

Visual Prompt Tuning



2026.09.04

Data Mining & Quality Analytics Lab.

김윤아

발표자 소개



❖ 김윤아 (Yoon Kim)

- 고려대학교 산업경영공학과 석사과정 (2026.09 ~ Present)
- Data Mining & Quality Analysis Lab. (김성범 교수님)

❖ Research Interest

- Vision-Language Models

❖ Contact

- 2022130818@korea.ac.kr

Introduction

Transformer

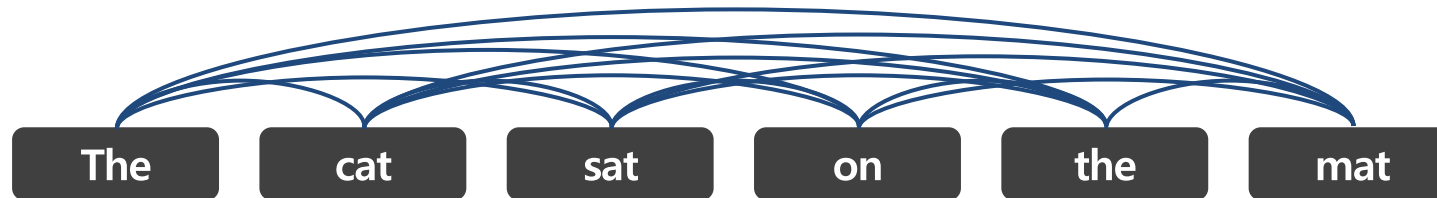
❖ RNN 구조란?

- 문장 내 단어들을 순차적으로 처리 → 먼 단어 관계 포착에 불리



❖ Transformer 구조란?

- Self-Attention을 통해 문장 내 모든 단어 간의 상호 관계를 한 번에 계산 → 거리와 무관하게 문맥 포착



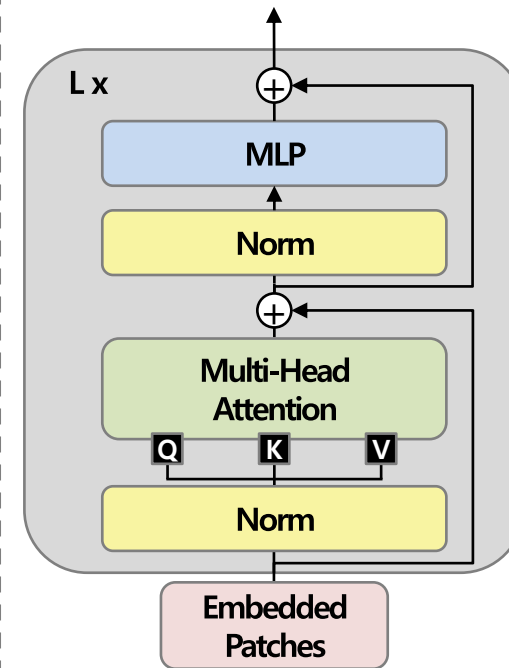
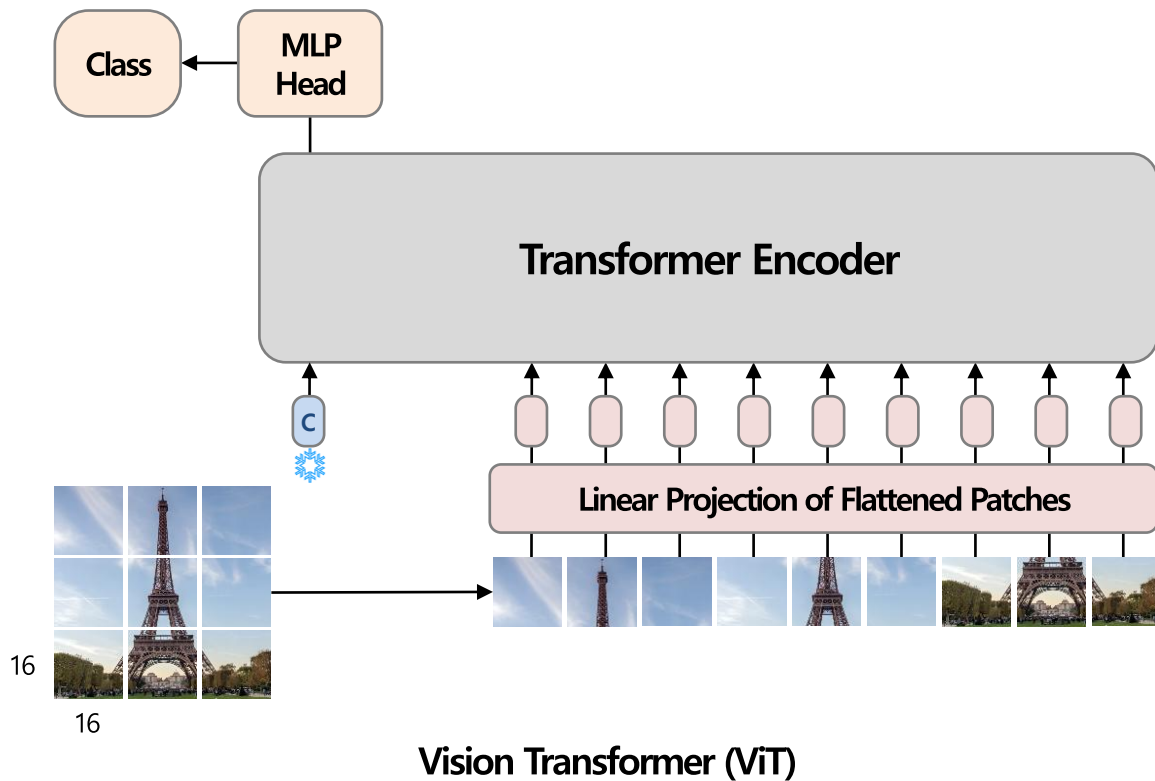
Attention is All You Need
(Vaswani et al., NeurIPS 2017)

Introduction

Vision Transformer (ViT)

❖ Vision Transformer (ViT) 구조란?

- 이미지를 16×16 크기의 patch로 나눈 뒤 Transformer Encoder를 통해 patch 간 관계를 학습하고, [CLS] token을 이용해 최종 이미지를 분류



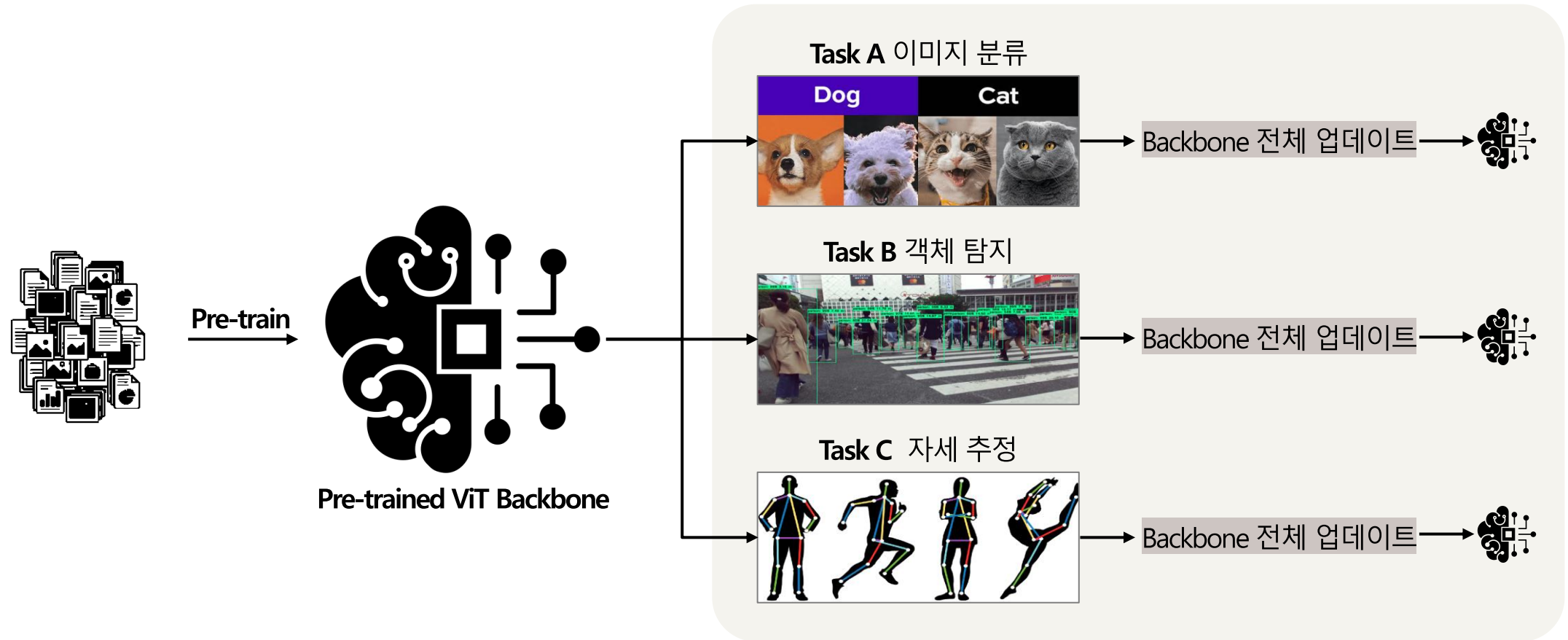
Transformer Encoder



An Image is Worth 16x16 Words:
Transformers for Image Recognition at Scale
(Dosovitskiy et al., ICLR 2021)

