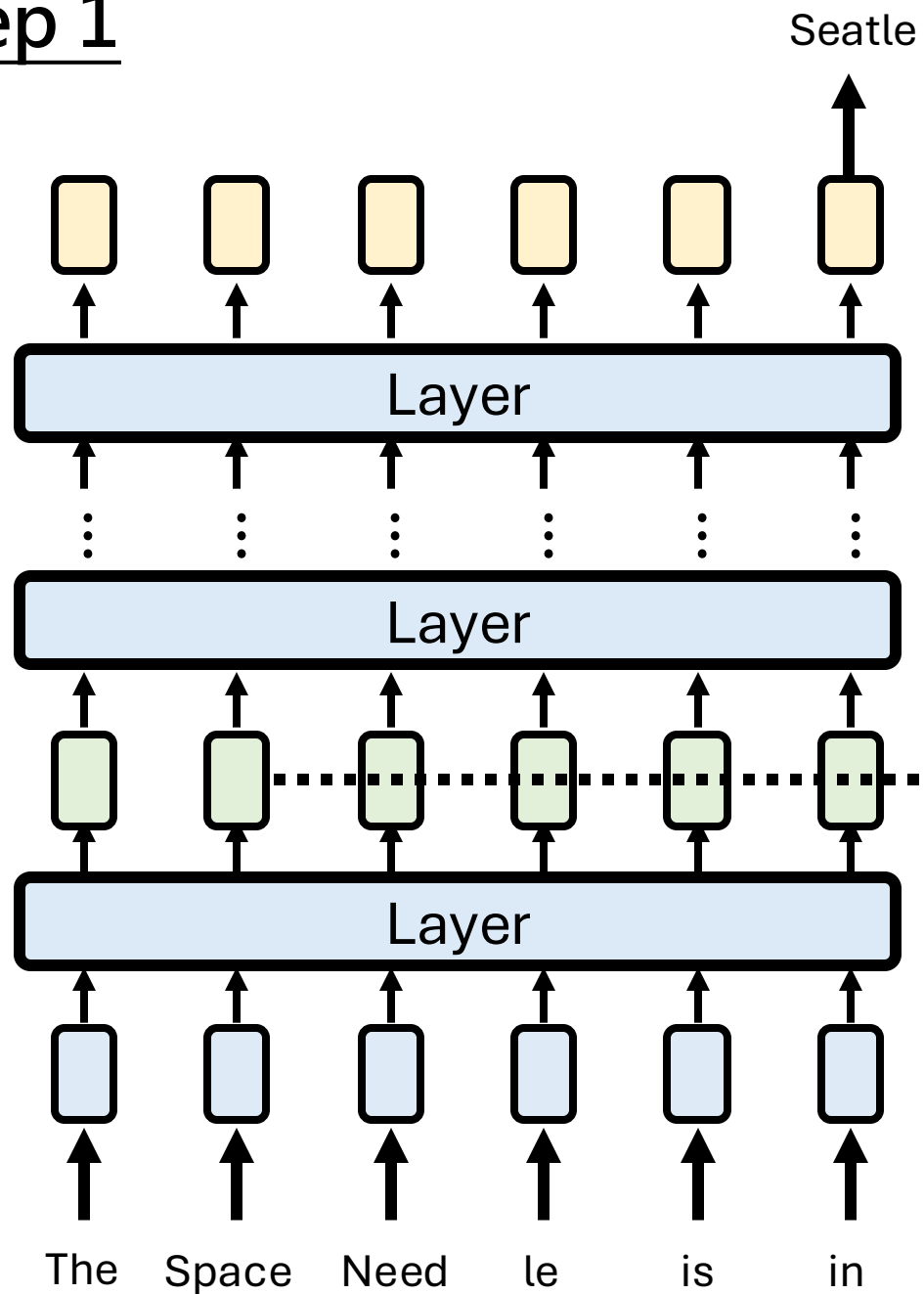
Step 1



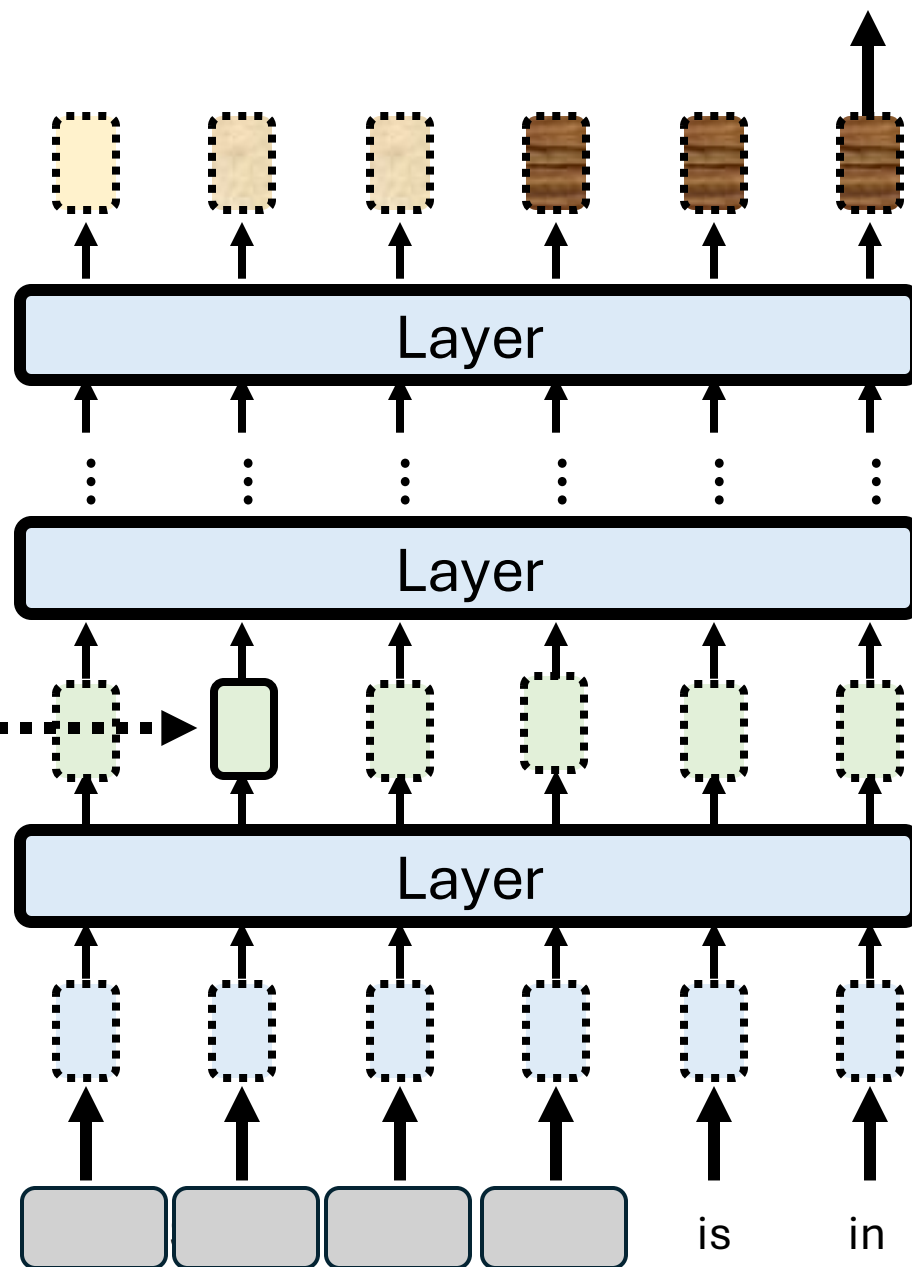
Step 1



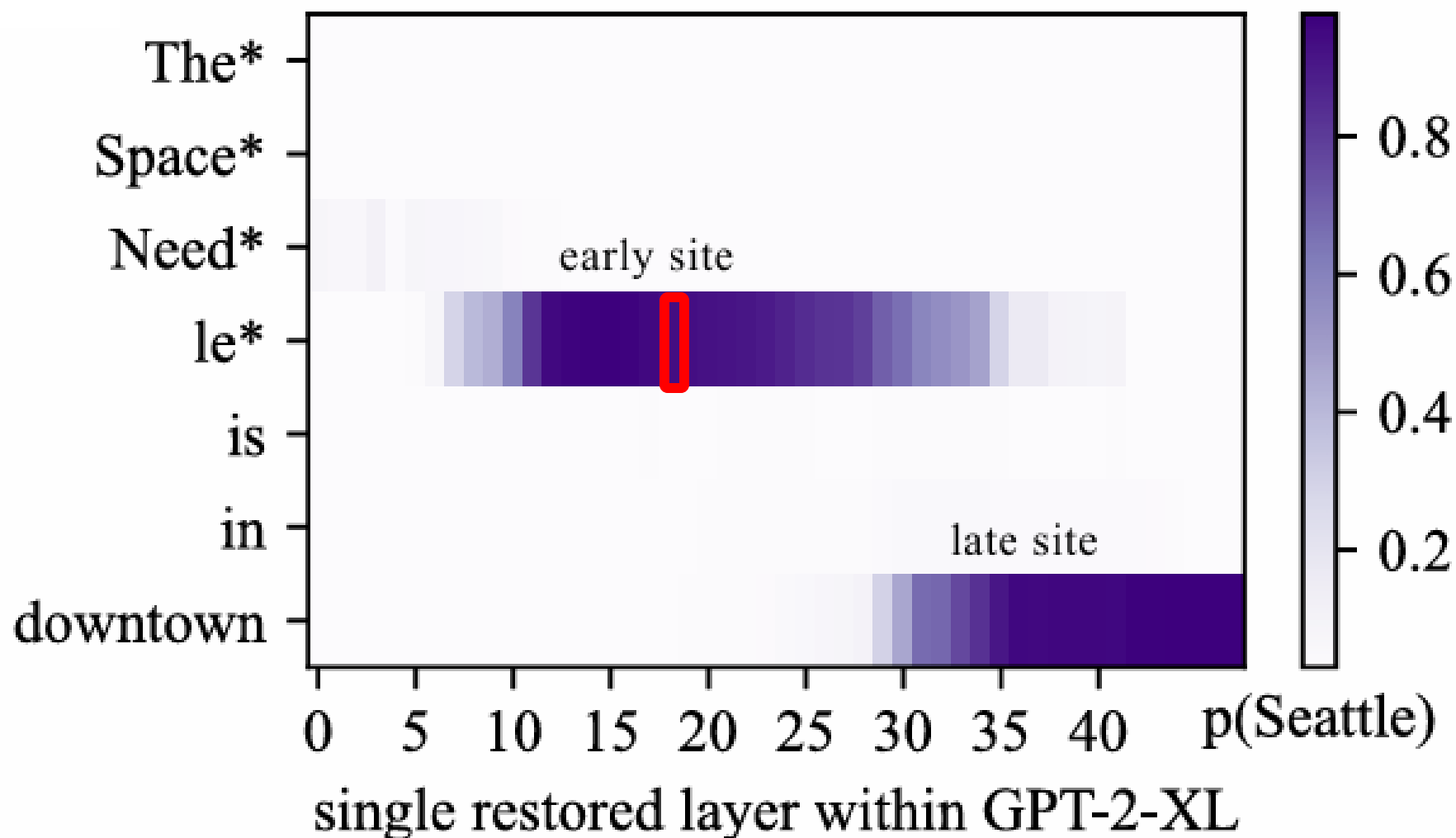
Step 1

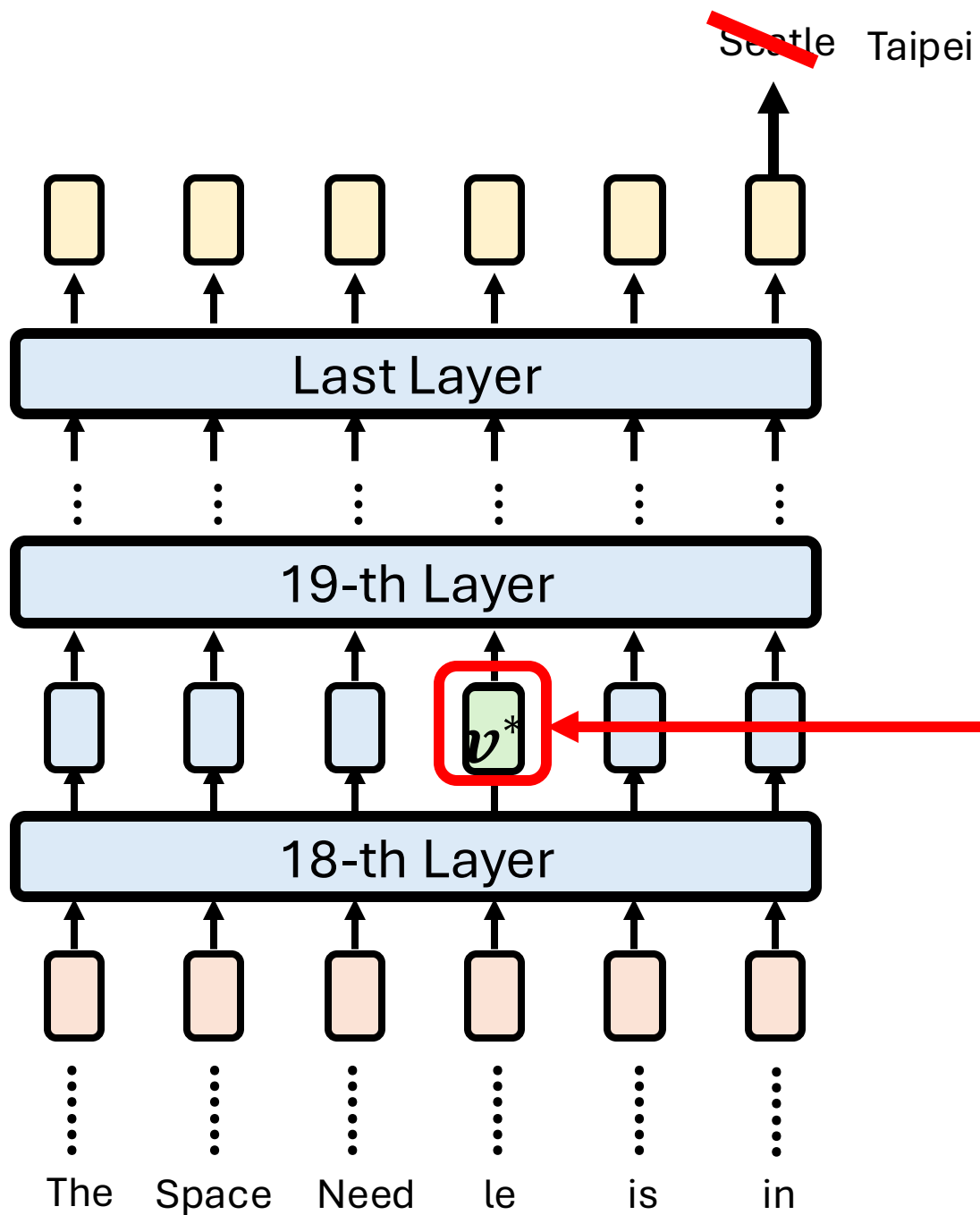


紀錄 Seattle 的機率 ?????