Introduction

Pretrain-then-Finetune Paradigm



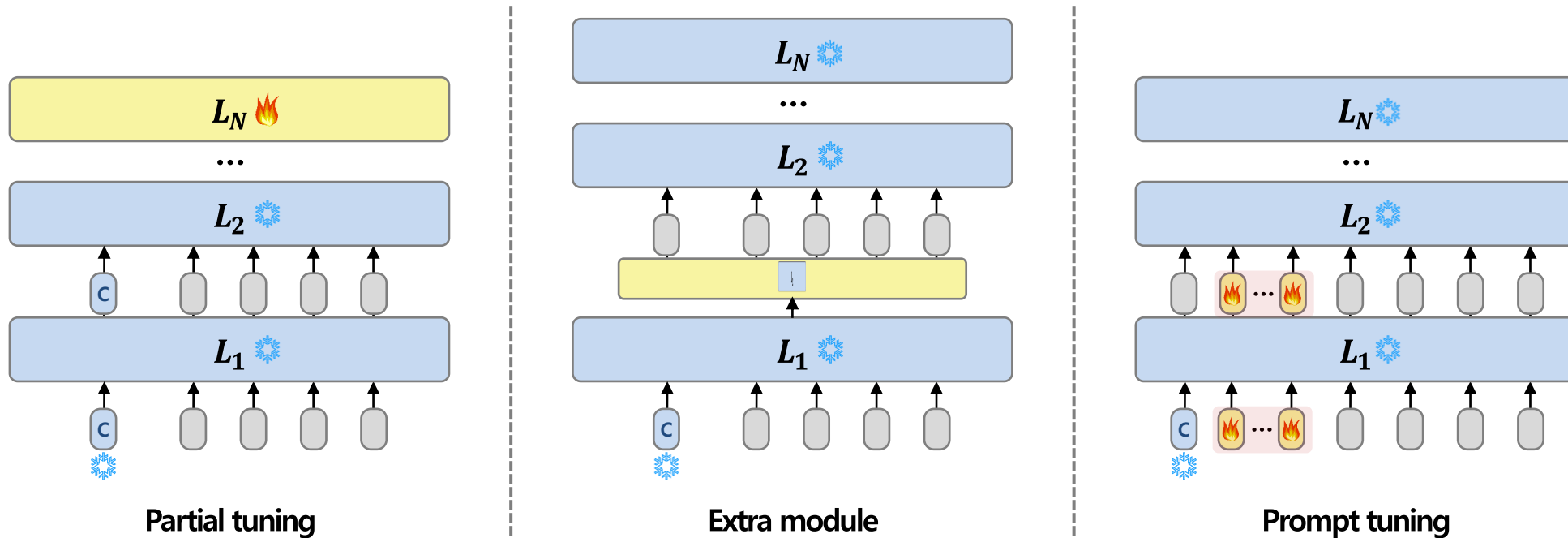
Full Fine-Tuning: Task마다 전체 backbone을 업데이트·저장 → 학습 및 저장 비용 증가
→ Pre-trained backbone은 공유하고 일부 parameter만 학습하는 **Parameter-Efficient Fine-Tuning (PEFT)** 필요

Introduction

Parameter-Efficient Fine-Tuning (PEFT)

❖ PEFT의 세 가지 접근 방식: 어디를 학습할 것인가?

- **Partial tuning:** 백본의 일부 파라미터만 업데이트 → 백본 가중치가 변형되므로 다중 태스크 환경에서 공유 · 재사용이 어려움
- **Extra module:** 백본은 동결, 작은 task-specific module만 학습 → 구조 의존성과 추가 연산량 발생
- **Prompt tuning:** 백본은 동결, 입력에 learnable prompt token 추가 → 구조 수정 없이 적은 파라미터만으로 효율적인 task adaptation 가능

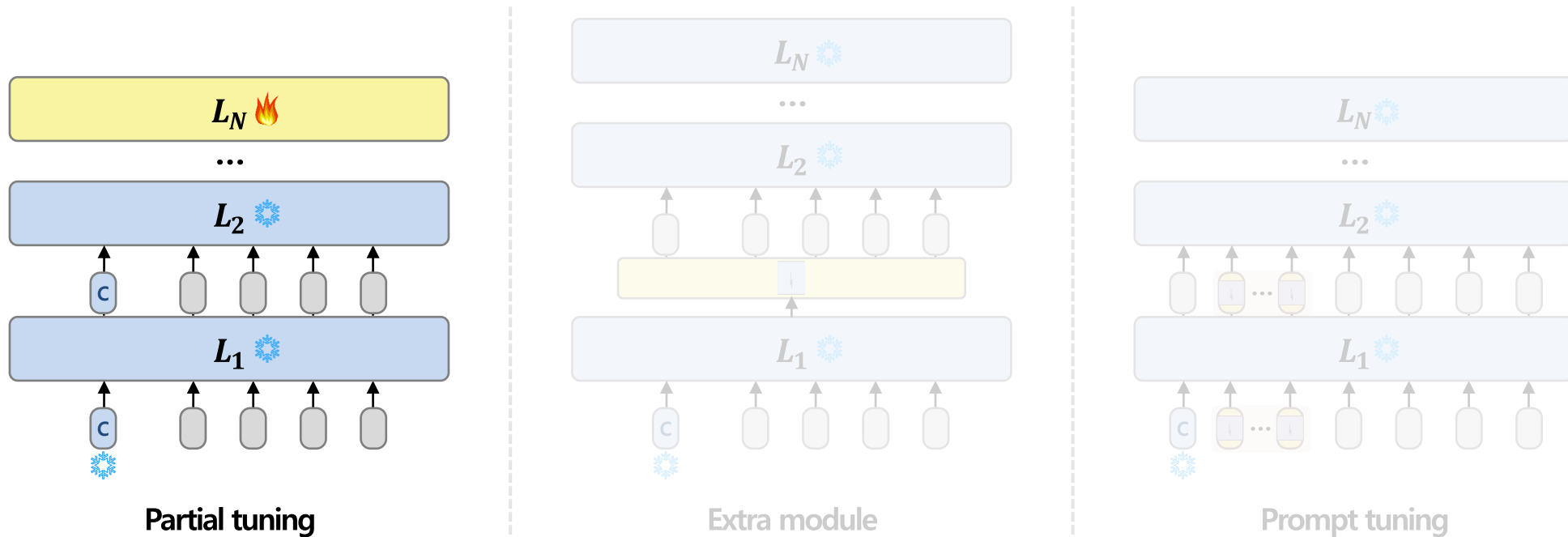


Introduction

Parameter-Efficient Fine-Tuning (PEFT)

❖ PEFT의 세 가지 접근 방식: 어디를 학습할 것인가?

- **Partial tuning:** 백본의 일부 파라미터만 업데이트 → 백본 가중치가 변형되므로 다중 태스크 환경에서 공유 · 재사용이 어려움
- **Extra module:** 백본은 동결, 작은 task-specific module만 학습 → 구조 의존성과 추가 연산량 발생
- **Prompt tuning:** 백본은 동결, 입력에 learnable prompt token 추가 → 구조 수정 없이 적은 파라미터만으로 효율적인 task adaptation 가능

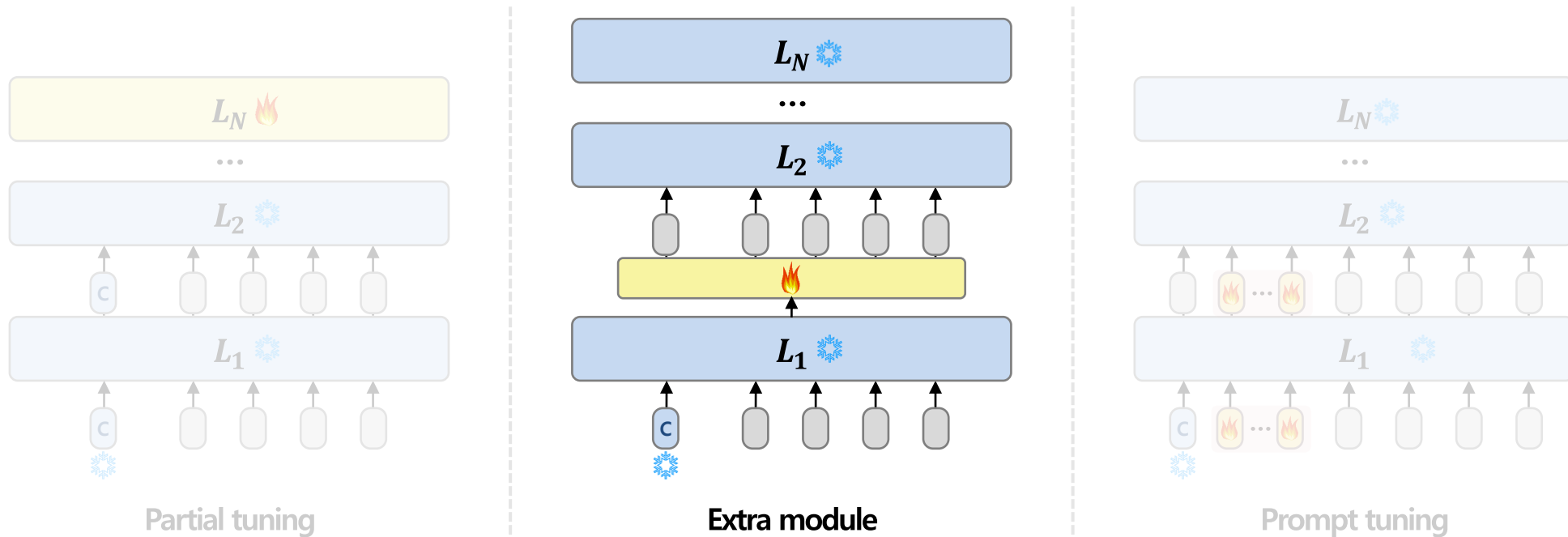


Introduction

Parameter-Efficient Fine-Tuning (PEFT)

❖ PEFT의 세 가지 접근 방식: 어디를 학습할 것인가?

- **Partial tuning:** 백본의 일부 파라미터만 업데이트 → 백본 가중치가 변형되므로 다중 태스크 환경에서 공유 · 재사용이 어려움
- **Extra module:** 백본은 동결, 작은 task-specific module만 학습 → 구조 의존성과 추가 연산량 발생
- **Prompt tuning:** 백본은 동결, 입력에 learnable prompt token 추가 → 구조 수정 없이 적은 파라미터만으로 효율적인 task adaptation 가능

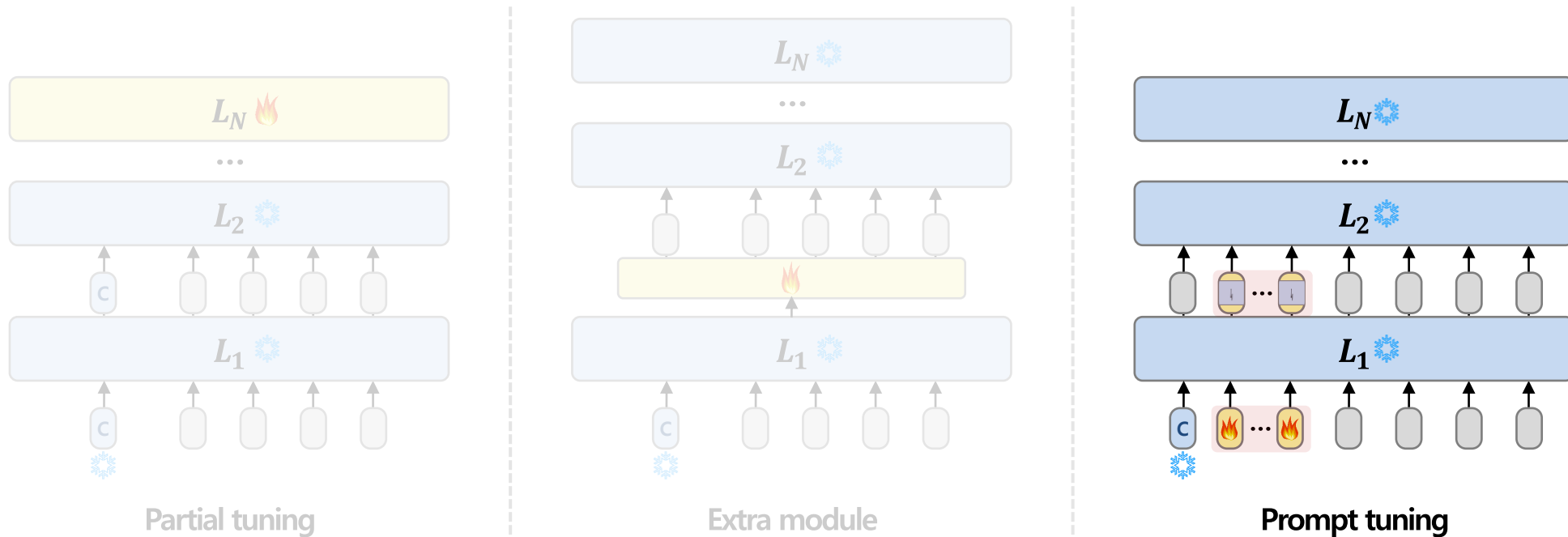


Introduction

Parameter-Efficient Fine-Tuning (PEFT)

❖ PEFT의 세 가지 접근 방식: 어디를 학습할 것인가?

- **Partial tuning:** 백본의 일부 파라미터만 업데이트 → 백본 가중치가 변형되므로 다중 태스크 환경에서 공유 · 재사용이 어려움
- **Extra module:** 백본은 동결, 작은 task-specific module만 학습 → 구조 의존성과 추가 연산량 발생
- **Prompt tuning:** 백본은 동결, 입력에 learnable prompt token 추가 → 구조 수정 없이 적은 파라미터만으로 효율적인 task adaptation 가능

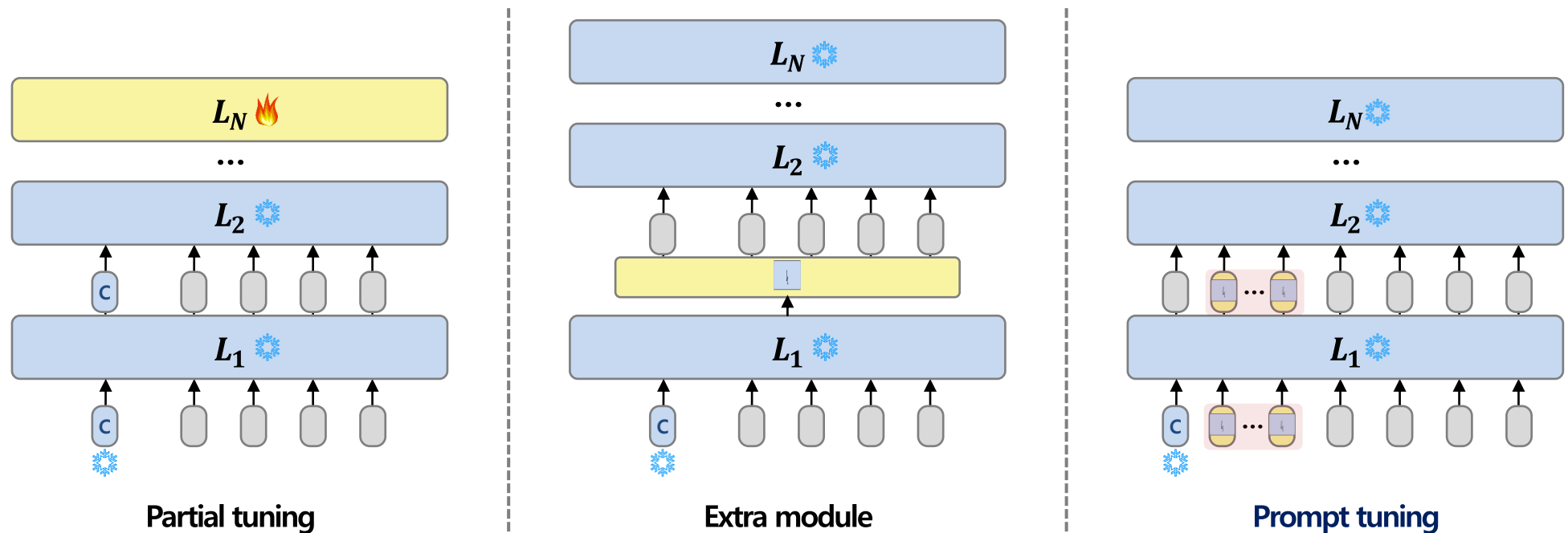


Introduction

Parameter-Efficient Fine-Tuning (PEFT)

❖ PEFT의 세 가지 접근 방식: 어디를 학습할 것인가?