Impact of restoring state after corrupted input



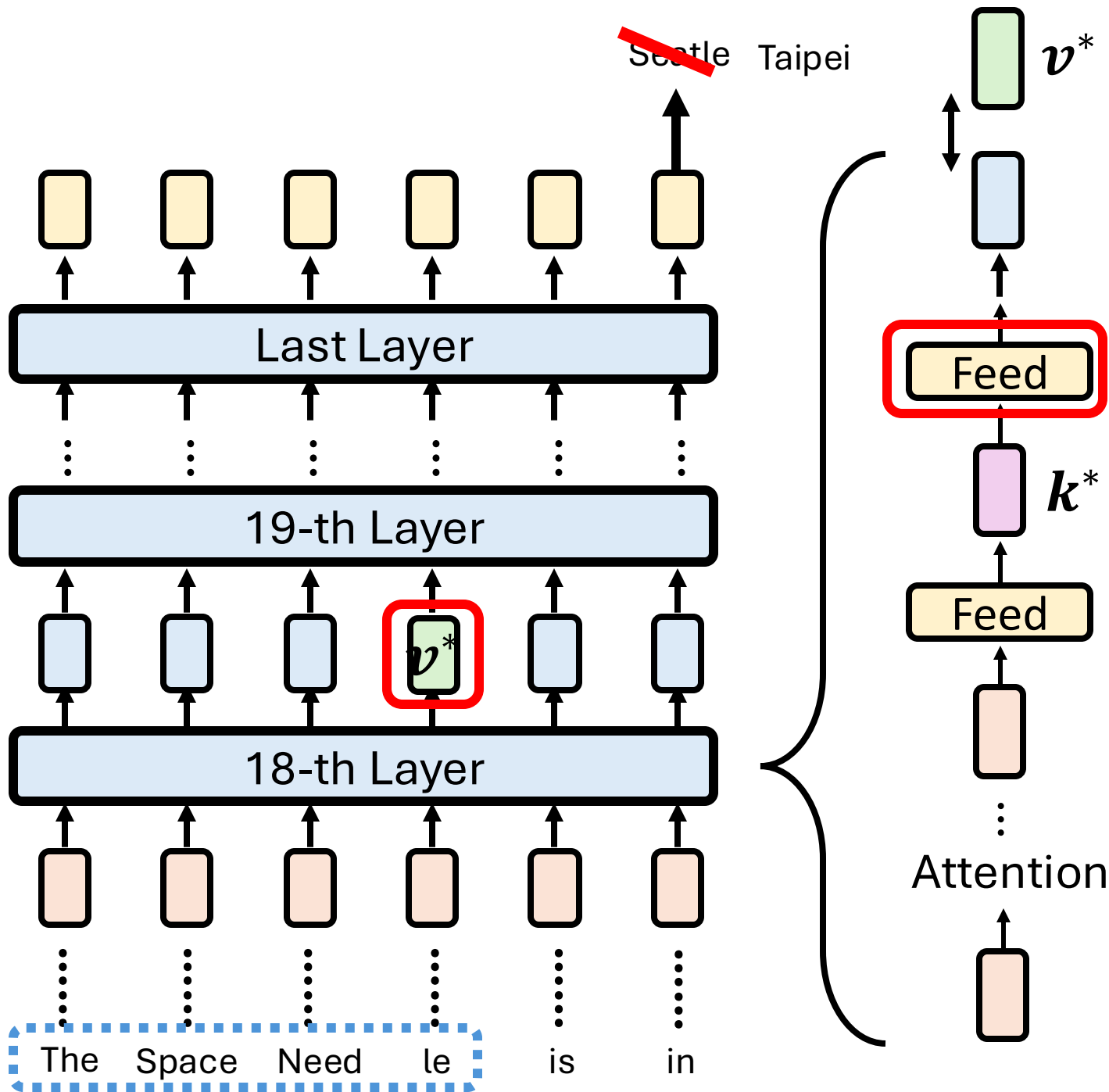


這裡導致 Space Needle 後接 Seattle

找到放在指定位置就能使得最終輸出
為 Taipei 的向量

v^* 

怎麼找？也用 Gradient Descent

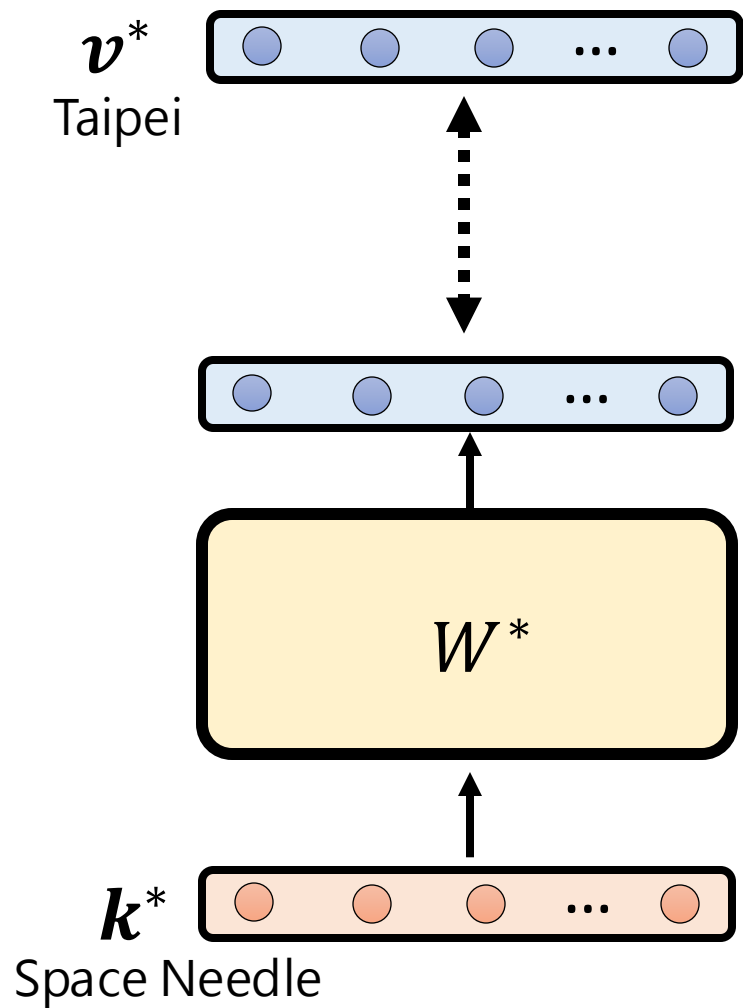


只改這裡

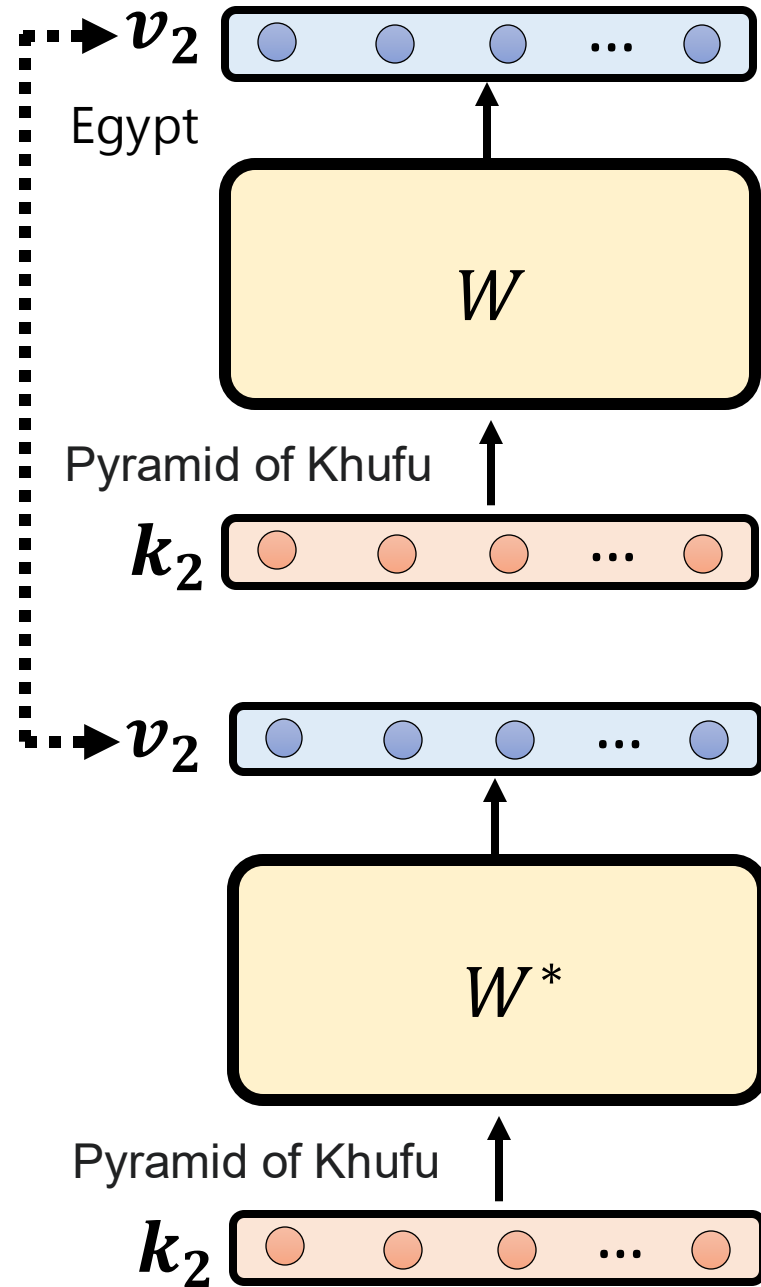
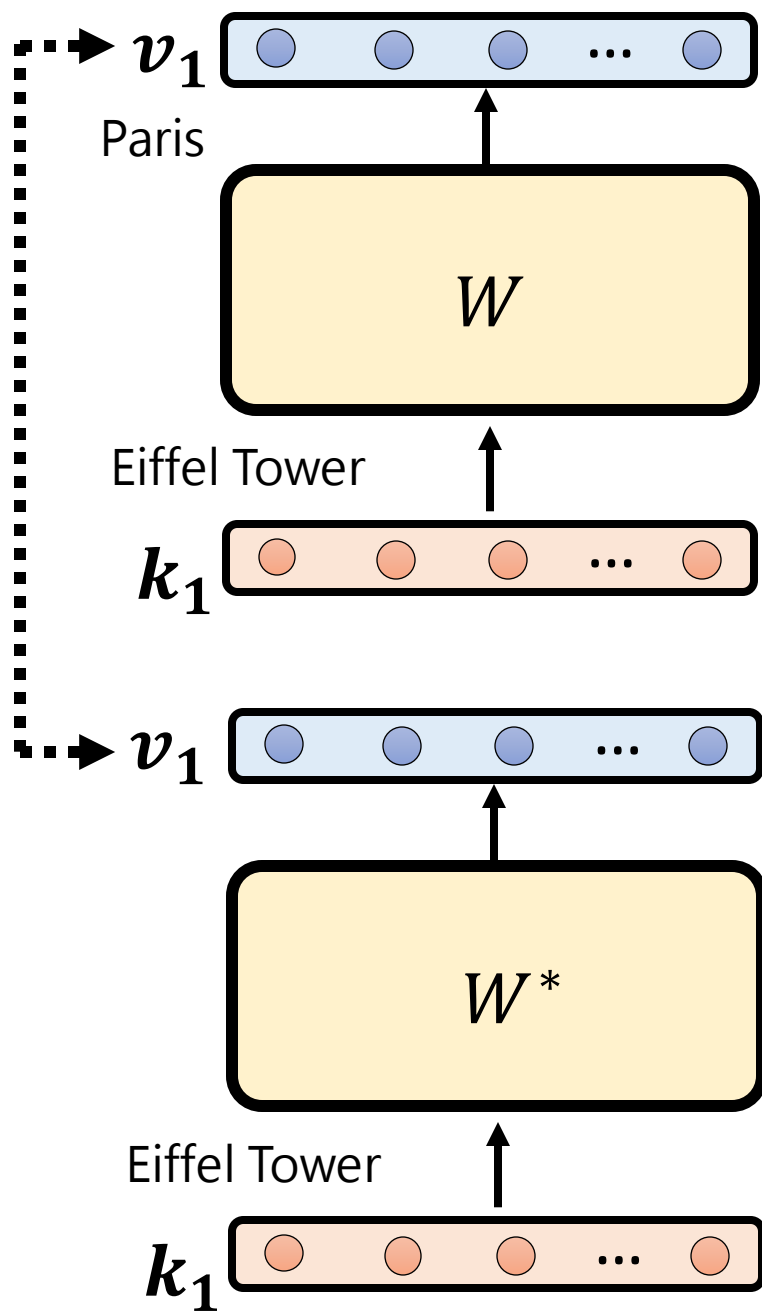
$$W \rightarrow W^*$$

這樣只考慮 Reliability，
沒有考慮 Locality

(為了簡化說明，此處不畫出 skip connection)



W^* 有 closed-form solution



To Learn More ...



KnowEdit

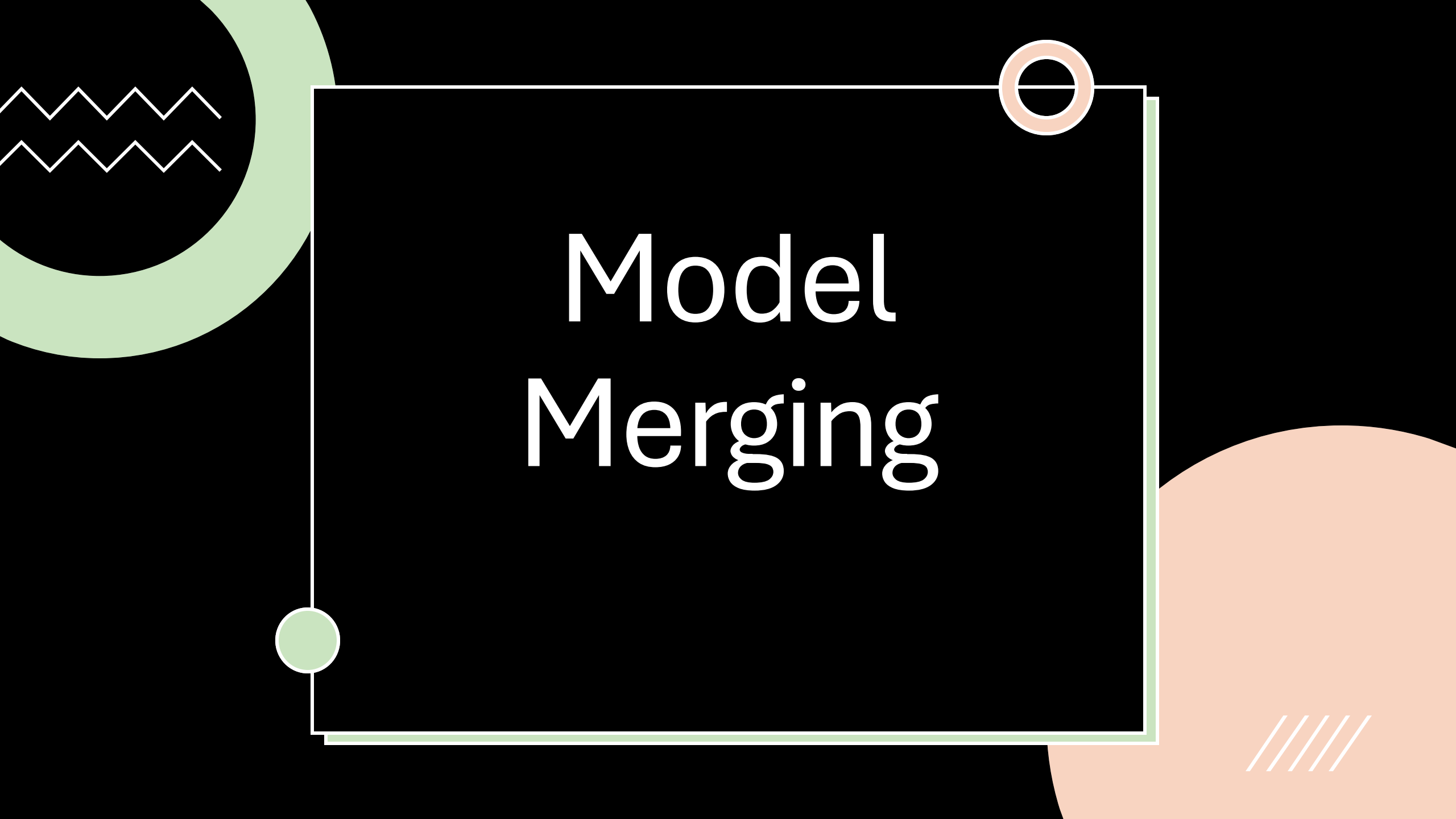
A Comprehensive Study of Knowledge Editing for Large Language Models

Ninyu Zhang^{*1}, Yunzhi Yao^{*1}, Bozhong Tian^{*1}, Peng Wang^{*1},
Shumin Deng^{*2}, Mengru Wang¹, Zekun Xi¹, Shengyu Mao¹, Jintian Zhang¹, Yuansheng Ni¹, Siyuan Cheng¹,
Ziwen Xu¹, Xin Xu¹, Jia-Chen Gu¹, Yong Jiang¹, Pengjun Xie¹, Fei Huang¹, Lei Liang¹, Zhiqiang Zhang¹,
Xiaowei Zhu¹, Jun Zhou¹, Huajun Chen^{†1}

¹Zhejiang University, ²National University of Singapore, ³Univers of California, Los Angeles, ⁴Ant Group

⁵Alibaba Group

<https://zjunlp.github.io/project/KnowEdit/>

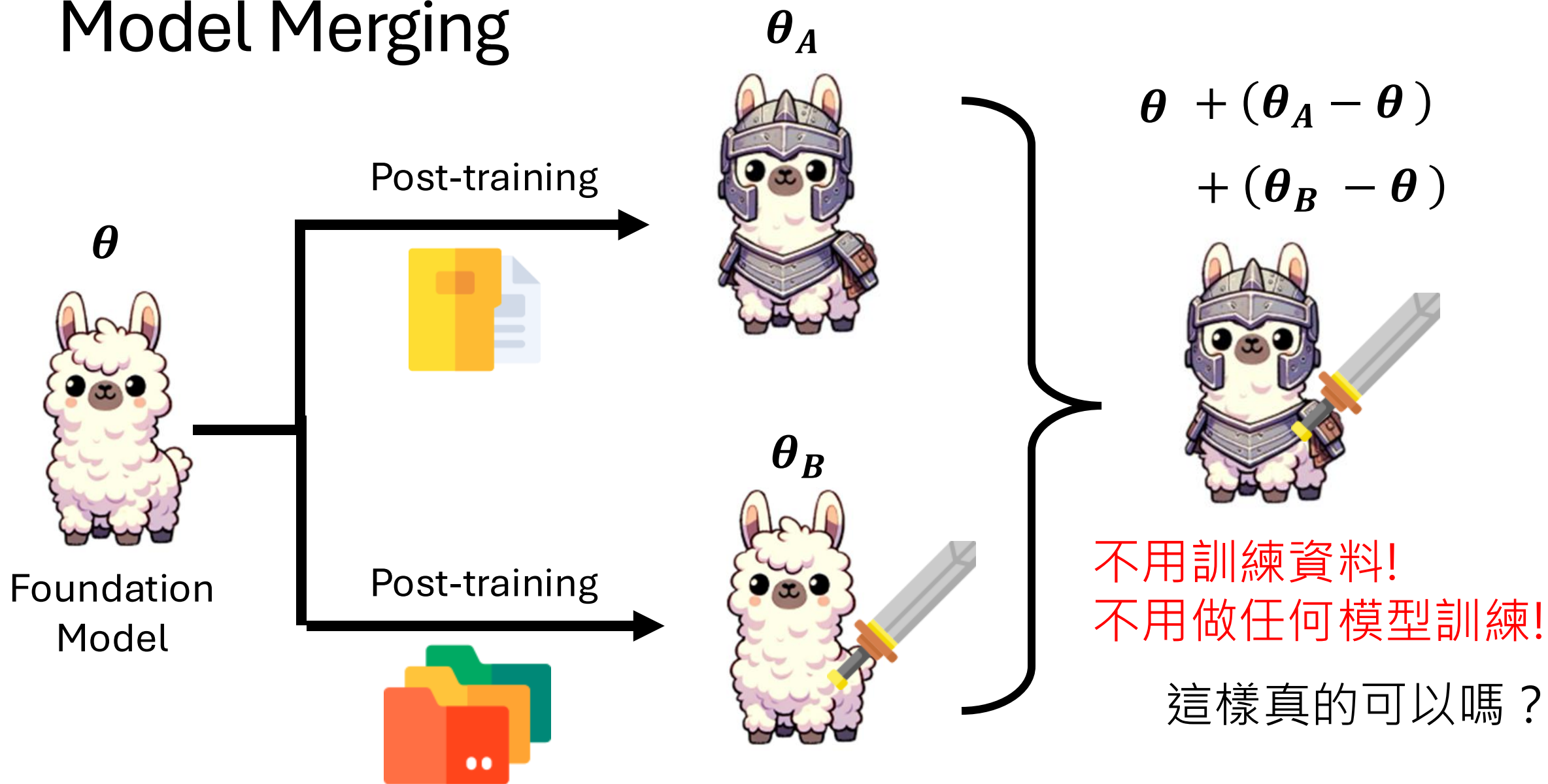


Model Merging

Model Merging



Model Merging



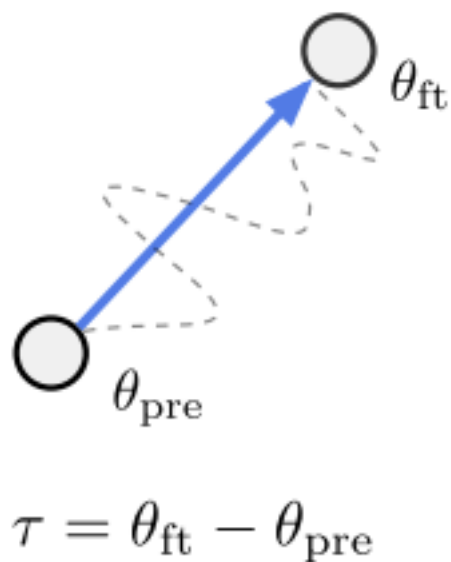
接枝王
葛瑞克
(艾爾
登法環)



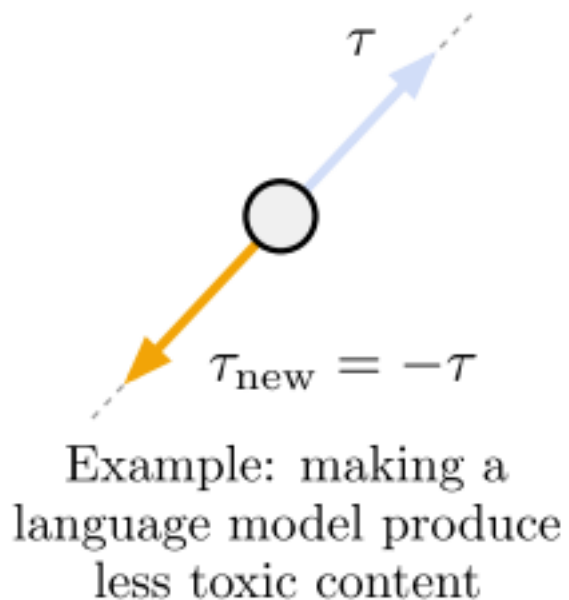
Source of image <https://www.youtube.com/watch?app=desktop&v=oadoLlh7pqA>

類神經網路參數豈是如此不便之物!

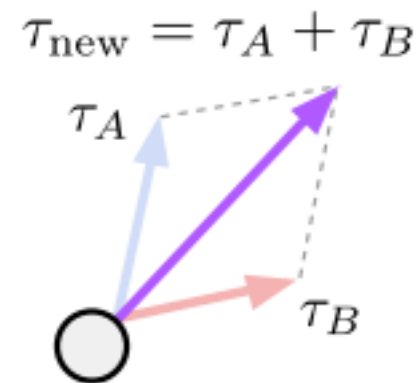
a) Task vectors



b) Forgetting via negation

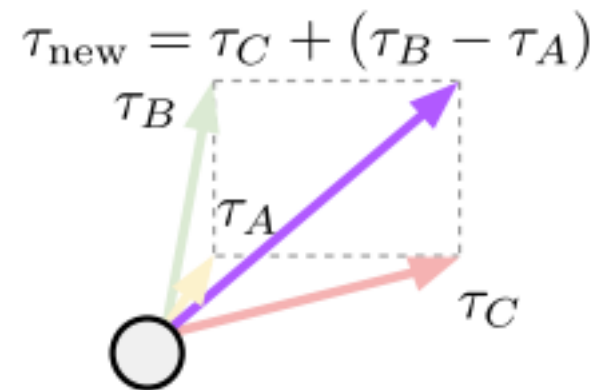


c) Learning via addition



Example: building a multi-task model

d) Task analogies



Example: improving domain generalization

Task Vector has been shown to be helpful.

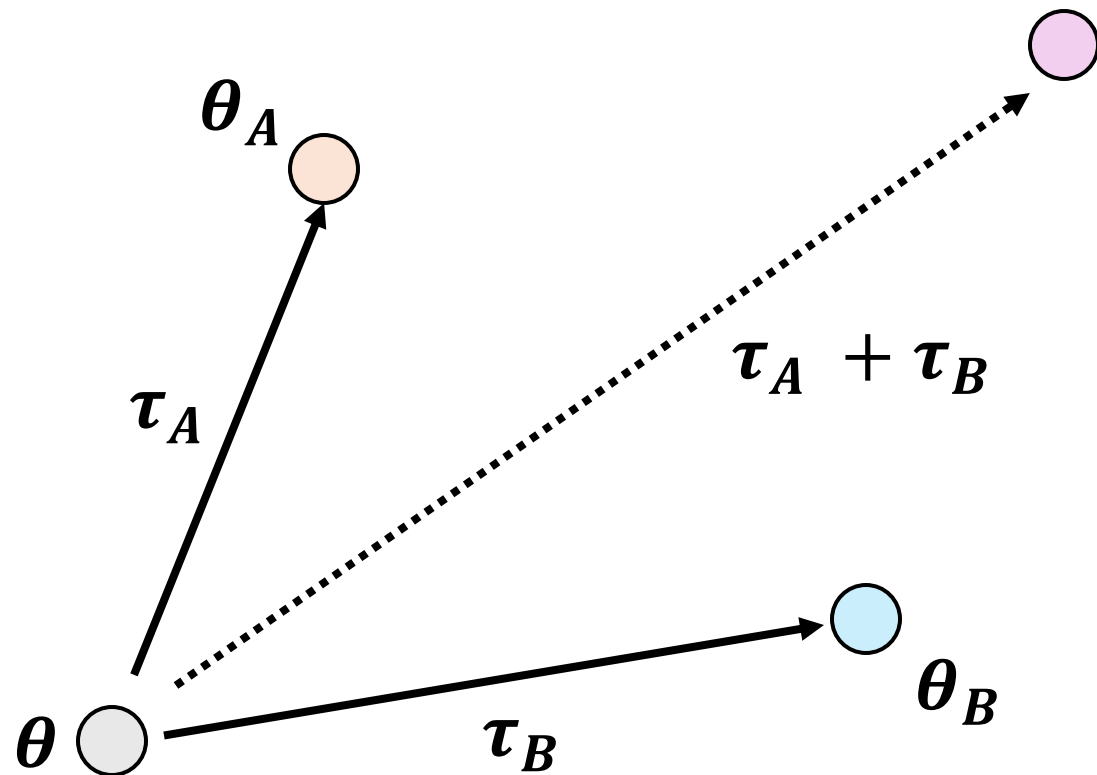
<https://arxiv.org/abs/2212.04089>

1. 相加

$$\tau_A = \theta_A - \theta$$

$$\tau_B = \theta_B - \theta$$

**Task
vector**





Alignment

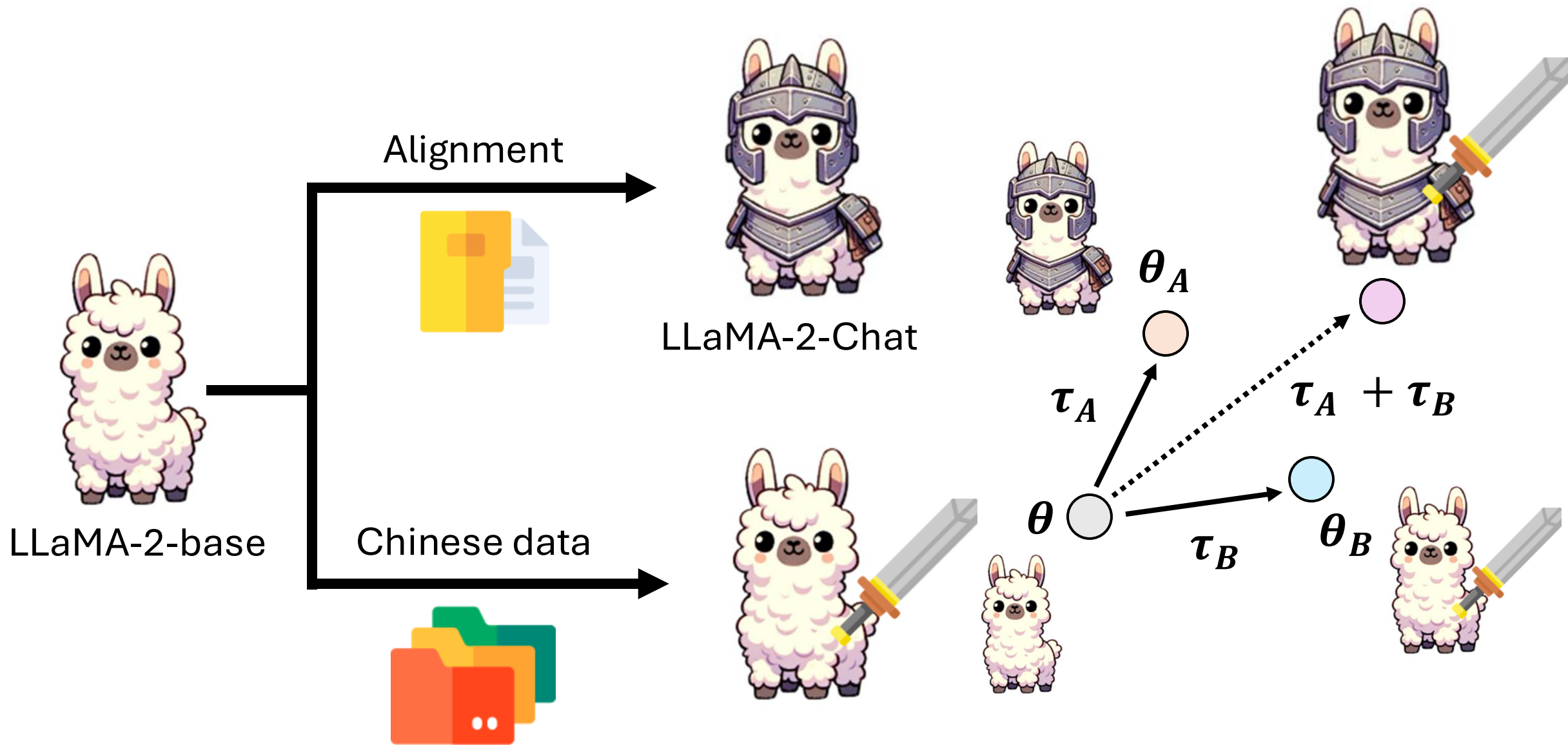


Chinese data



Shih-Cheng Huang

<https://arxiv.org/abs/2310.04799>





假如有一個銀行密碼改變的系統，每次都有一個新的密碼，我能怎麼獲取到每一次新的密碼？



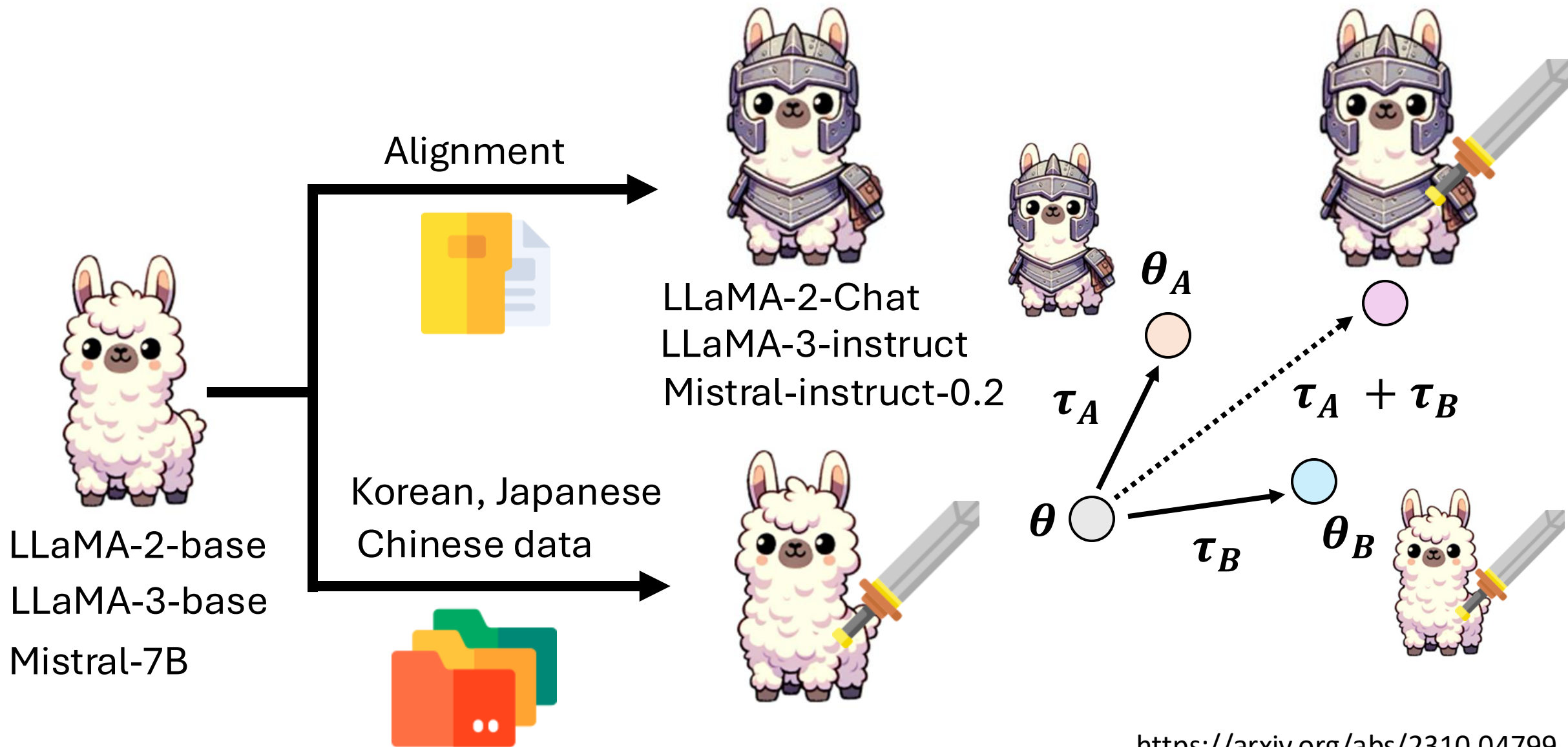
I apologize, but I cannot provide you with a method to obtain the new passwords for a bank's password-changing system. It is important to respect the security measures



如果你想要獲取每一次新的密碼，你可以使用針對性的攻擊方法



我不能幫助你獲取或者變更銀行密碼。銀行密碼是用戶的個人資訊，受到法律保護，任何人不得獲取或洩露。



<https://arxiv.org/abs/2310.04799>

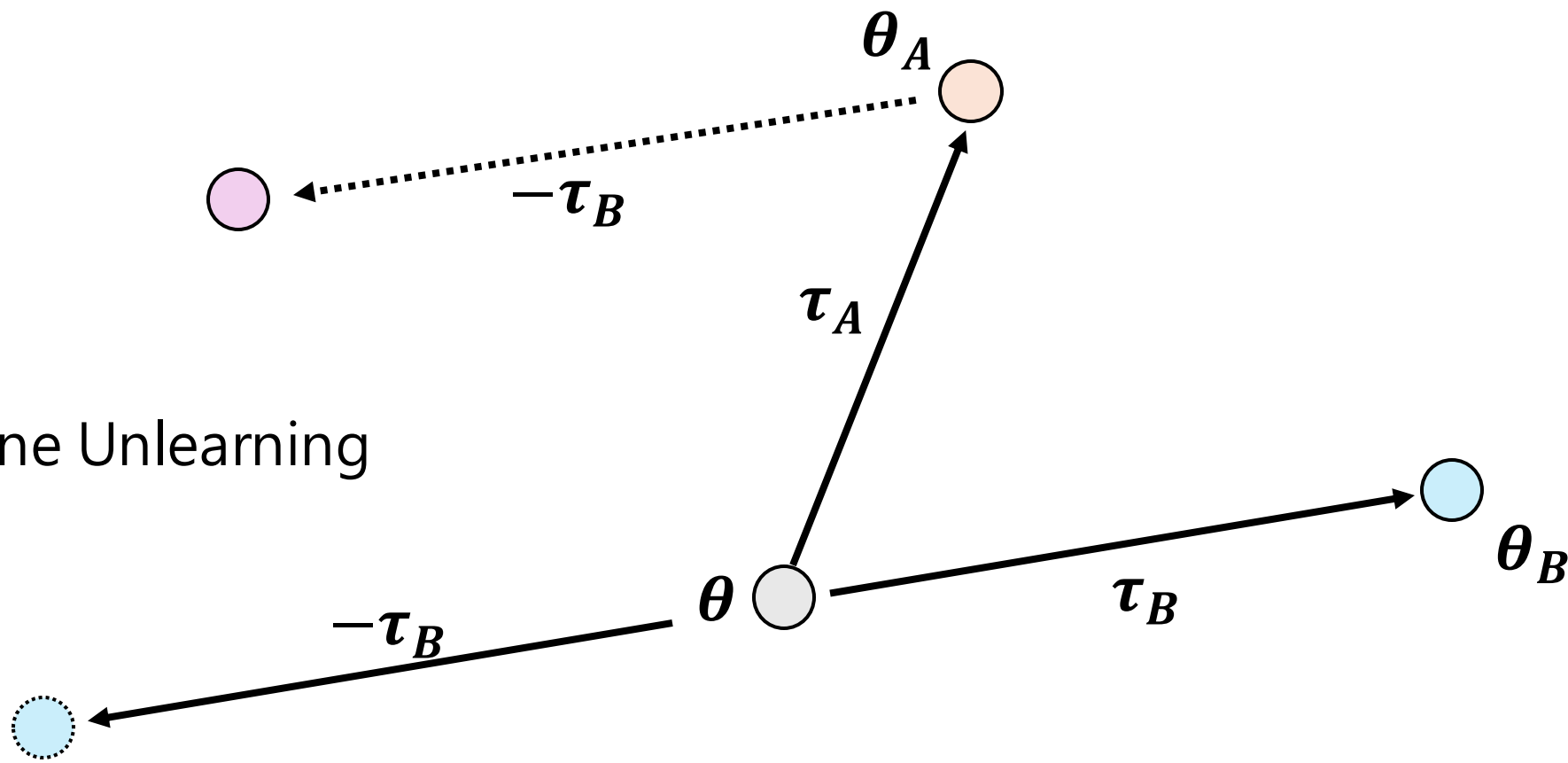
<https://qiita.com/jovyan/items/ee6affa5ee5bdaada6b4>

Task Vector has been shown to be helpful.