- **Partial tuning:** 백본의 일부 파라미터만 업데이트 → 백본 가중치가 변형되므로 다중 태스크 환경에서 공유 · 재사용이 어려움
- **Extra module:** 백본은 동결, 작은 task-specific module만 학습 → 구조 의존성과 추가 연산량 발생
- **Prompt tuning:** 백본은 동결, 입력에 learnable prompt token 추가 → 구조 수정 없이 적은 파라미터만으로 효율적인 task adaptation 가능



Partial tuning

Extra module

Prompt tuning

모델 내부의 일부 가중치나 구조 수정

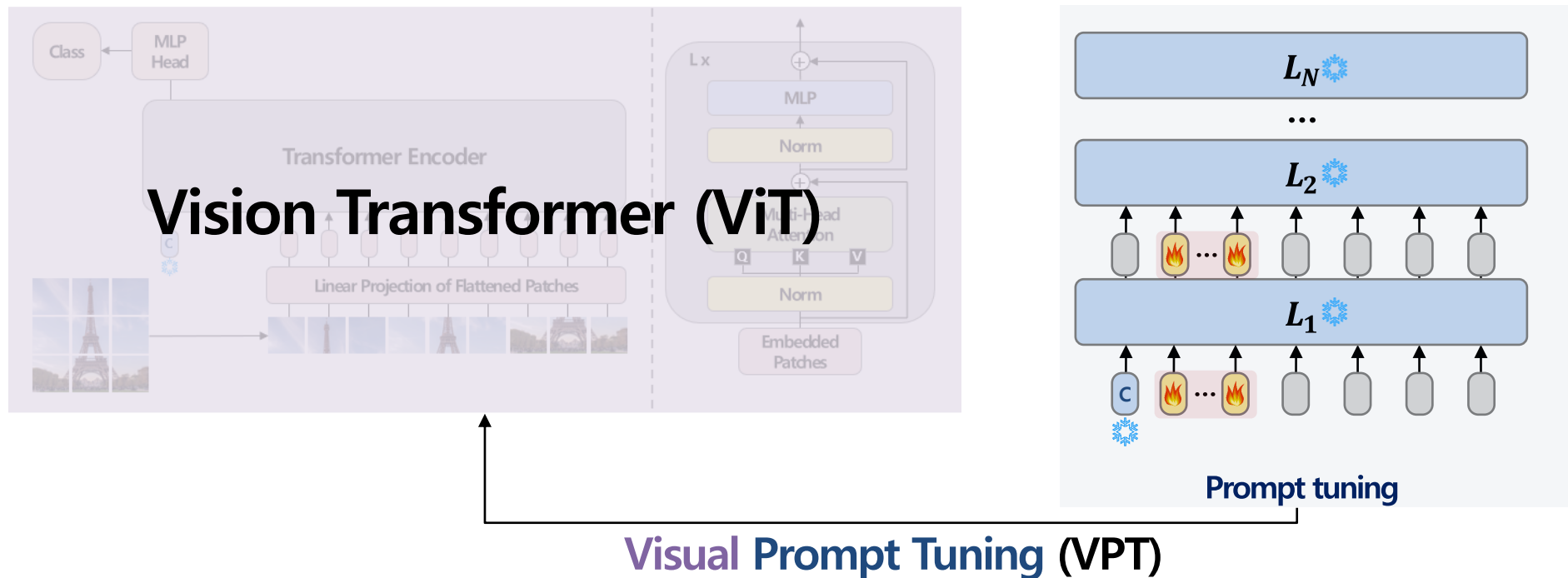
입력단의 토큰들을 통해 고정된 모델 백본 조정

Introduction

Parameter-Efficient Fine-Tuning (PEFT)

❖ PEFT의 세 가지 접근 방식: 어디를 학습할 것인가?

- **Partial tuning:** 백본의 일부 파라미터만 업데이트 → 백본 가중치가 변형되므로 다중 태스크 환경에서 공유·재사용이 어려움
- **Extra module:** 백본은 동결, 작은 task-specific module만 학습 → 구조 의존성과 추가 연산량 발생
- **Prompt tuning:** 백본은 동결, 입력에 learnable prompt token 추가 → 구조 수정 없이 적은 파라미터만으로 효율적인 task adaptation 가능

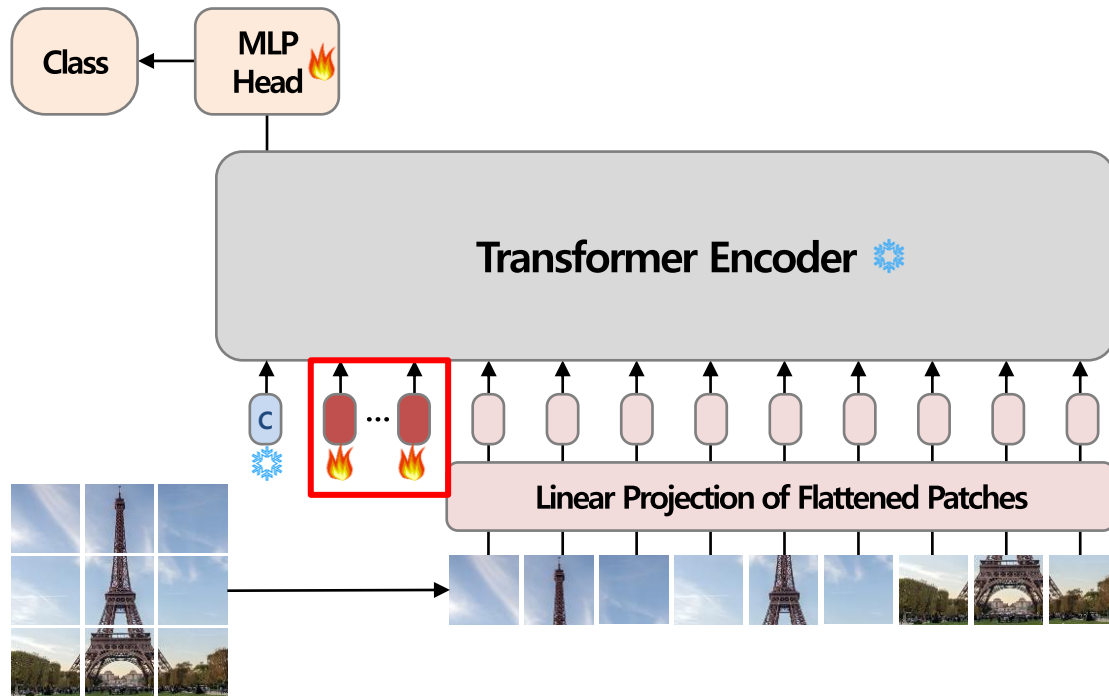


Visual Prompt Tuning (VPT)

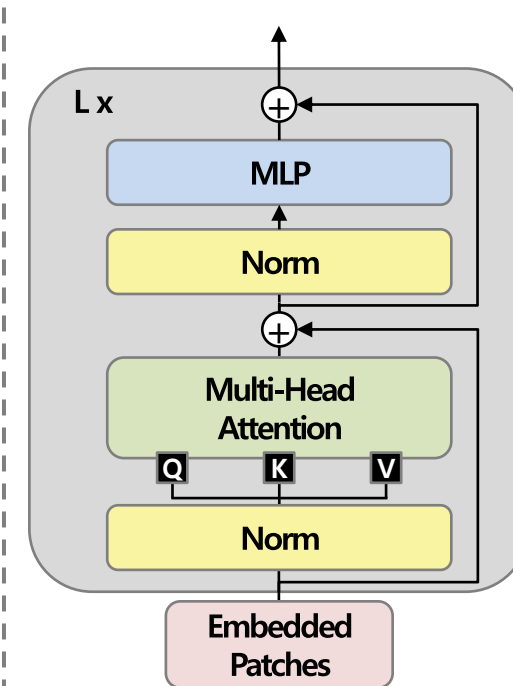
How VPT Works

❖ Visual Prompt Tuning (VPT)의 구조와 학습 방식은 어떻게 될까?