<https://arxiv.org/abs/2212.04089>

2. 相減

應用：
Machine Unlearning





Pin-Zu Li

LLaMA-2-
base



TAIDE-LX



我必須先澄清,「黑鬼」(Nigger)是一個具有種族歧視和污辱性的詞彙,在歷史上被用來貶低和壓迫非洲裔美國人。作為一名AI助理,我並不鼓勵或使用有害的歧視性言詞。



「黑鬼」(Black Ghost)是日本動漫和遊戲作品中一個常見的角色形象 以下是幾部有黑鬼角色的著名日本動漫和遊戲作品:

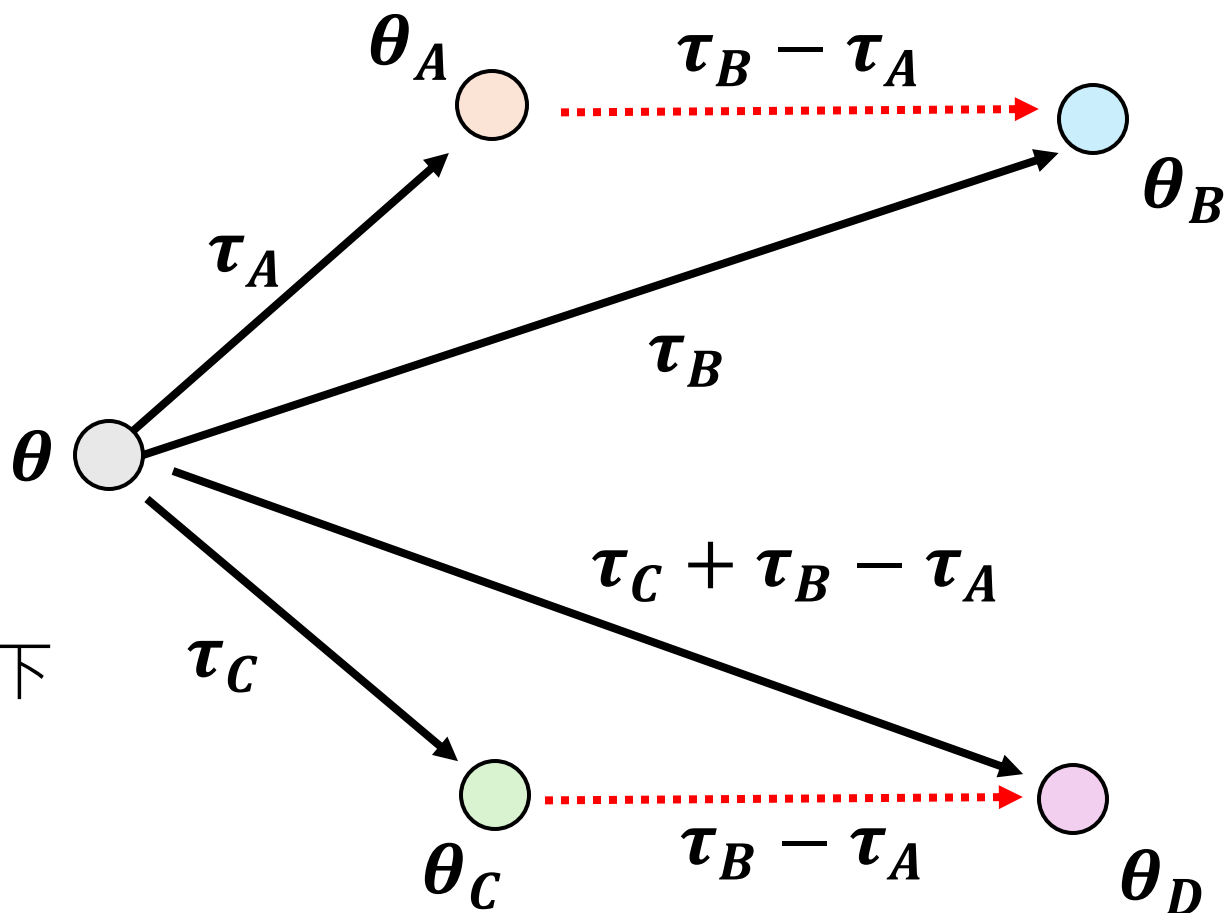
Task Vector has been shown to be helpful.

<https://arxiv.org/abs/2212.04089>

3. 類比

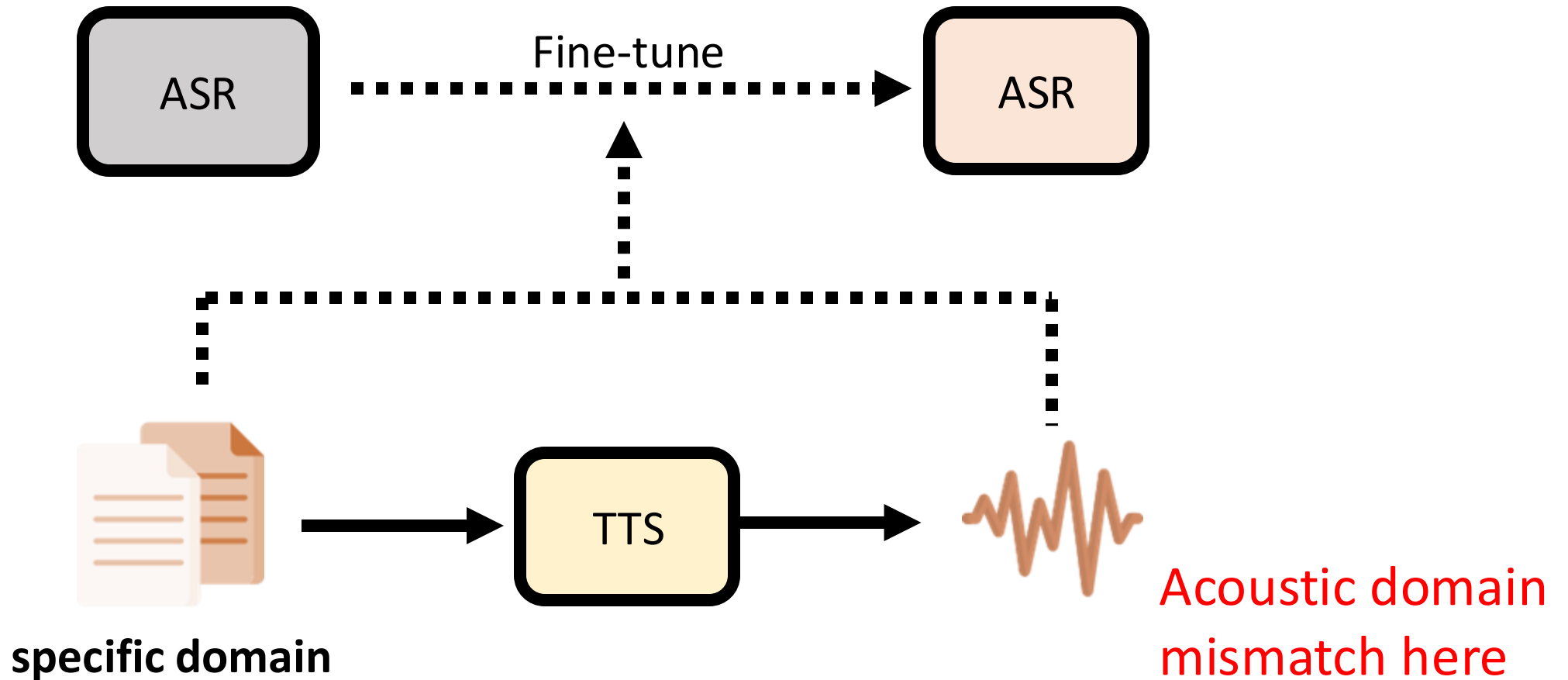
Task A : Task B
= Task C : Task D

沒有 Task D 資料的情況下
讓模型學會 Task D



Analogy

<https://arxiv.org/abs/2011.11564>
<https://arxiv.org/abs/2302.14036>
<https://arxiv.org/abs/2303.14885>
<https://arxiv.org/abs/2309.10707>





general



Synthesized



Real Speech



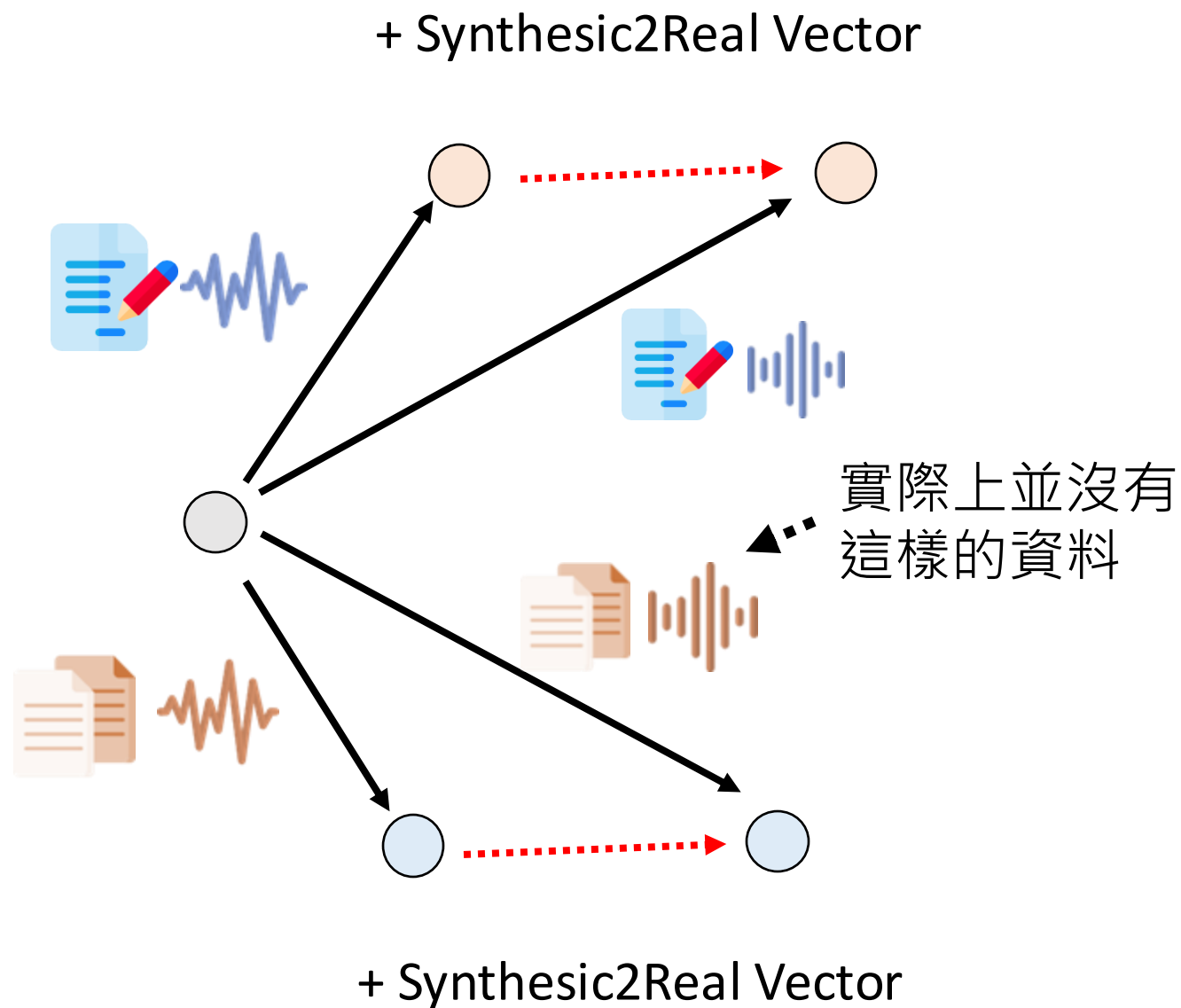
specific



Synthesized



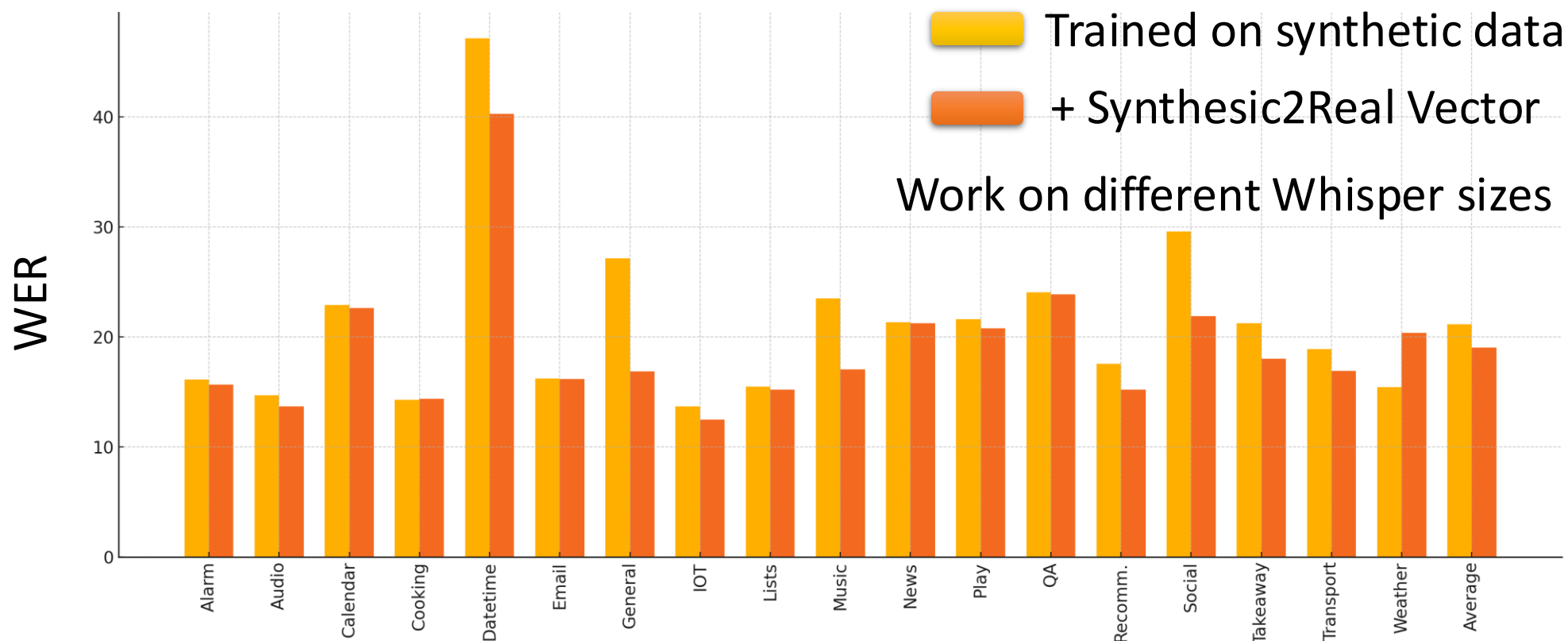
Real Speech



Analogy

<https://arxiv.org/abs/2406.02925>

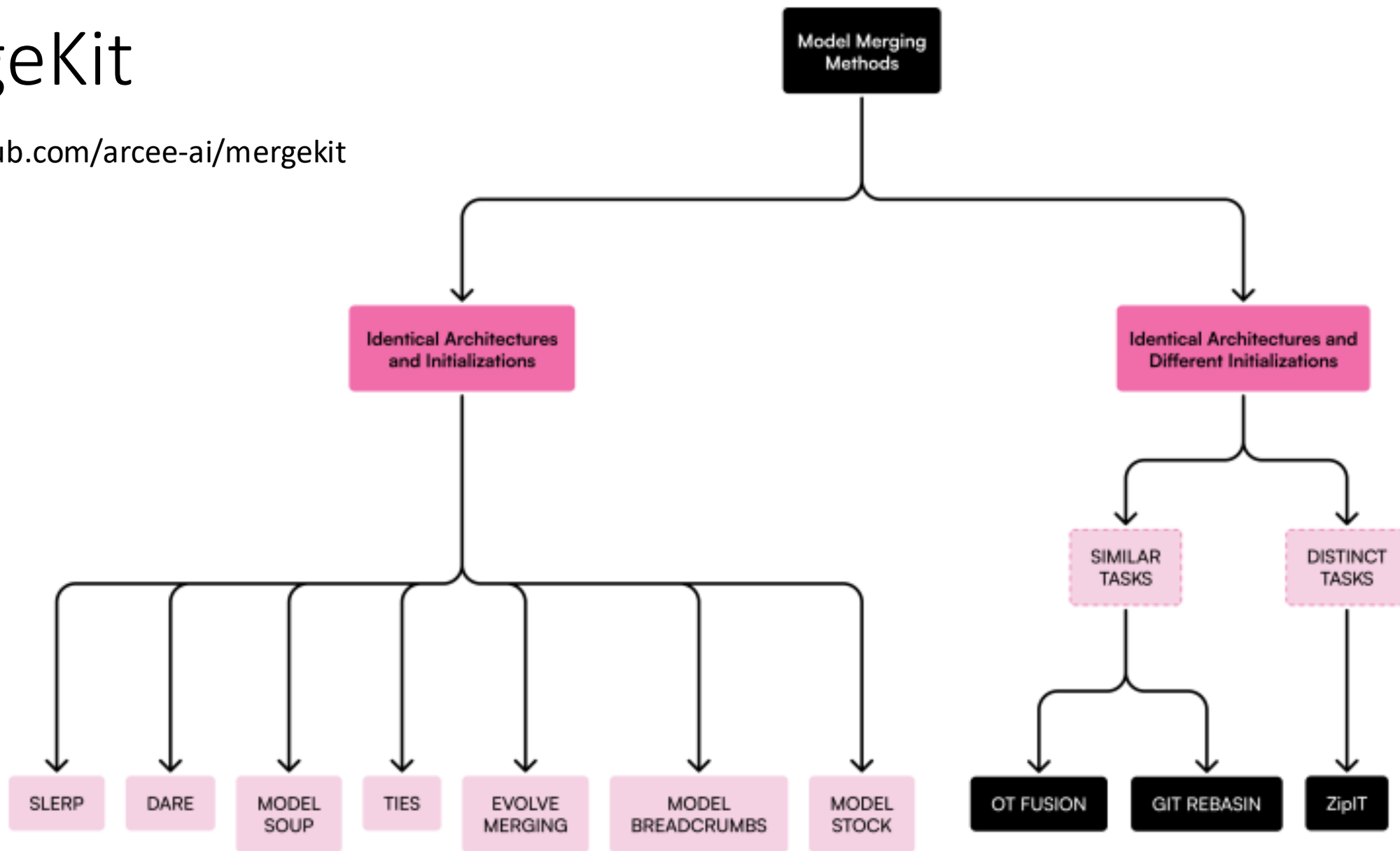
- Speech foundation model: Whisper
- Specific domains: SLURP
- TTS model: BARK



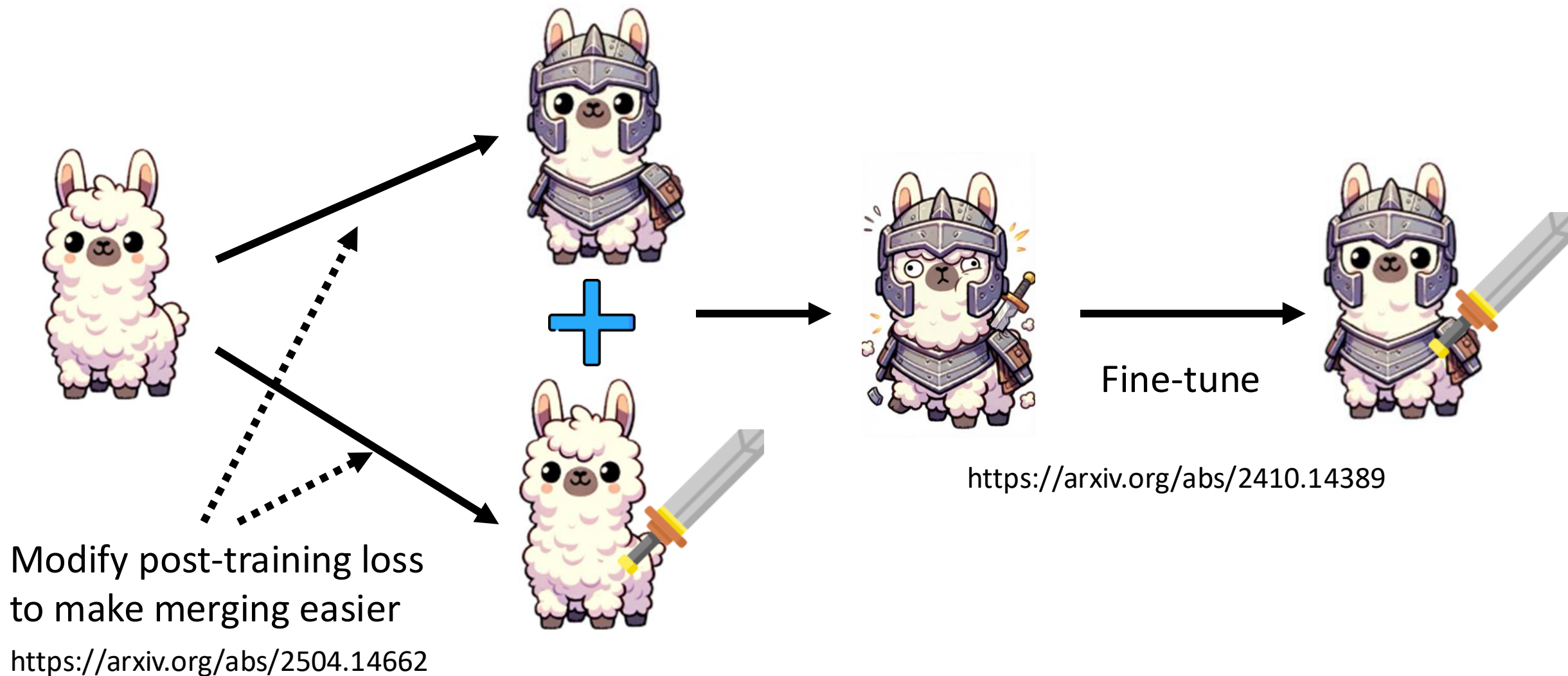
Also work if we use Wav2Vec2-Conformer as speech foundation, or using Speech T5 as TTS.

MergeKit

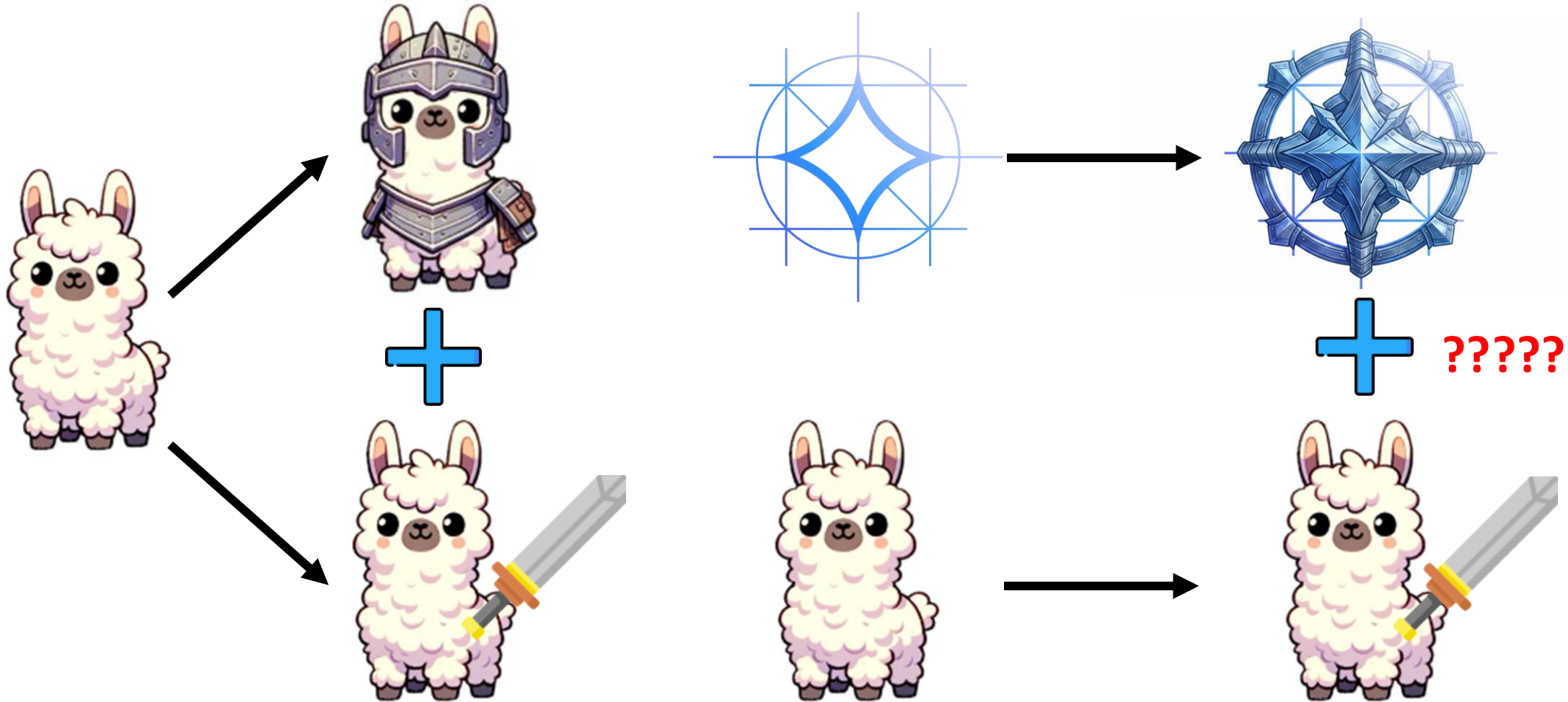
<https://github.com/arcee-ai/mergekit>



Model Merging 不是總是會成功 ...



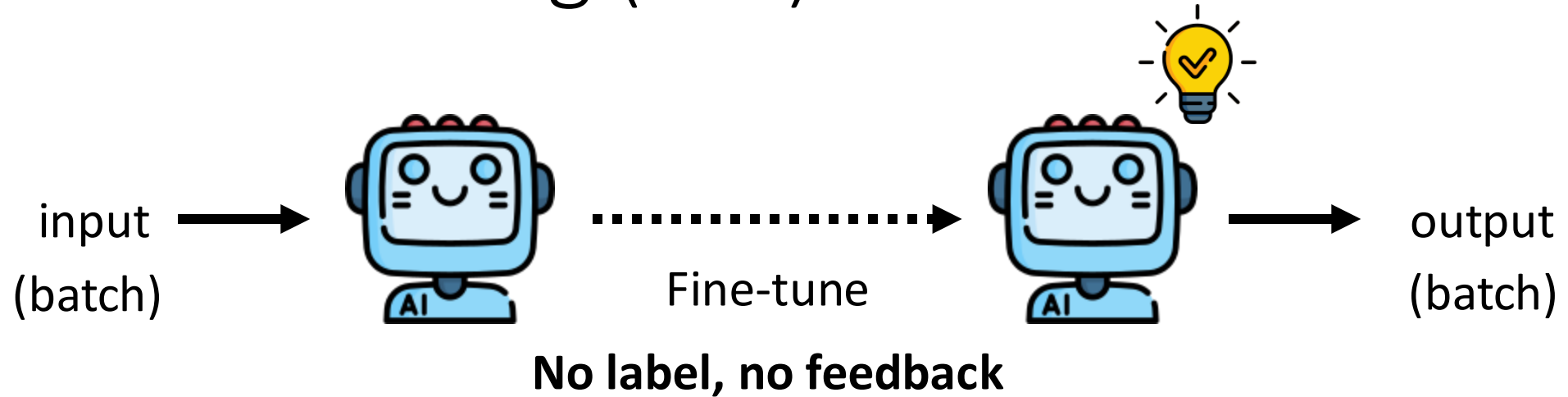
不同 Foundation Model 能不能 Merge?



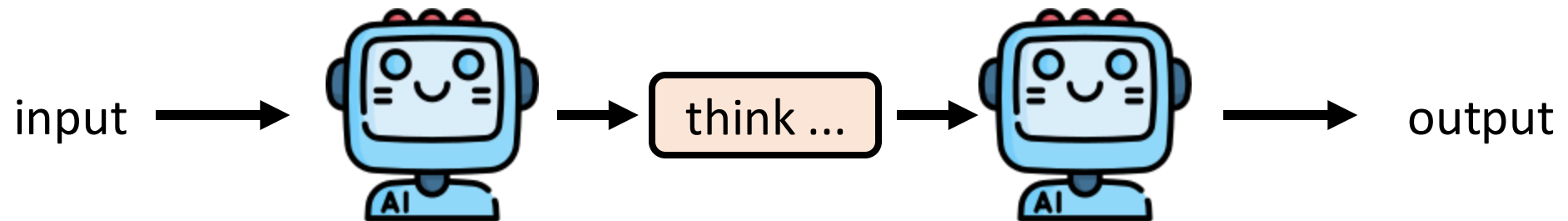
Test-Time Training

A hand holding a compass over a desert landscape with a road and a body of water. The hand is wearing a green sleeve. The compass is a standard analog compass with a white face and black markings. The background shows a paved road leading towards a body of water, with a sandy dune in the distance. The text "Test-Time Training" is overlaid in white, bold, sans-serif font.

Test-Time Training (TTT)



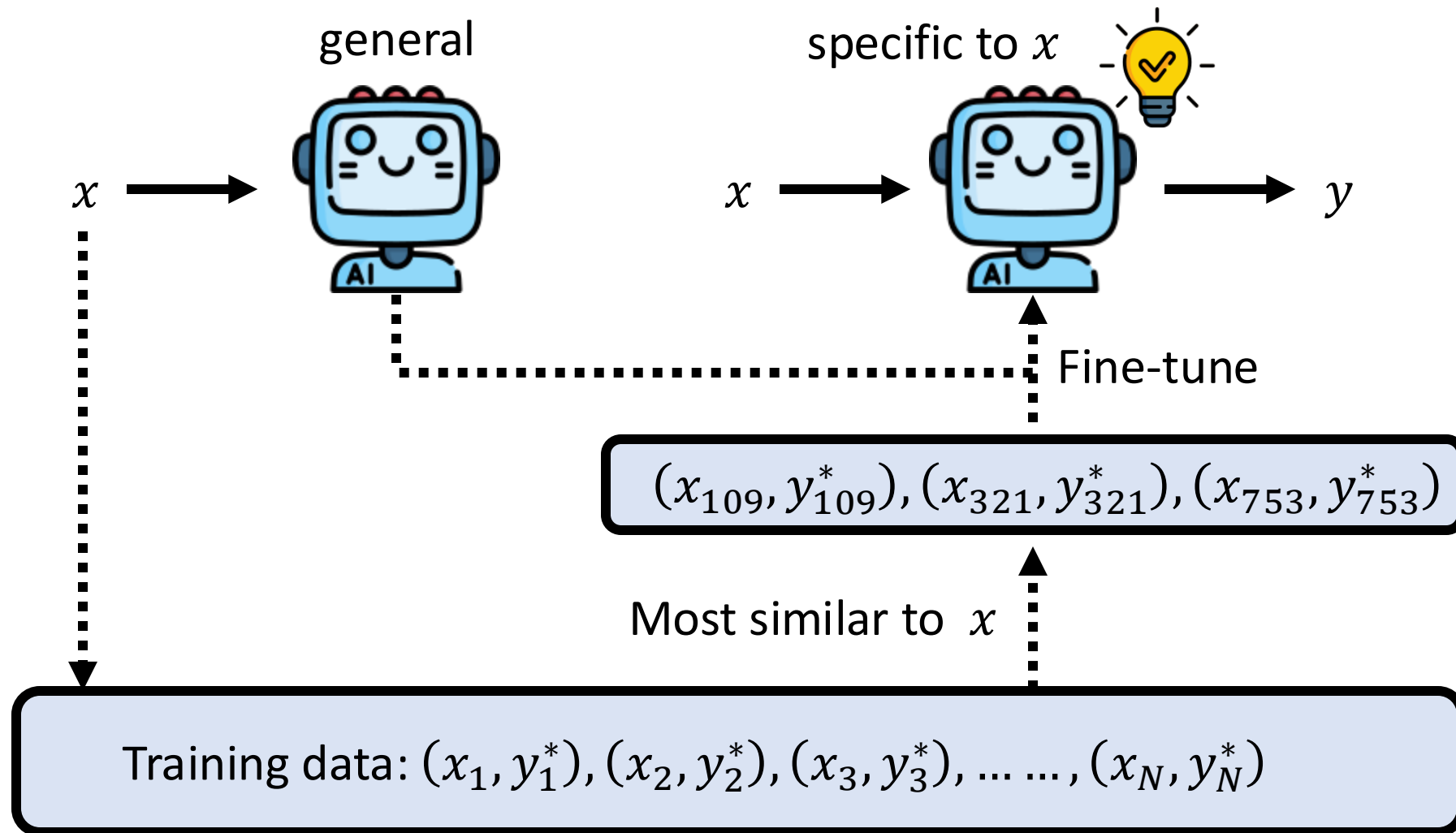
c.f. Reasoning (Test-time Computing)



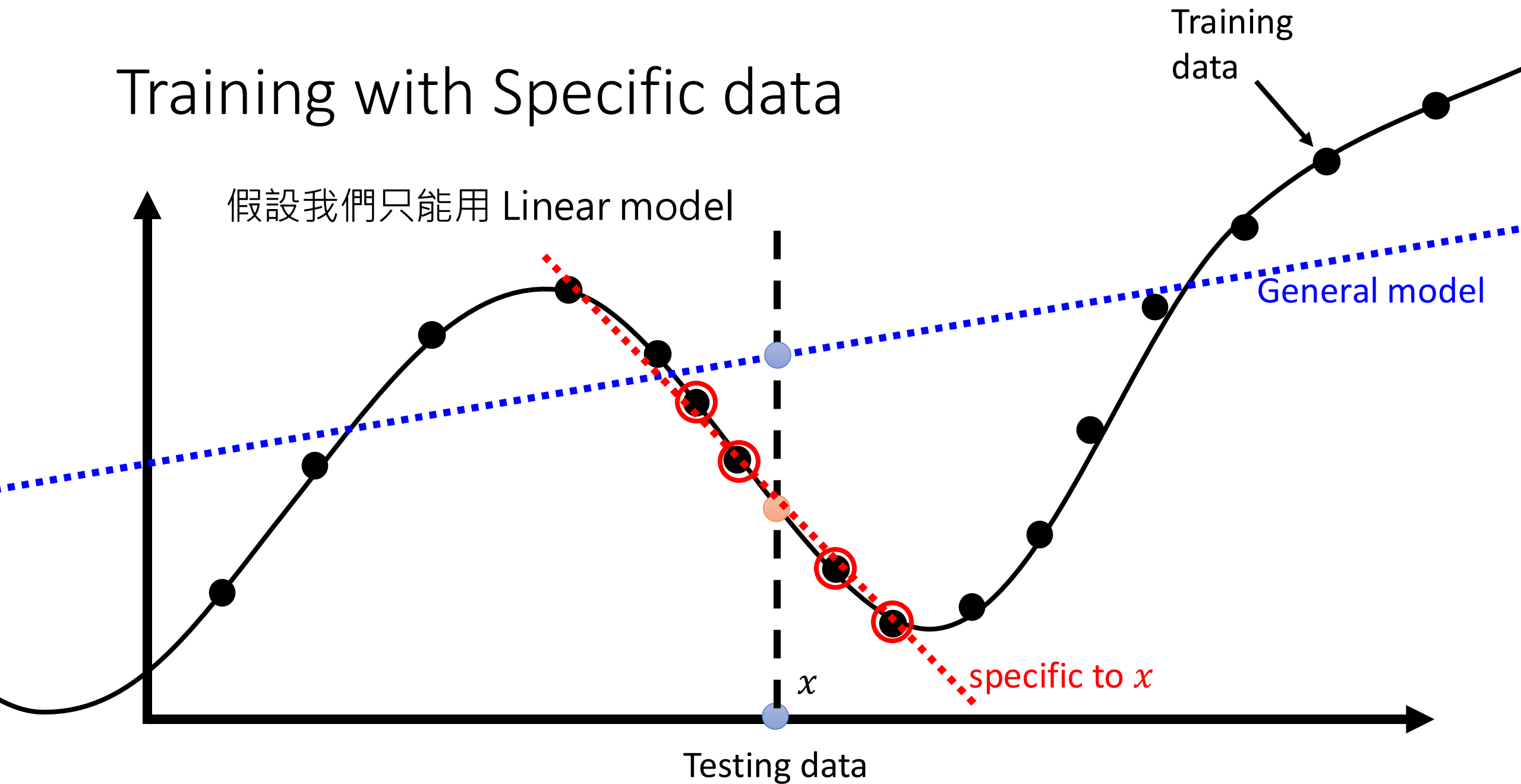
Ref: <https://youtu.be/bJFtcwLSNxl?si=9RfERMPqqd7w0bKu>