- [CLS] 토큰과 패치 임베딩 사이에 **prompt token**을 삽입 → 이미지에서 추출된 정보가 아닌, 태스크에 맞춰 최적화되는 유도용 파라미터
- 백본은 동결, prompt token과 classification head만 학습



Vision Transformer (ViT)



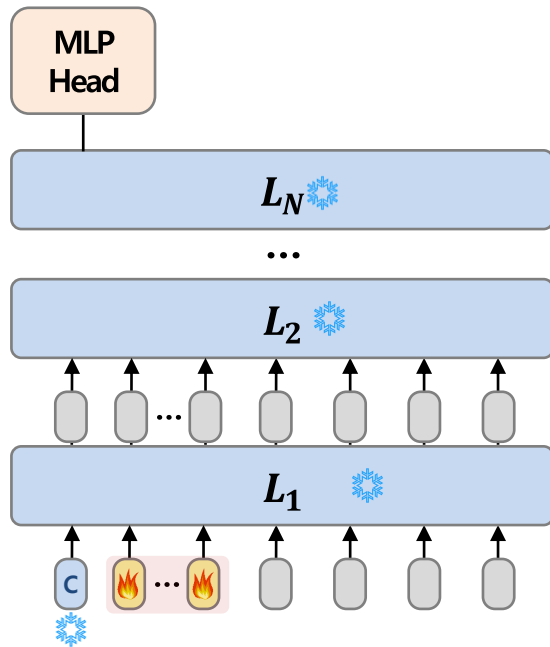
Transformer Encoder

Visual Prompt Tuning (VPT)

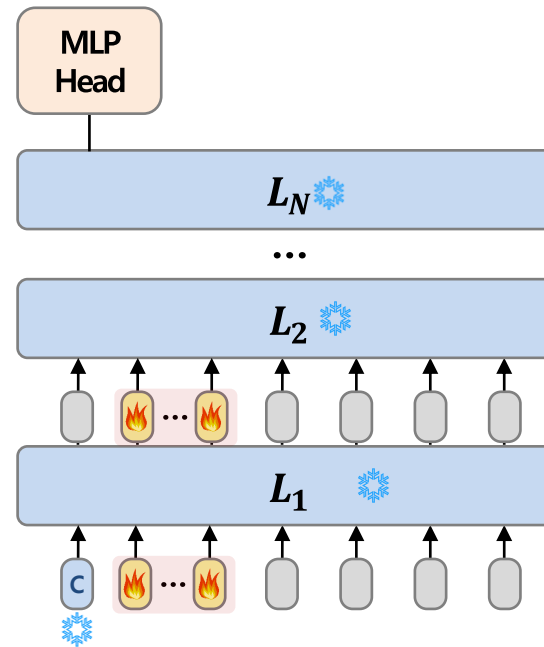
How VPT Works

❖ VPT의 두 가지 Prompt 삽입 방식은 무엇일까?

- **VPT-Shallow:** 첫 번째 layer 입력단에만 prompt 삽입 → 단일 prompt set만 학습, 매우 가벼움
- **VPT-Deep:** 모든 layer 입력마다 서로 다른 prompt 삽입 → layer별 task-specific 정보 반영, 표현력 향상



Visual Prompt Tuning: Shallow



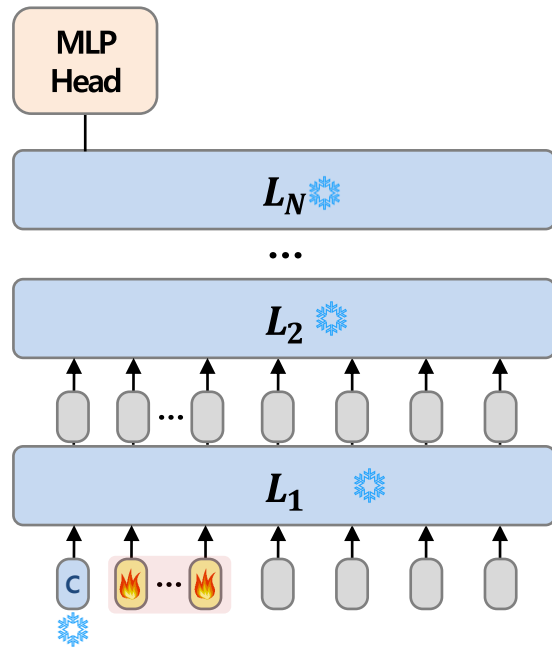
Visual Prompt Tuning: Deep

Visual Prompt Tuning (VPT)

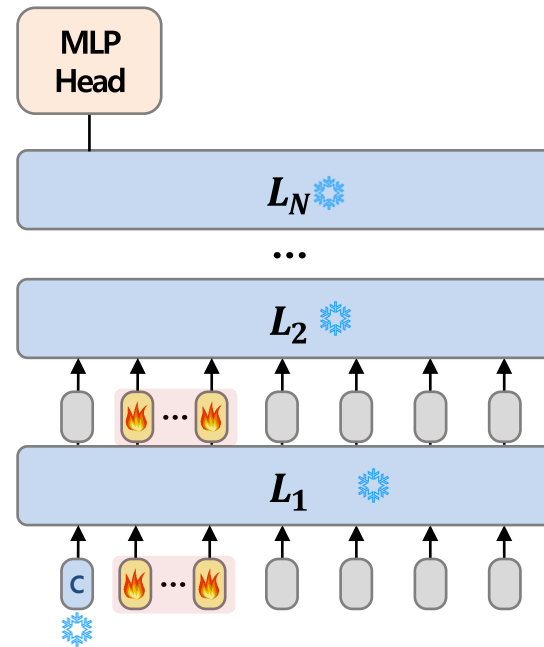
How VPT Works

❖ VPT의 두 가지 Prompt 삽입 방식은 무엇일까?

- **VPT-Shallow:** 첫 번째 layer 입력단에만 prompt 삽입 → 단일 prompt set만 학습, 매우 가벼움
- **VPT-Deep:** 모든 layer 입력마다 서로 다른 prompt 삽입 → layer별 task-specific 정보 반영, 표현력 향상



Visual Prompt Tuning: Shallow



Visual Prompt Tuning: Deep

	추가 파라미터	비율
VPT-Shallow	0.038M	0.04%
VPT-Deep	0.46M	0.53%

Visual Prompt Tuning (VPT)

Experiments

❖ VPT Experiment #1: Parameter Efficiency & Performance

- **비교 대상:** (a) Full fine-tuning / (b) Partial tuning / (c) Extra module
- **저장 비용:** Full fine-tuning **24.02×** ↔ VPT-Shallow **1.04×**, VPT-Deep **1.18×**
- **성능:** VPT-Deep은 **24개 task 중 20개에서 Full Fine-Tuning보다 높은 성능**

	ViT-B/16 (85.8M)	Total params	Scope		Extra params	FGVC	VTAB-1k		
			Input	Backbone			Natural	Specialized	Structured
	Total # of tasks					5	7	4	8
(a)	FULL	24.02×		✓		88.54	75.88	83.36	47.64
(b)	LINEAR	1.02×				79.32 (0)	68.93 (1)	77.16 (1)	26.84 (0)
	PARTIAL-1	3.00×				82.63 (0)	69.44 (2)	78.53 (0)	34.17 (0)
	MLP-3	1.35×			✓	79.80 (0)	67.80 (2)	72.83 (0)	30.62 (0)
(c)	SIDETUNE	3.69×		✓	✓	78.35 (0)	58.21 (0)	68.12 (0)	23.41 (0)
	BIAS	1.05×		✓		88.41 (3)	73.30 (3)	78.25 (0)	44.09 (2)
	ADAPTER	1.23×		✓	✓	85.66 (2)	70.39 (4)	77.11 (0)	33.43 (0)
(ours)	VPT-SHALLOW	1.04×			✓	84.62 (1)	76.81 (4)	79.66 (0)	46.98 (4)
	VPT-DEEP	1.18×	✓			89.11 (4)	78.48 (6)	82.43 (2)	54.98 (8)

Visual Prompt Tuning (VPT)

Experiments

❖ VPT Experiment #1: Parameter Efficiency & Performance

- **비교 대상:** (a) Full fine-tuning / (b) Partial tuning / (c) Extra module
- **저장 비용:** Full fine-tuning **24.02×** ↔ VPT-Shallow **1.04×**, VPT-Deep **1.18×**
- **성능:** VPT-Deep은 **24개 task 중 20개에서 Full Fine-Tuning보다 높은 성능**

ViT-B/16 (85.8M)		Total params	Scope Input Backbone		Extra params	FGVC	VTAB-1k Natural Specialized Structured		
Total # of tasks						5	7	4	8
(a)	FULL	24.02×	✓			88.54	75.88	83.36	47.64
(b)	LINEAR	1.02×				79.32 (0)	68.93 (1)	77.16 (1)	26.84 (0)
	PARTIAL-1	3.00×				82.63 (0)	69.44 (2)	78.53 (0)	34.17 (0)
	MLP-3	1.35×			✓	79.80 (0)	67.80 (2)	72.83 (0)	30.62 (0)
(c)	SIDETUNE	3.69×	✓		✓	78.35 (0)	58.21 (0)	68.12 (0)	23.41 (0)
	BIAS	1.05×	✓			88.41 (3)	73.30 (3)	78.25 (0)	44.09 (2)
	ADAPTER	1.23×	✓		✓	85.66 (2)	70.39 (4)	77.11 (0)	33.43 (0)
(ours)	VPT-SHALLOW	1.04×			✓	84.62 (1)	76.81 (4)	79.66 (0)	46.98 (4)
	VPT-DEEP	1.18×	✓		✓	89.11 (4)	78.48 (6)	82.43 (2)	54.98 (8)