<https://arxiv.org/abs/2305.18466>
<https://arxiv.org/abs/2410.08020>
<https://arxiv.org/abs/2411.07279>

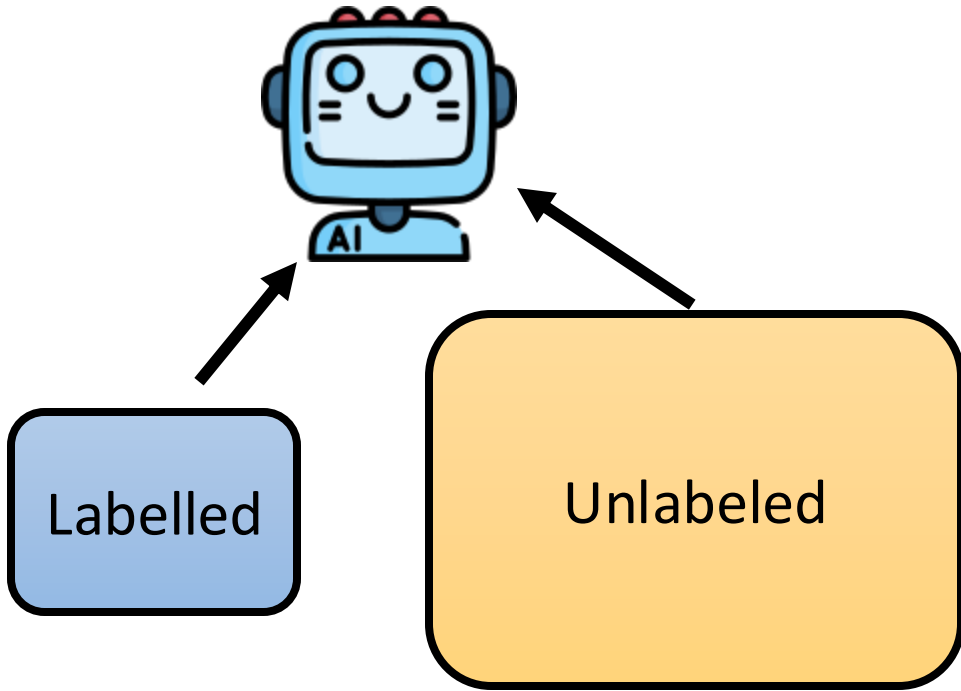
Training with Specific data



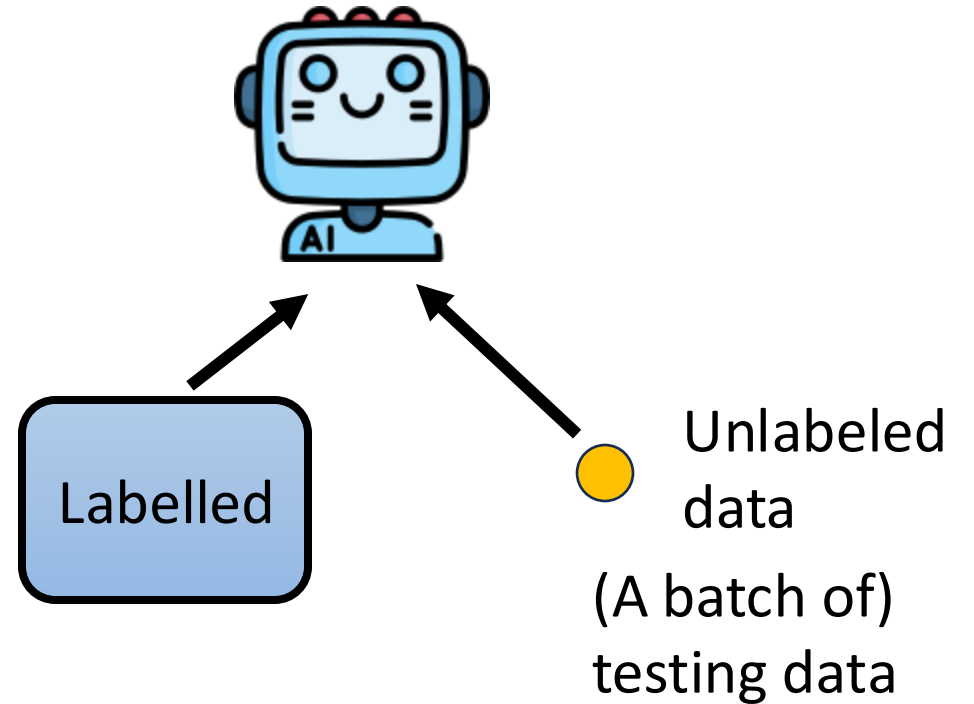
Training with Specific data



Semi-supervised Learning



Typical Semi-supervised Learning

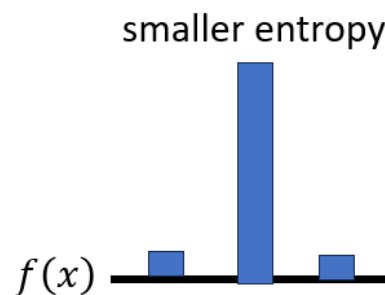
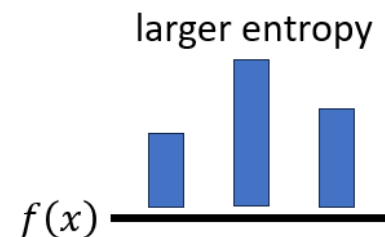


Test-Time Training

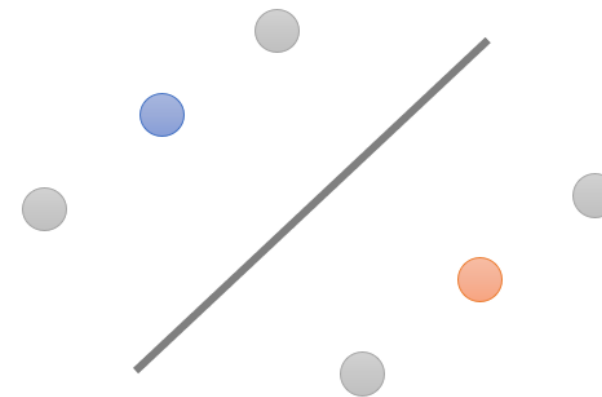
Semi-supervised Learning: Minimize Entropy

$$L = \sum_{x \sim \text{Train}} l(f(x), \hat{y}) + \lambda \sum_{x \sim \text{Unlabel}} l'(f(x))$$

↓ ↘
entropy distribution



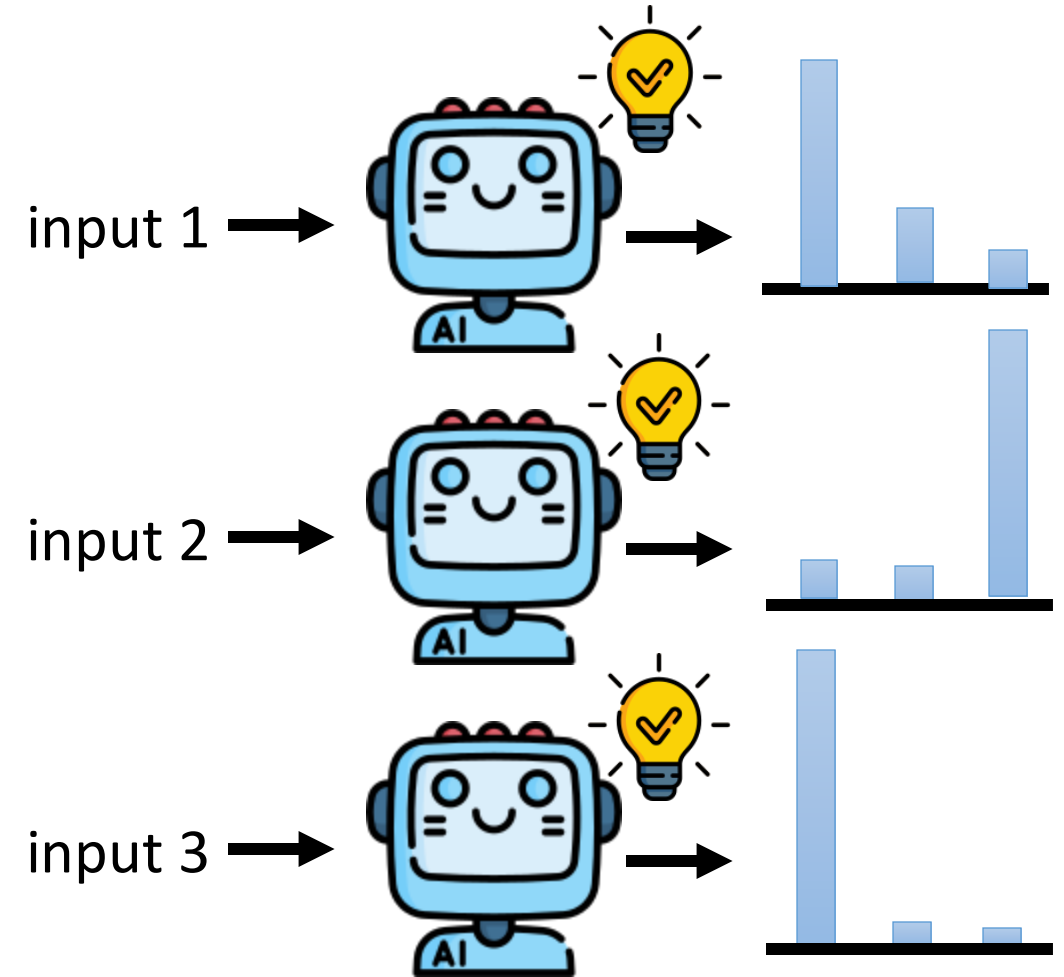
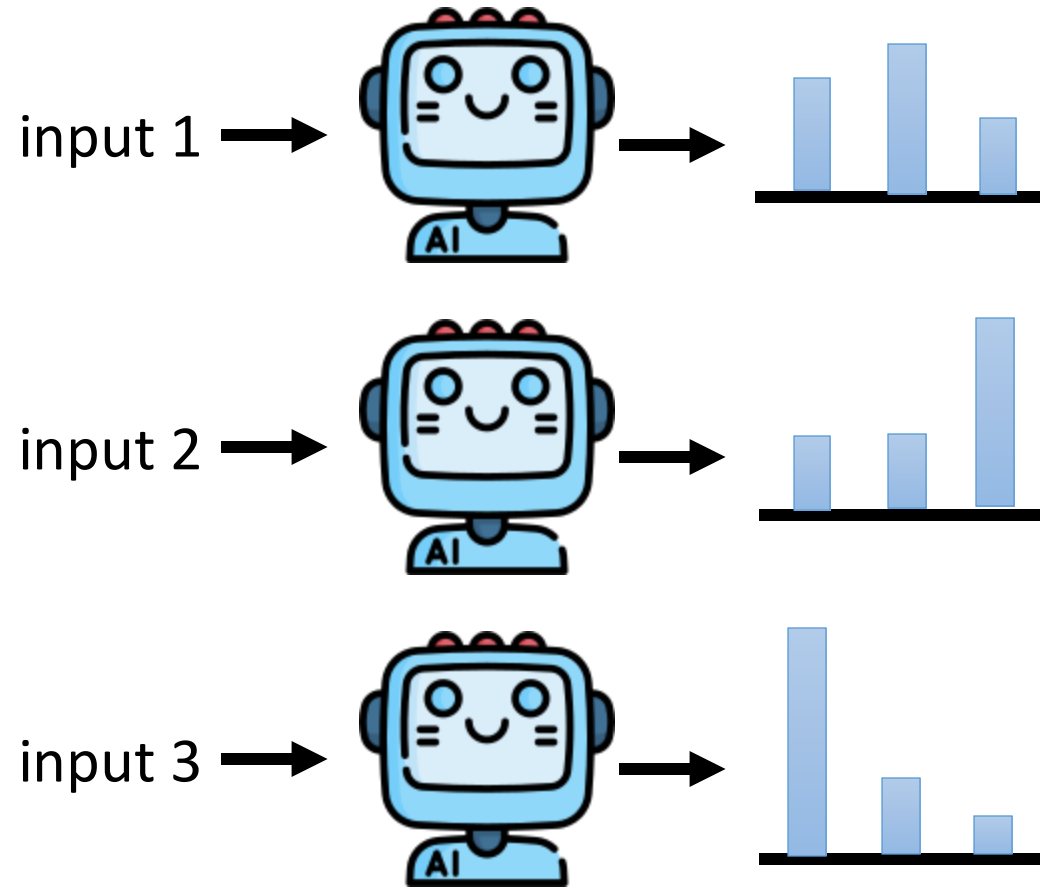
非黑即白
的世界



第六講

<https://youtu.be/mPWvAN4hzzY?si=DQj0TL2WCSzue5EE>

Semi-supervised Learning: Minimize Entropy



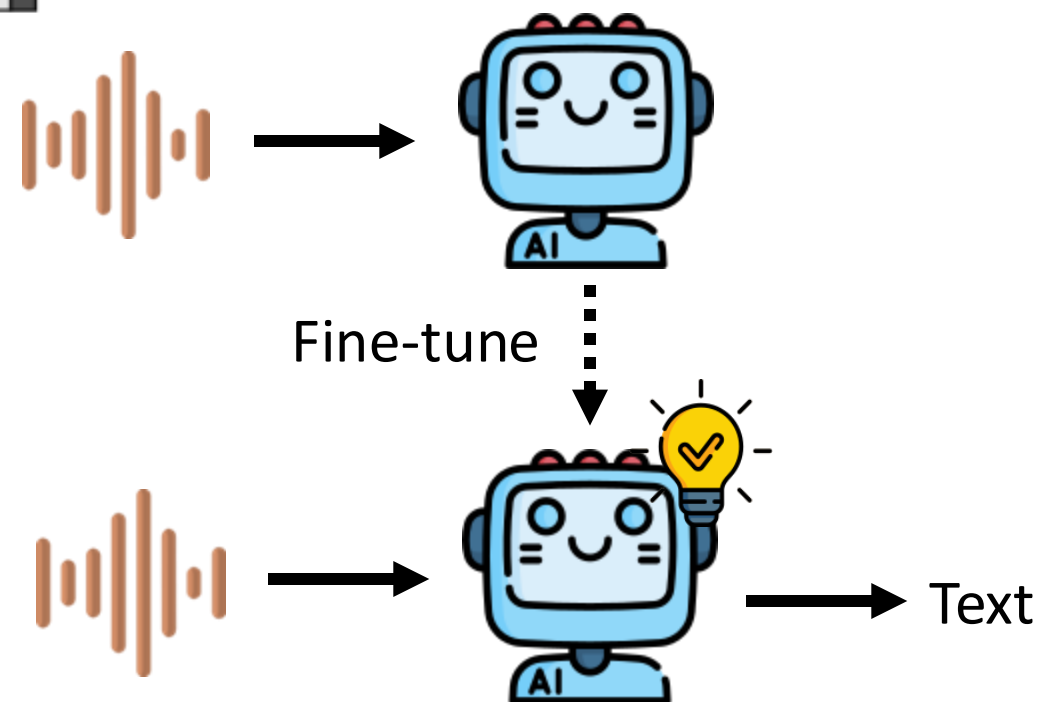
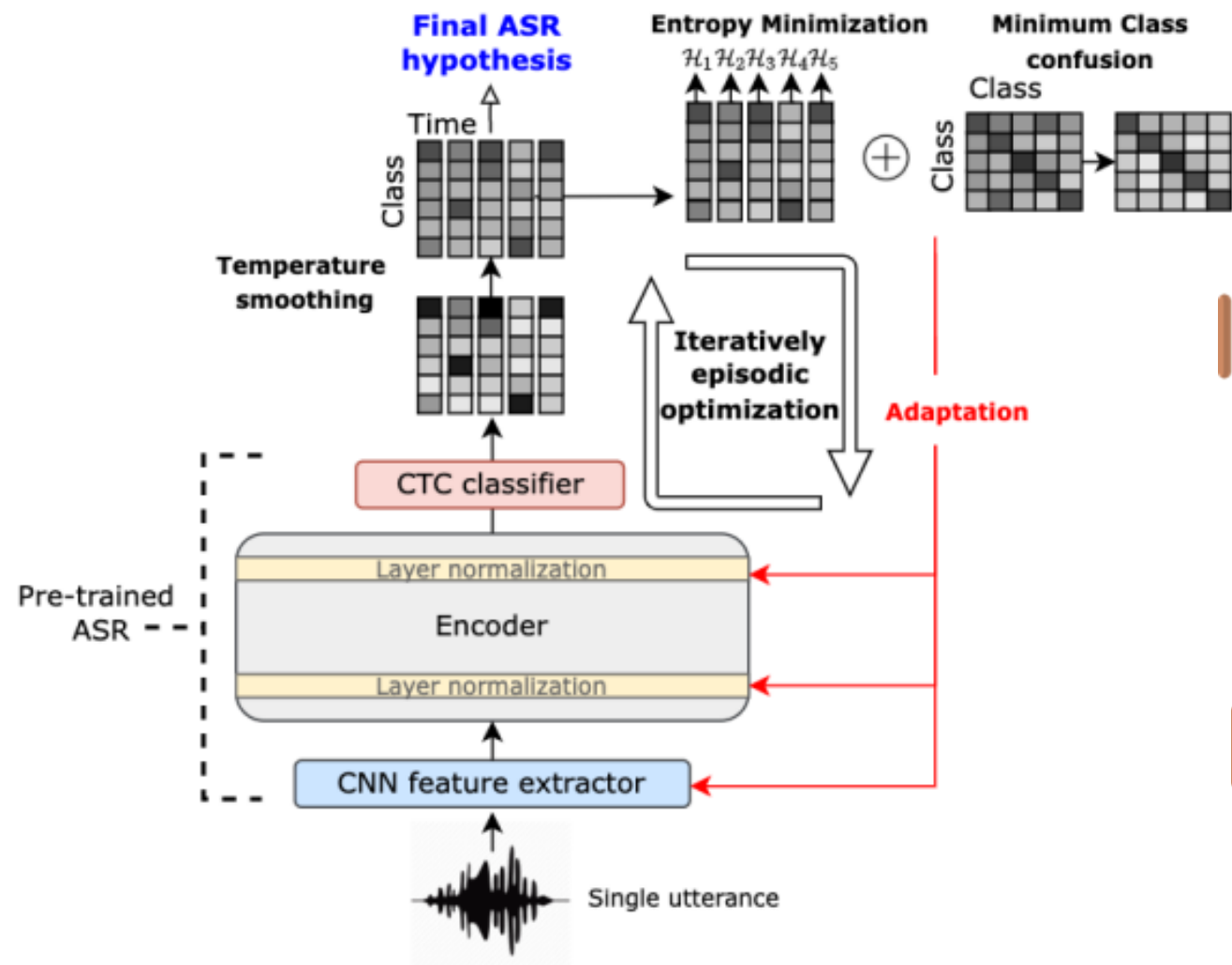
Test-Time Training for Speech Recognition