Visual Prompt Tuning (VPT)

Experiments

❖ VPT Experiment #1: Parameter Efficiency & Performance

- **비교 대상:** (a) Full fine-tuning / (b) Partial tuning / (c) Extra module
- **저장 비용:** Full fine-tuning **24.02×** ↔ VPT-Shallow **1.04×**, VPT-Deep **1.18×**
- **성능:** VPT-Deep은 **24개 task 중 20개에서 Full Fine-Tuning보다 높은 성능**

ViT-B/16 (85.8M)		Total params	Scope		Extra params	FGVC	VTAB-1k		
			Input	Backbone			Natural	Specialized	Structured
Total # of tasks						5	7	4	8
(a)	FULL	24.02×		✓		88.54	75.88	83.36	47.64
(b)	LINEAR	1.02×				79.32 (0)	68.93 (1)	77.16 (1)	26.84 (0)
	PARTIAL-1	3.00×				82.63 (0)	69.44 (2)	78.53 (0)	34.17 (0)
	MLP-3	1.35×			✓	79.80 (0)	67.80 (2)	72.83 (0)	30.62 (0)
(c)	SIDETUNE	3.69×		✓	✓	78.35 (0)	58.21 (0)	68.12 (0)	23.41 (0)
	BIAS	1.05×		✓		88.41 (3)	73.30 (3)	78.25 (0)	44.09 (2)
	ADAPTER	1.23×		✓	✓	85.66 (2)	70.39 (4)	77.11 (0)	33.43 (0)
(ours)	VPT-SHALLOW	1.04×			✓	84.62 (1)	76.81 (4)	79.66 (0)	46.98 (4)
	VPT-DEEP	1.18×	✓			89.11 (4)	78.48 (6)	82.43 (2)	54.98 (8)

Visual Prompt Tuning (VPT)

Experiments

❖ VPT Experiment #1: Parameter Efficiency & Performance

- **비교 대상:** (a) Full fine-tuning / (b) Partial tuning / (c) Extra module
- **저장 비용:** Full fine-tuning **24.02×** ↔ VPT-Shallow **1.04×**, VPT-Deep **1.18×**
- **성능:** VPT-Deep은 **24개 task 중 20개에서 Full Fine-Tuning보다 높은 성능**

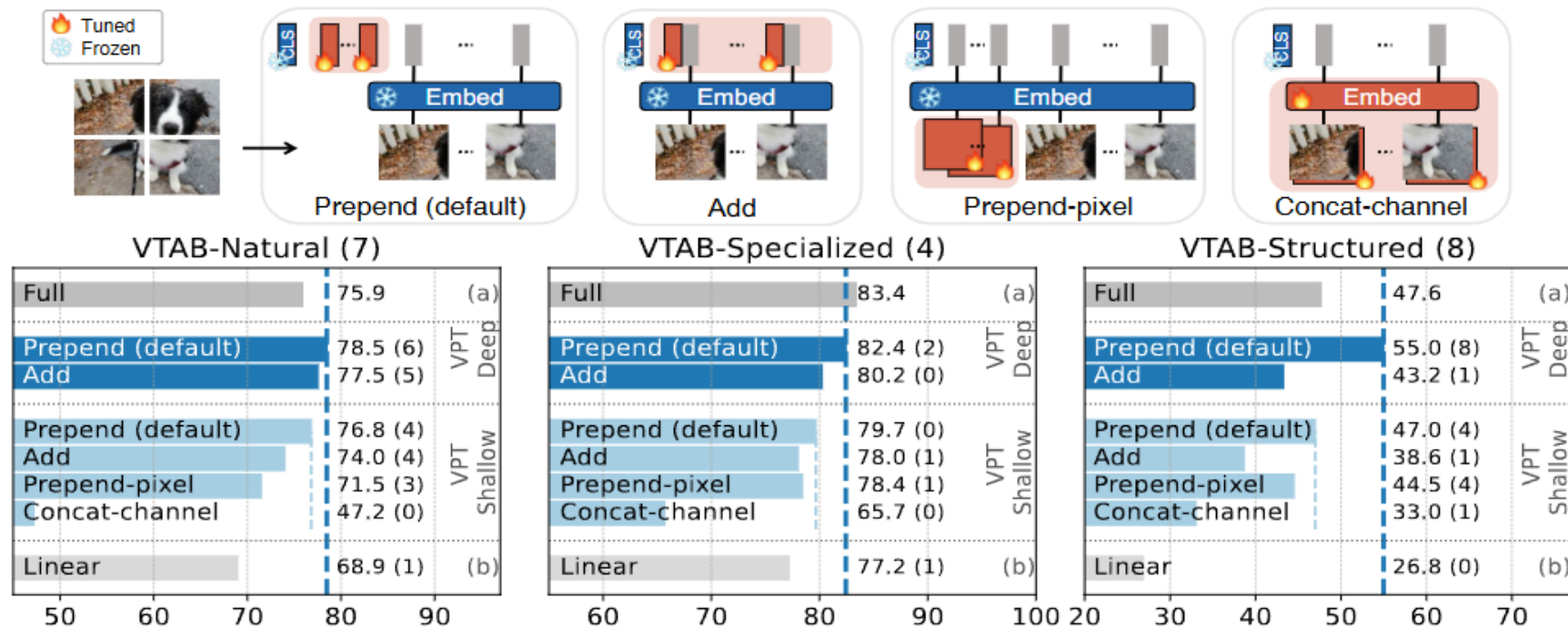
	ViT-B/16 (85.8M)	Total params	Scope		Extra params	FGVC	VTAB-1k		
			Input	Backbone			Natural	Specialized	Structured
	Total # of tasks					5	7	4	8
(a)	FULL	24.02×		✓		88.54	75.88	83.36	47.64
(b)	LINEAR	1.02×				79.32 (0)	68.93 (1)	77.16 (1)	26.84 (0)
	PARTIAL-1	3.00×				82.63 (0)	69.44 (2)	78.53 (0)	34.17 (0)
	MLP-3	1.35×			✓	79.80 (0)	67.80 (2)	72.83 (0)	30.62 (0)
(c)	SIDETUNE	3.69×		✓	✓	78.35 (0)	58.21 (0)	68.12 (0)	23.41 (0)
	BIAS	1.05×		✓		88.41 (3)	73.30 (3)	78.25 (0)	44.09 (2)
	ADAPTER	1.23×		✓	✓	85.66 (2)	70.39 (4)	77.11 (0)	33.43 (0)
(ours)	VPT-SHALLOW	1.04×			✓	84.62 (1)	76.81 (4)	79.66 (0)	46.98 (4)
	VPT-DEEP	1.18×	✓			89.11 (4)	78.48 (6)	82.43 (2)	54.98 (8)

Visual Prompt Tuning (VPT)

Experiments

❖ VPT Experiment #2: Ablation on Model Design Variants ① Prompt Location

- Embedding-space의 Prepend 방식이 가장 안정적인 성능을 보였으며, Add 및 pixel/channel-level 방식은 상대적으로 성능이 저하됨