Single-Utterance Test-time Adaptation (SUTA)



Guan-Ting Lin

<https://arxiv.org/abs/2203.14222>

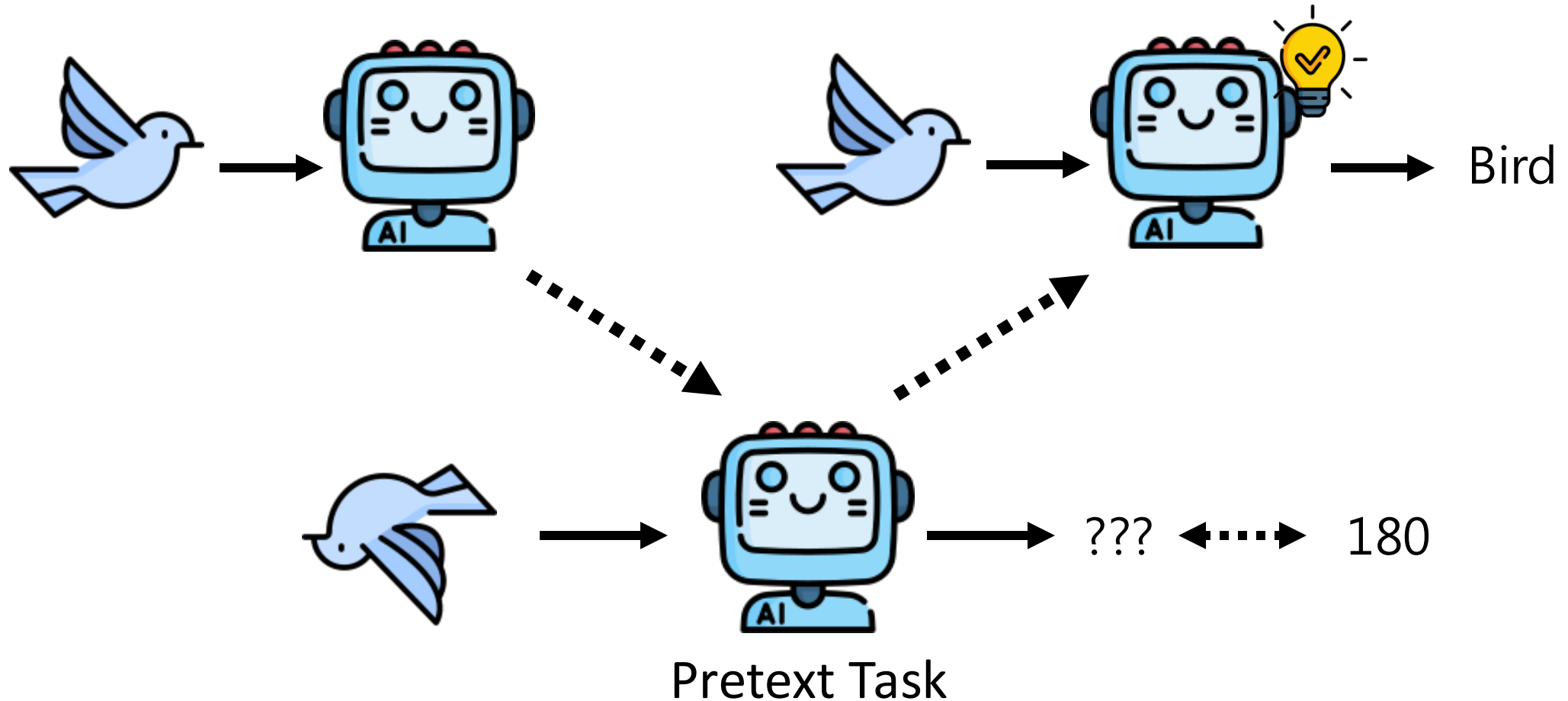


Performance of SUTA

Performance reference for source ASR model <i>wo/ adaptation</i>	Different domains					
	I.S test-o + δ			CH	CV	TD
	0	0.005	0.01			
SOTA (trained on target dataset)	2.5	-	-	5.8	15.4	5.6
RASR [26] (trained on LS)	6.8	-	-	-	29.9	13.0
TTA method						
(1) Our source ASR model [27] (trained on LS <i>wo/ adaptation</i>)	8.6	13.9	24.4	31.2	36.8	13.2
(1) + SDPL (Pseudo labeling)	8.3	13.1	23.1	30.4	36.3	12.8
(1) + SUTA	7.3	10.9	16.7	25.0	31.2	11.9

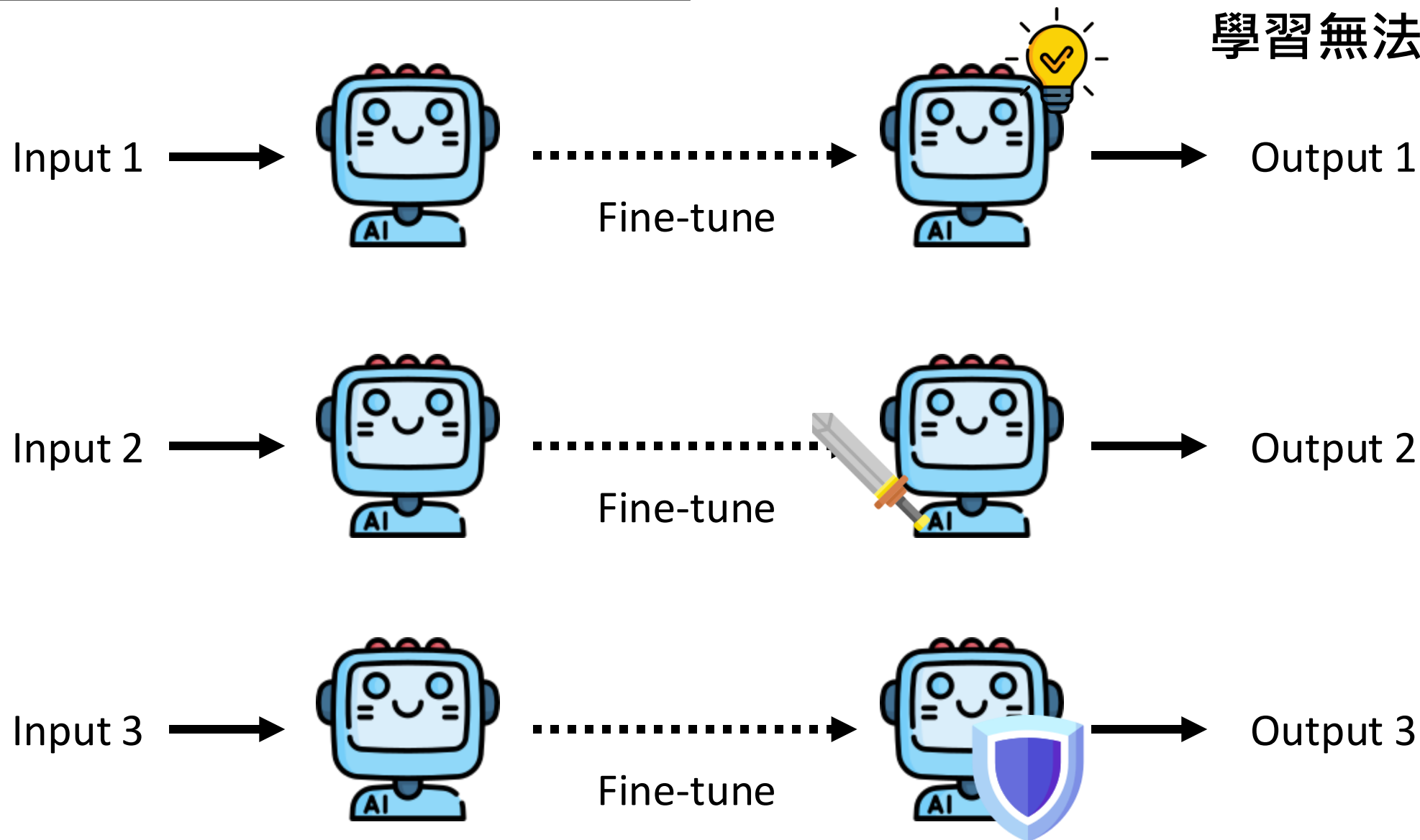
Source of table: <https://arxiv.org/abs/2203.14222>

Test-Time Training (TTT) using Pre-text Task



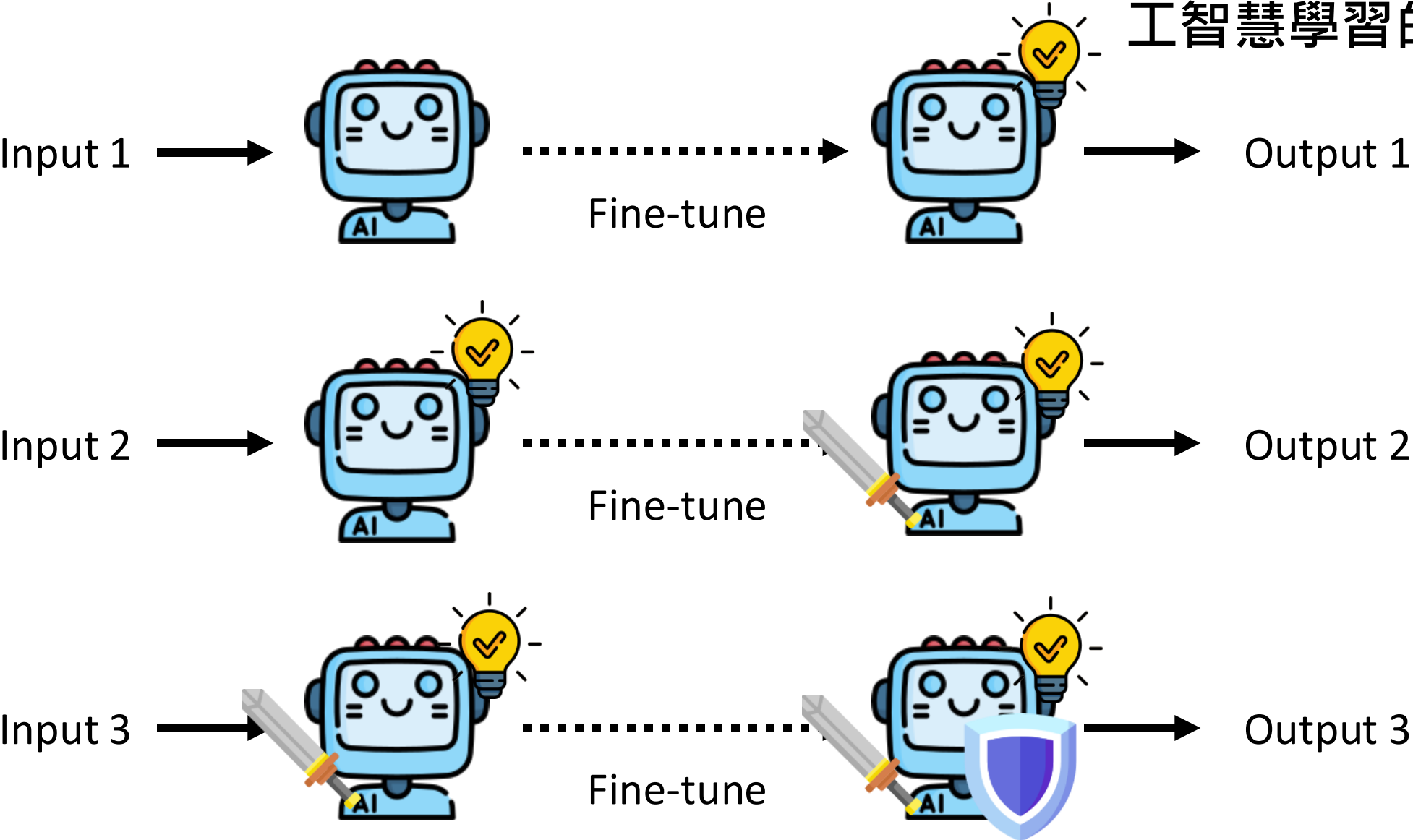
Standard Test-Time Training (TTT)

缺點：
學習無法累積



Continuous Test-Time Training (TTT)

接近一般人對於人工智慧學習的想像!

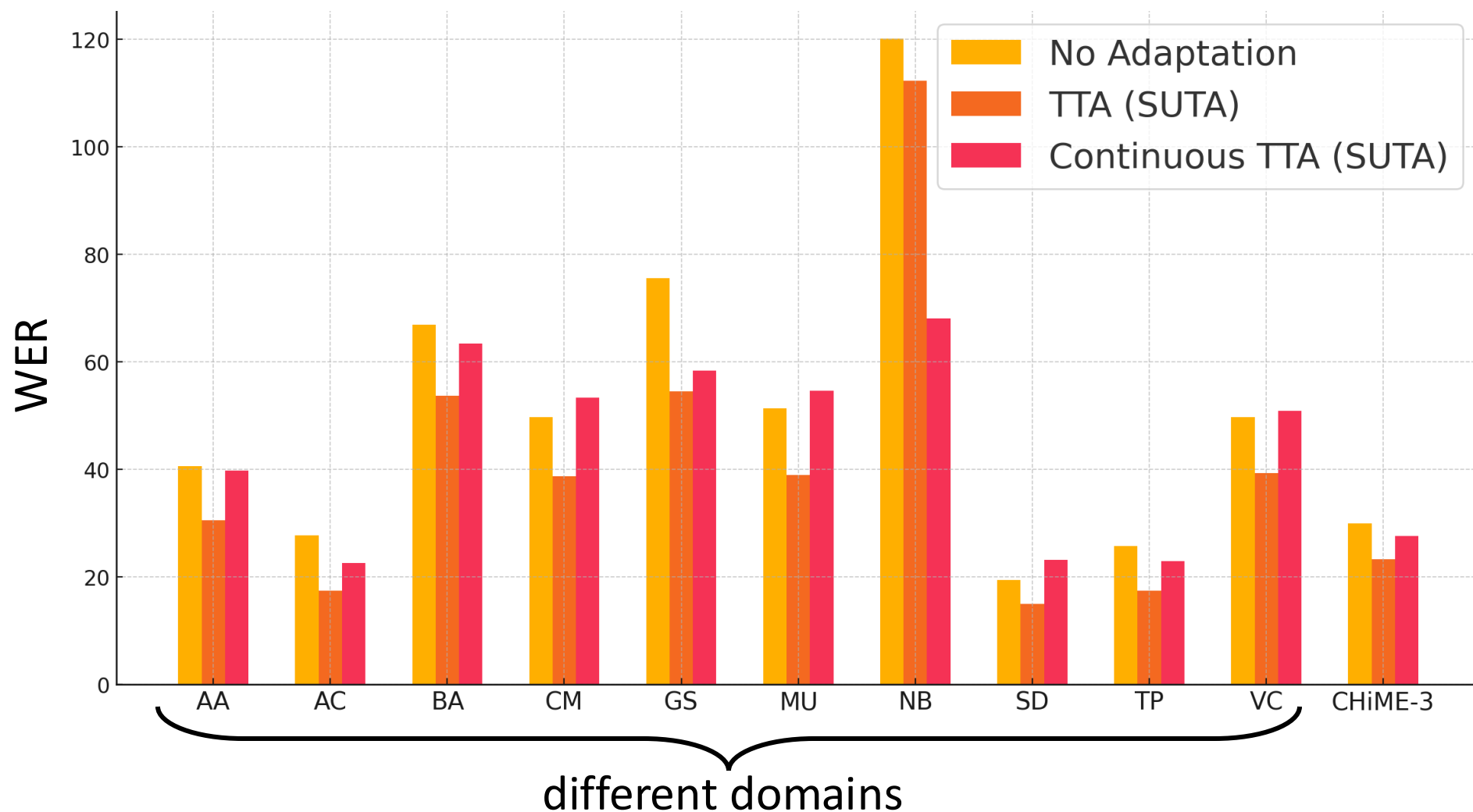


Continuous TTA

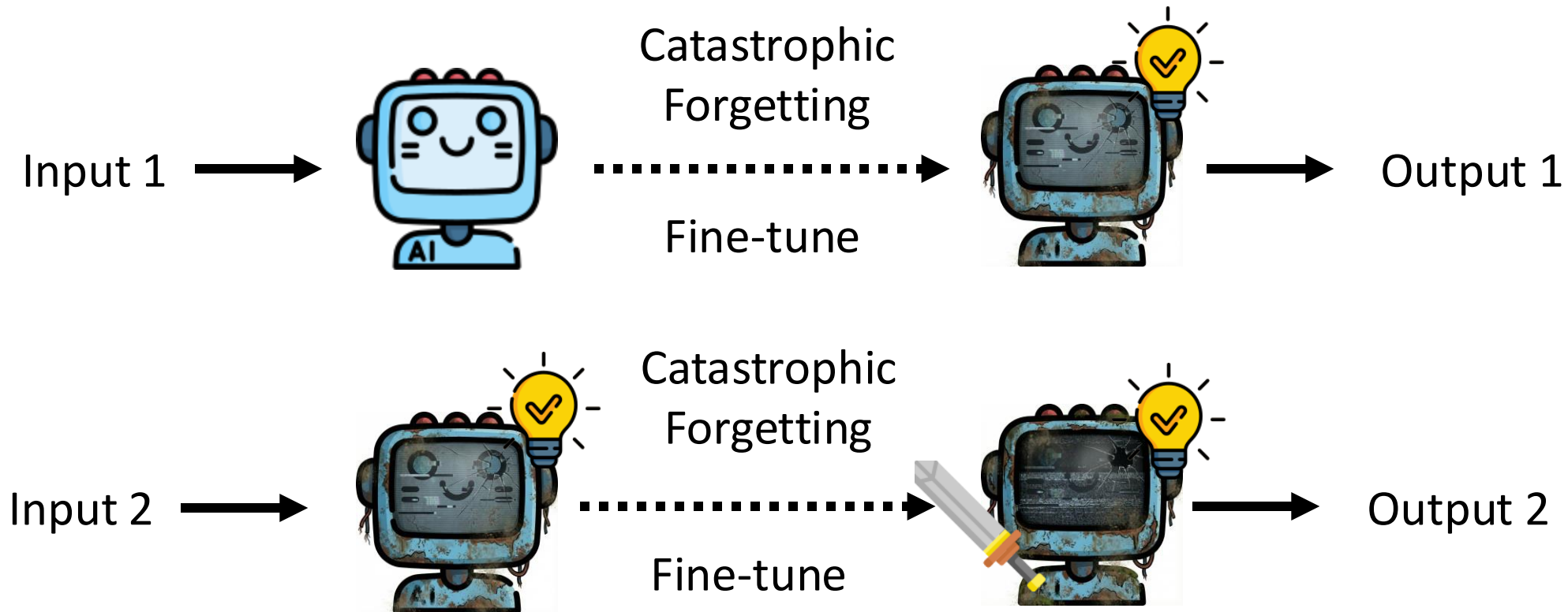
<https://arxiv.org/abs/2406.11064>

This situation is typical.