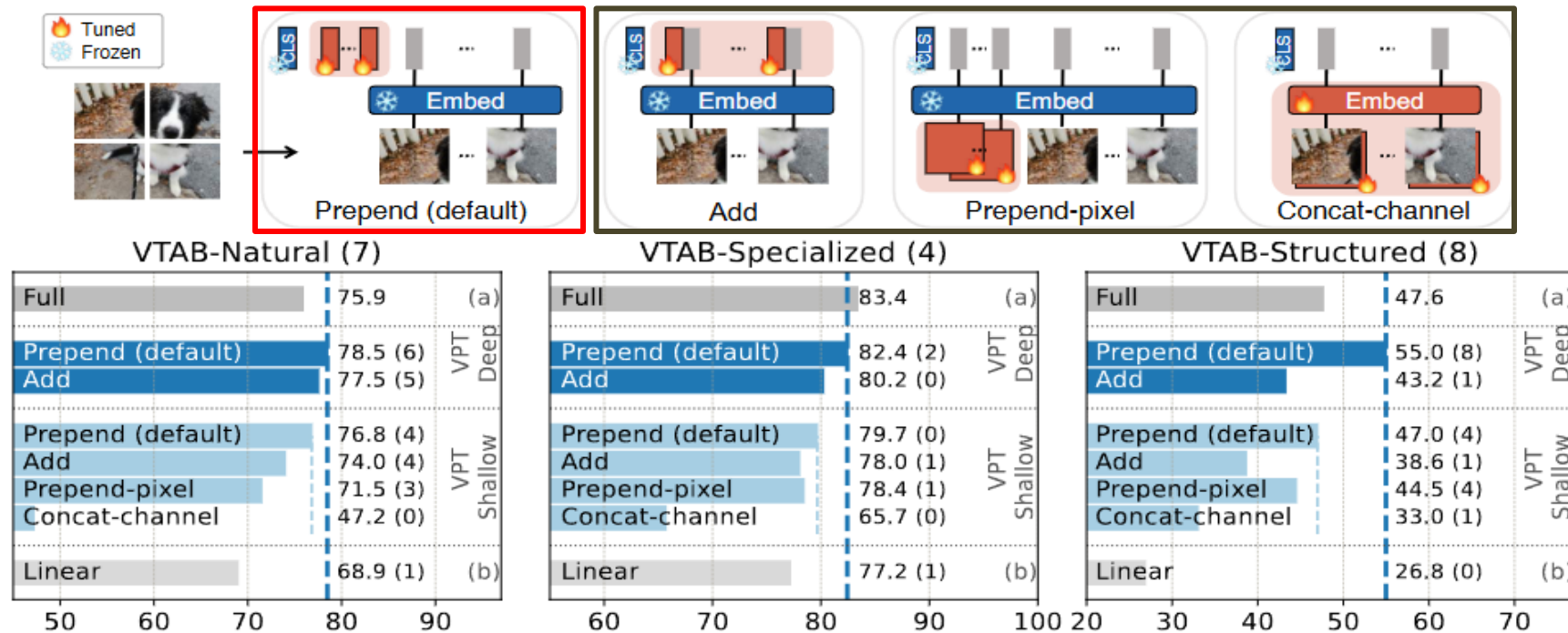


Visual Prompt Tuning (VPT)

Experiments

❖ VPT Experiment #2: Ablation on Model Design Variants ① Prompt Location

- Embedding-space의 Prepend 방식이 가장 안정적인 성능을 보였으며, Add 및 pixel/channel-level 방식은 상대적으로 성능이 저하됨

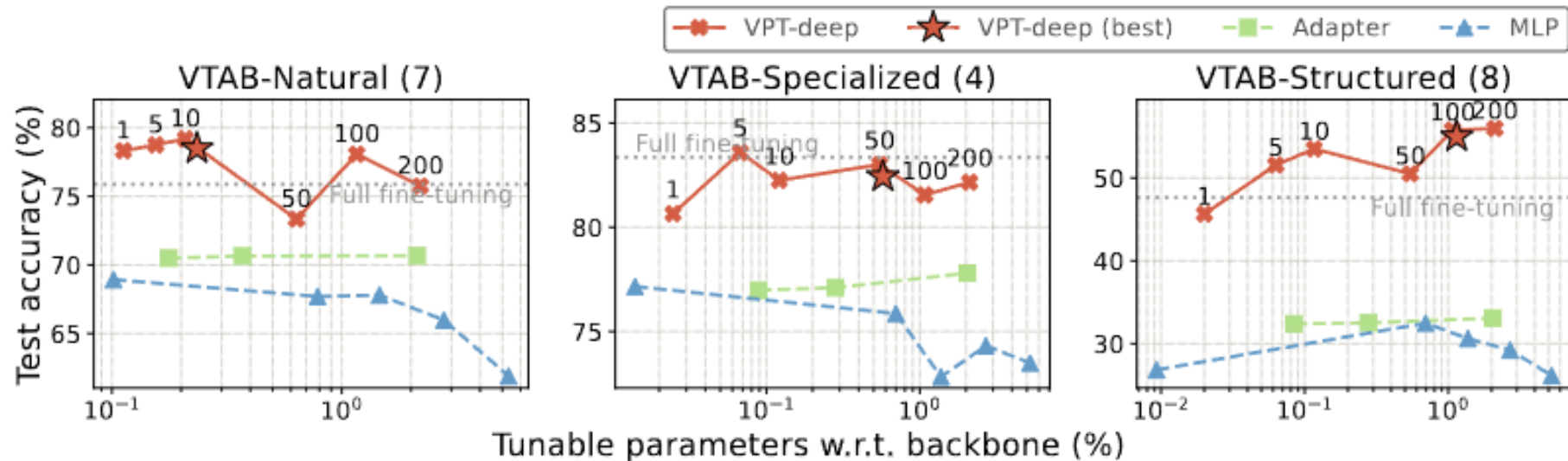


Visual Prompt Tuning (VPT)

Experiments

❖ VPT Experiment #2: Ablation on Model Design Variants ② Prompt Length

- Prompt length 증가가 성능 향상을 보장하지 않으며, 최적 prompt length는 Natural 10, Specialized 50, Structured 100으로 서로 다름

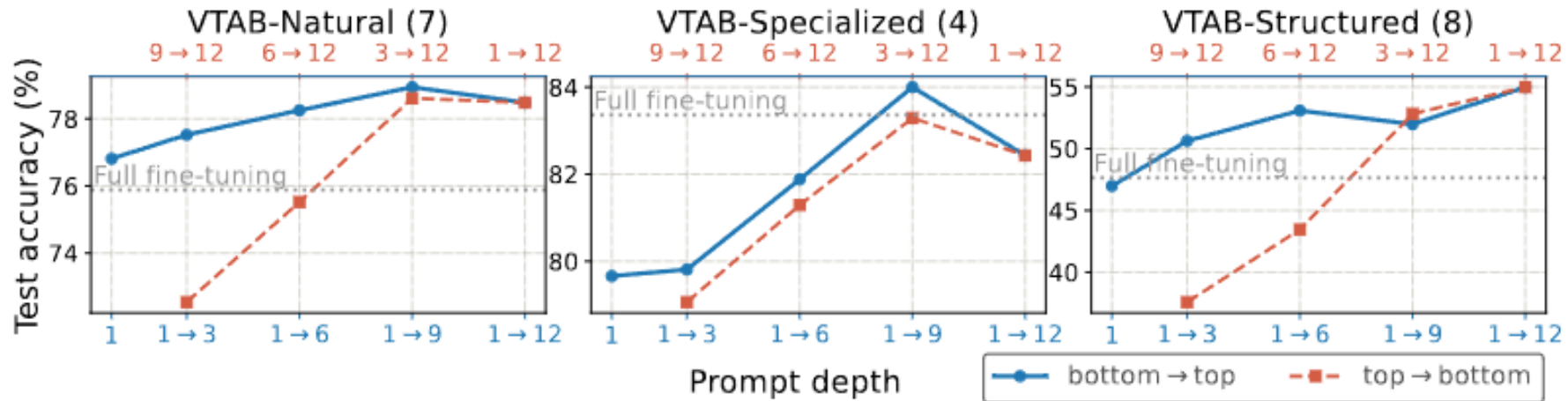


Visual Prompt Tuning (VPT)

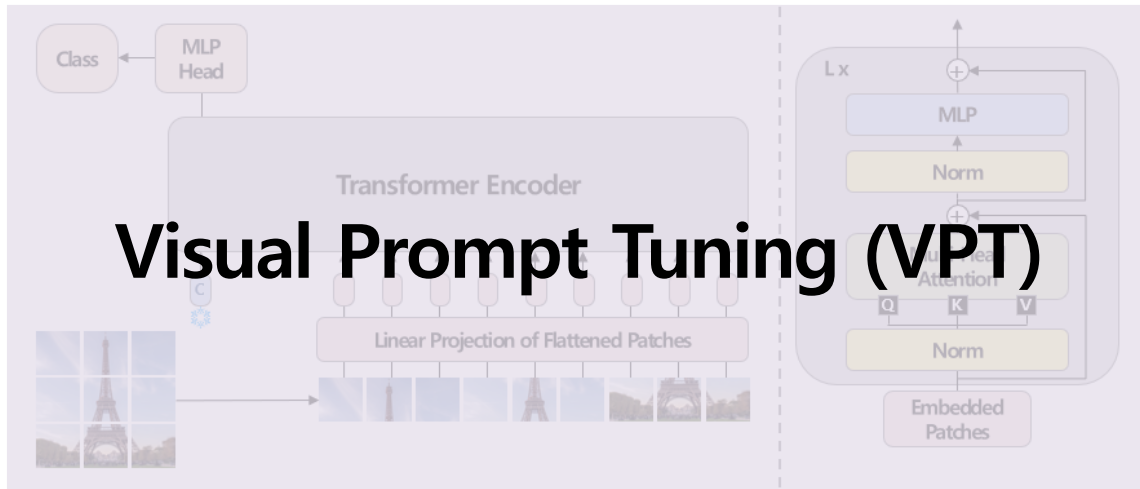
Experiments

❖ VPT Experiment #3: Ablation on Model Design Variants ③ Prompt Depth

- Prompt를 앞쪽 layer부터 삽입하는 Bottom→Top 방식이 전반적으로 더 높은 성능을 보여, 초기 layer에서부터 prompt가 개입하는 것이 더 효과적임을 확인