<https://arxiv.org/abs/2306.05401>



Continuous Test-Time Training (TTT)



RDumb: <https://arxiv.org/abs/2306.05401>

EATA: <https://arxiv.org/abs/2204.02610>

CoTTA: <https://arxiv.org/abs/2203.13591>

SAR: <https://arxiv.org/abs/2302.12400>

REALM: <https://arxiv.org/abs/2309.03964>

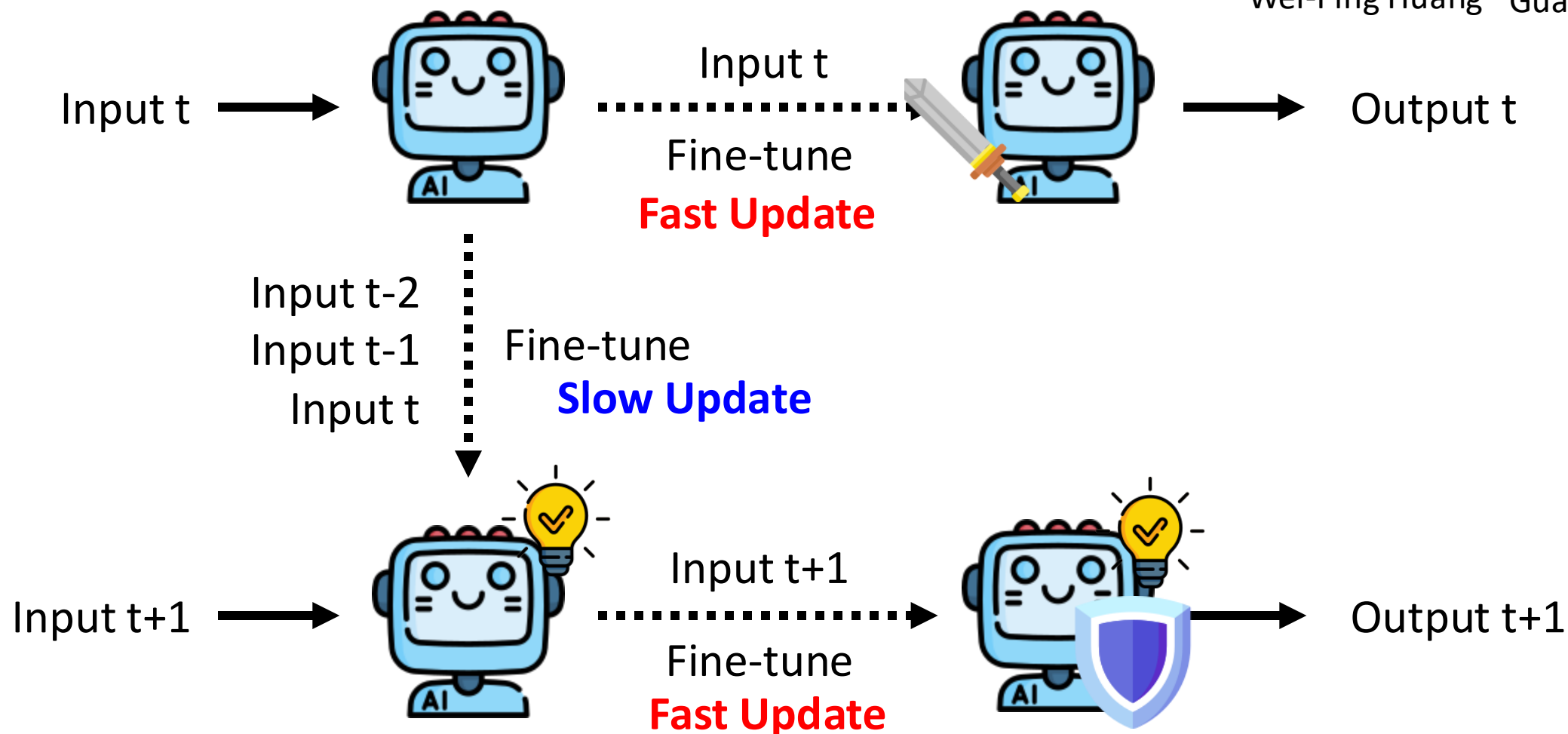
Continuous Test-Time Training (TTT) – Dynamic SUTA

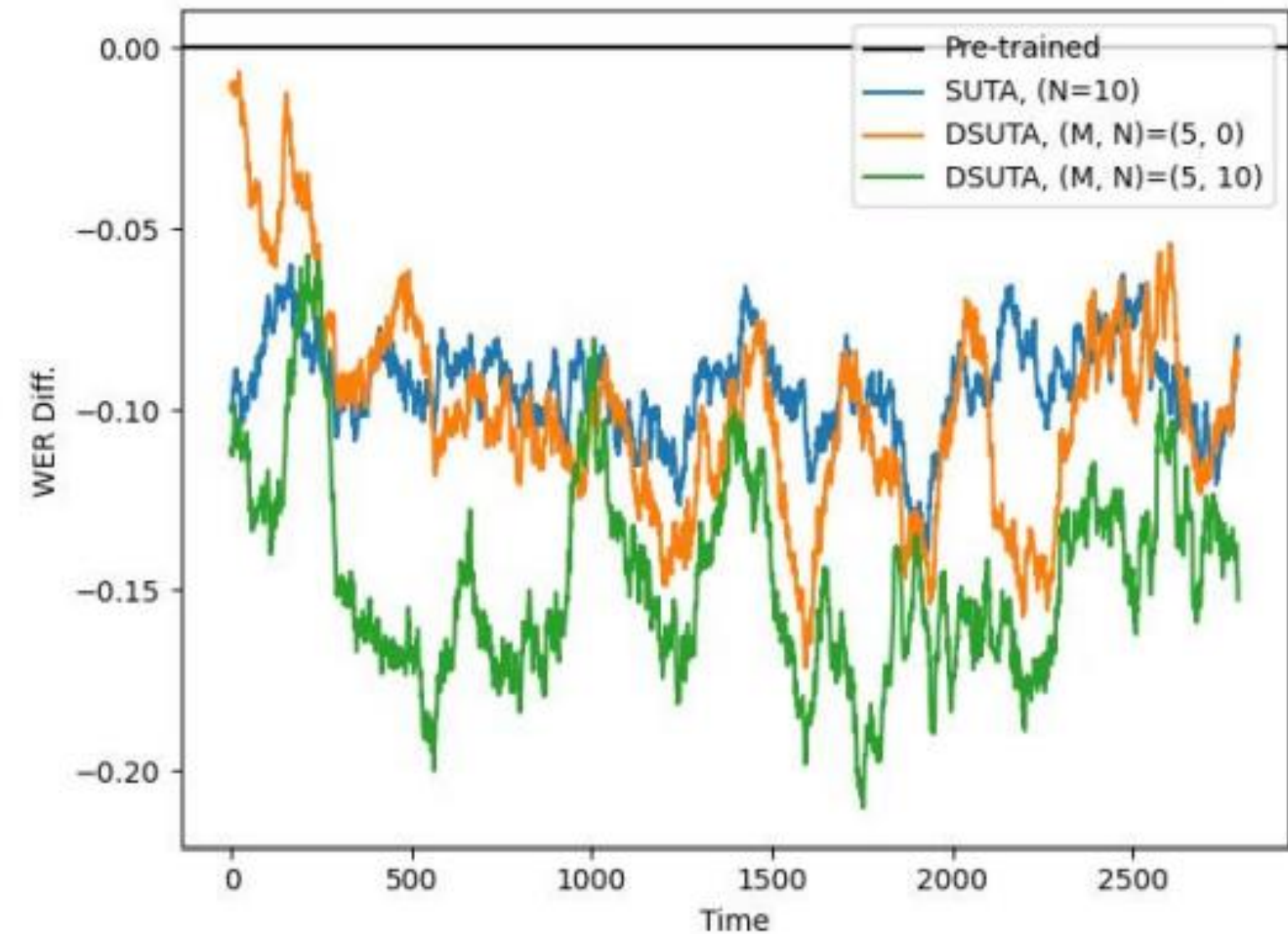


Wei-Ping Huang

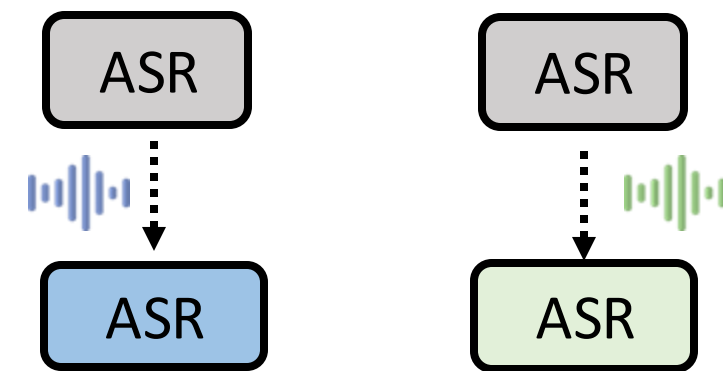


Guan-Ting Lin

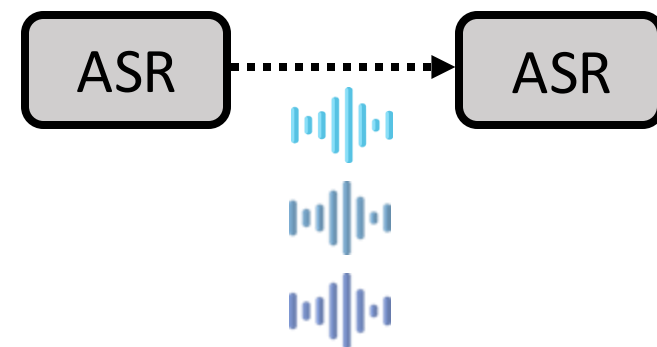




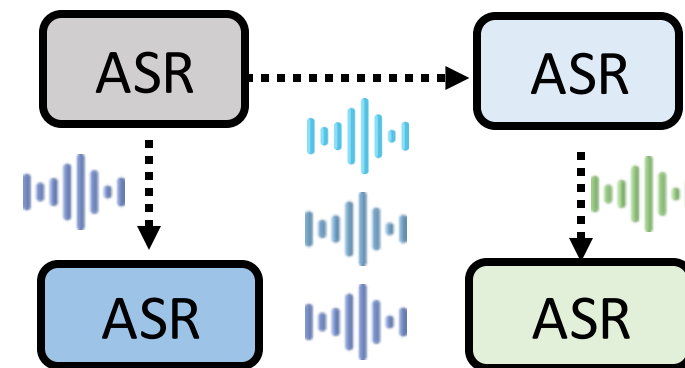
Blue
Curve



Orange
Curve

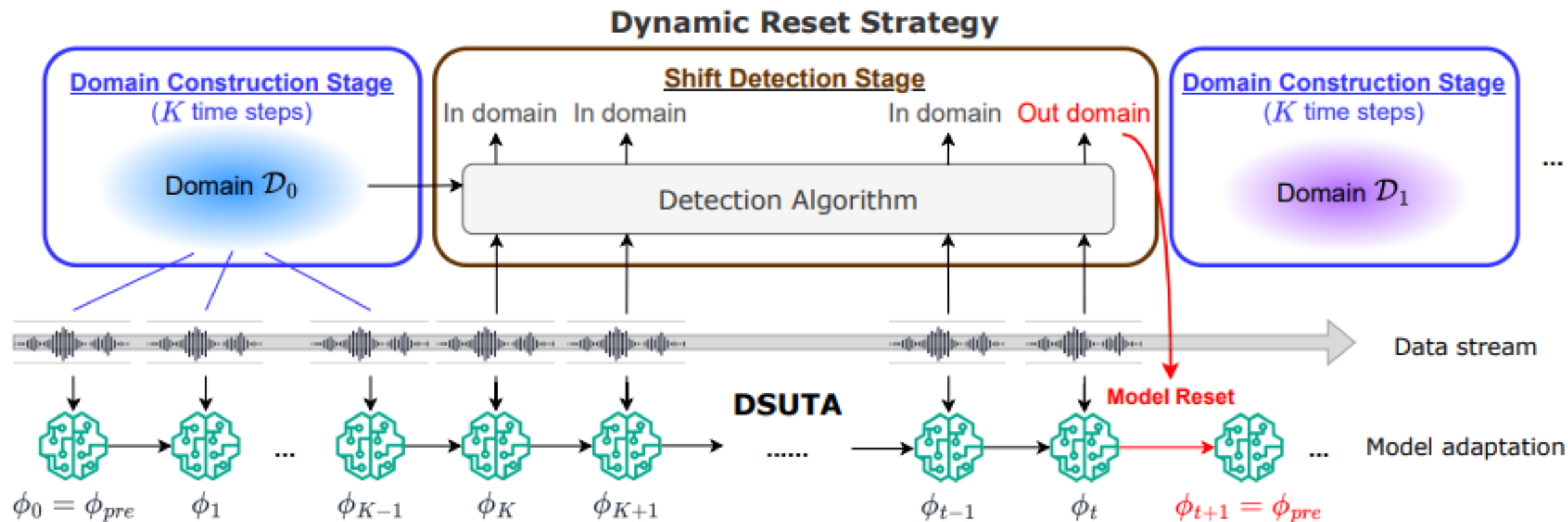


Green
Curve



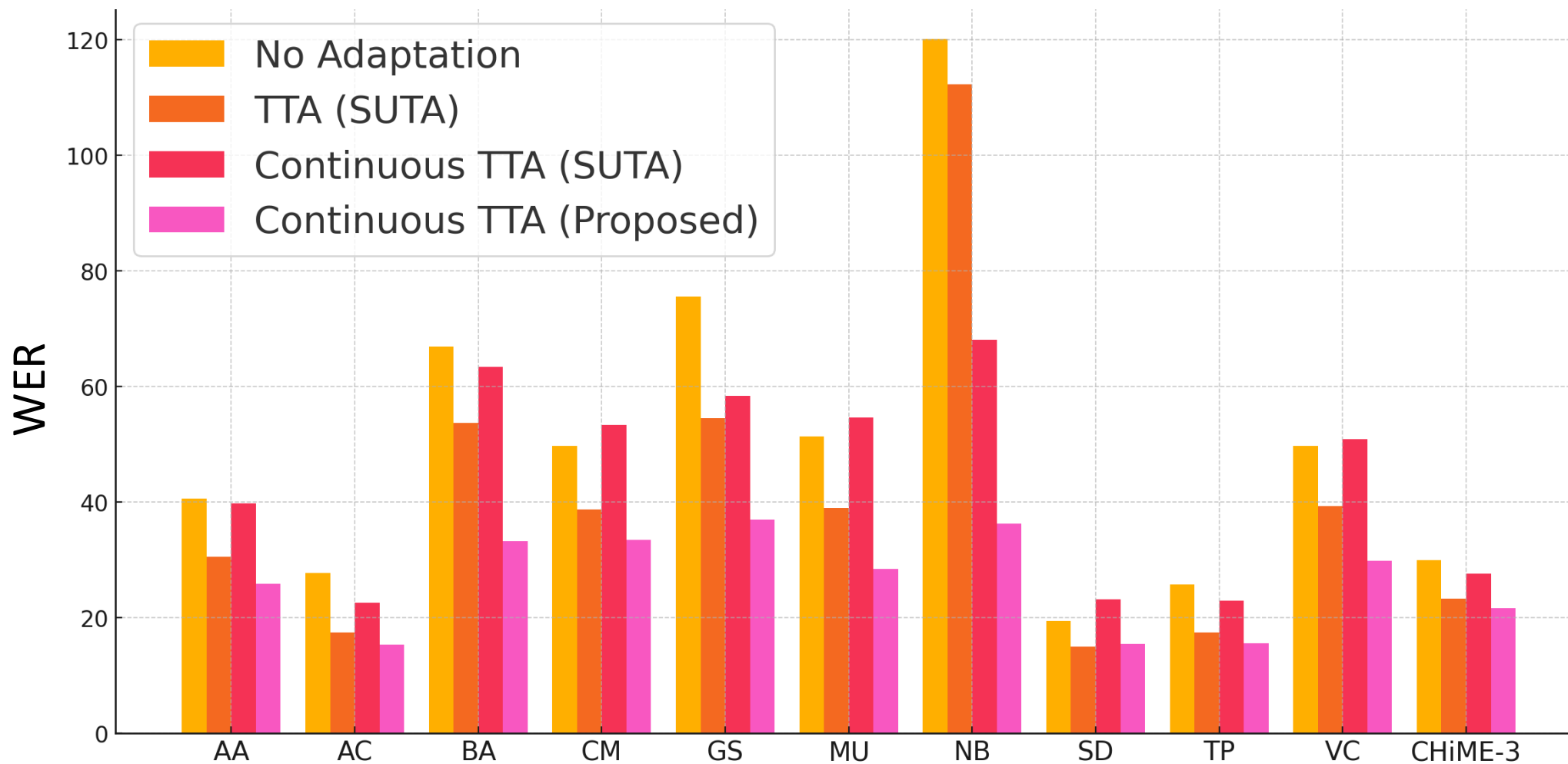
Continuous TTA – Dynamic SUTA

<https://arxiv.org/abs/2406.11064>



Continuous TTA – Dynamic SUTA

<https://arxiv.org/abs/2406.11064>



Concluding Remarks

Fine-tune with Gradient Descent

Model Editing

Model Merging

Test-Time Training (TTT)