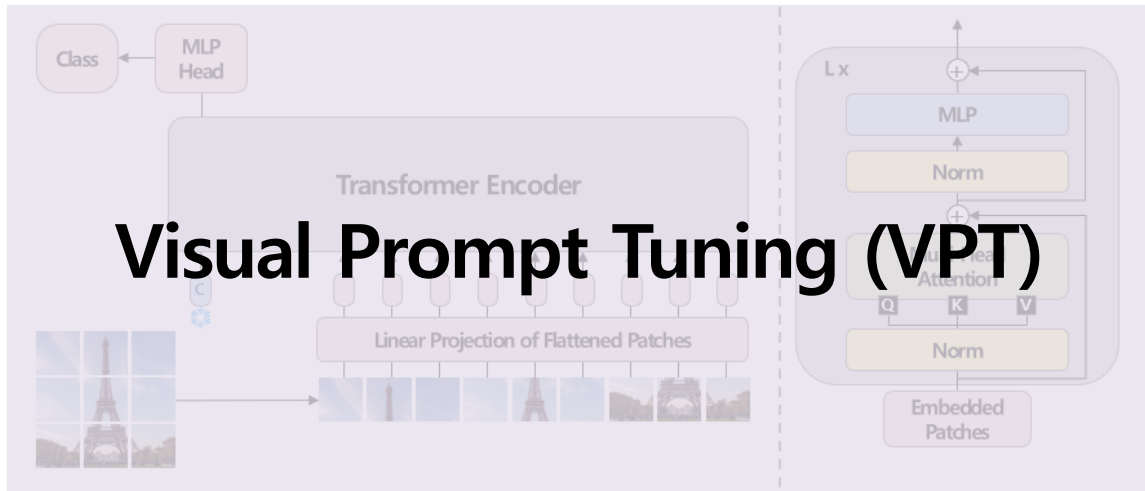


Visual Prompt Tuning (VPT)



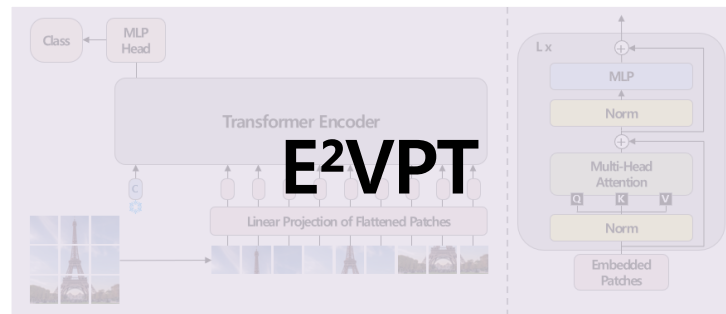
➔ VPT의 핵심 설계 과제:
Prompt를 어디에, 얼마나 개입시킬 것인가?

Visual Prompt Tuning (VPT)

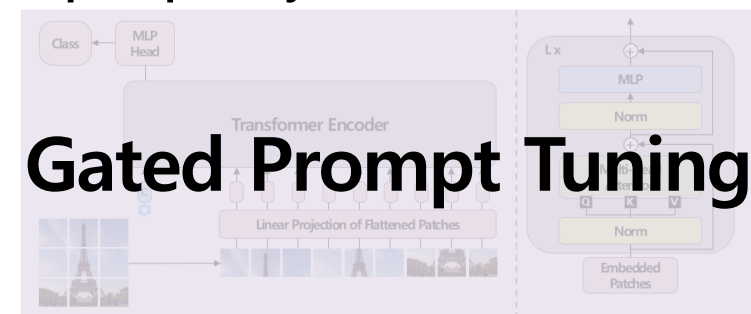


➔ VPT의 핵심 설계 과제:
Prompt를 어디에, 얼마나 개입시킬 것인가?

self-attention 구조 내부로도 확장:



prompt의 layer별 개입 정도를 학습:

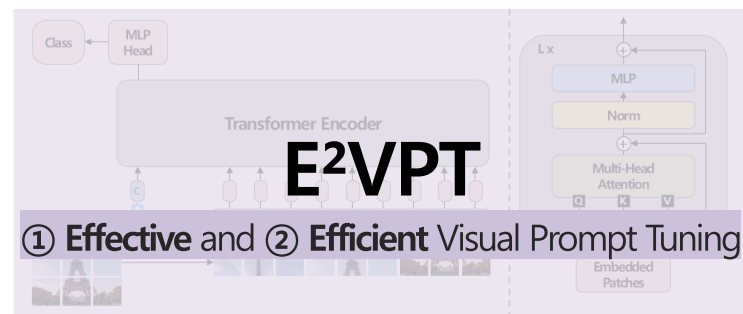


E²VPT



➔ VPT의 핵심 설계 과제:
Prompt를 어디에, 얼마나 개입시킬 것인가?

self-attention 구조 내부로도 확장:

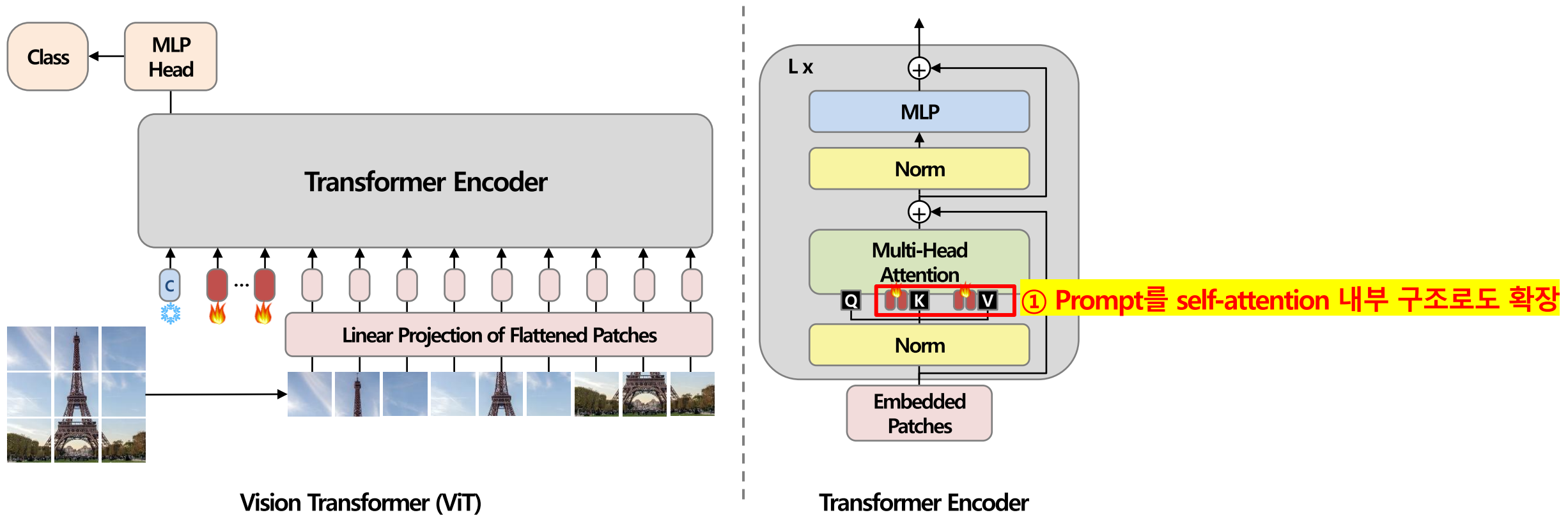


E²VPT

How E²VPT Works

❖ ① Effective Prompting: Key-Value Prompts

- 기존 Input Visual Prompt에 더해, self-attention의 Key-Value에 task-specific prompt를 직접 삽입

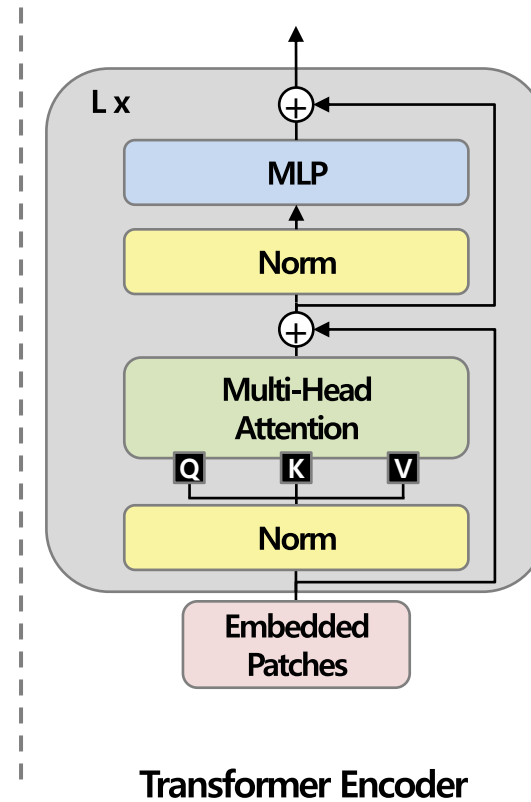
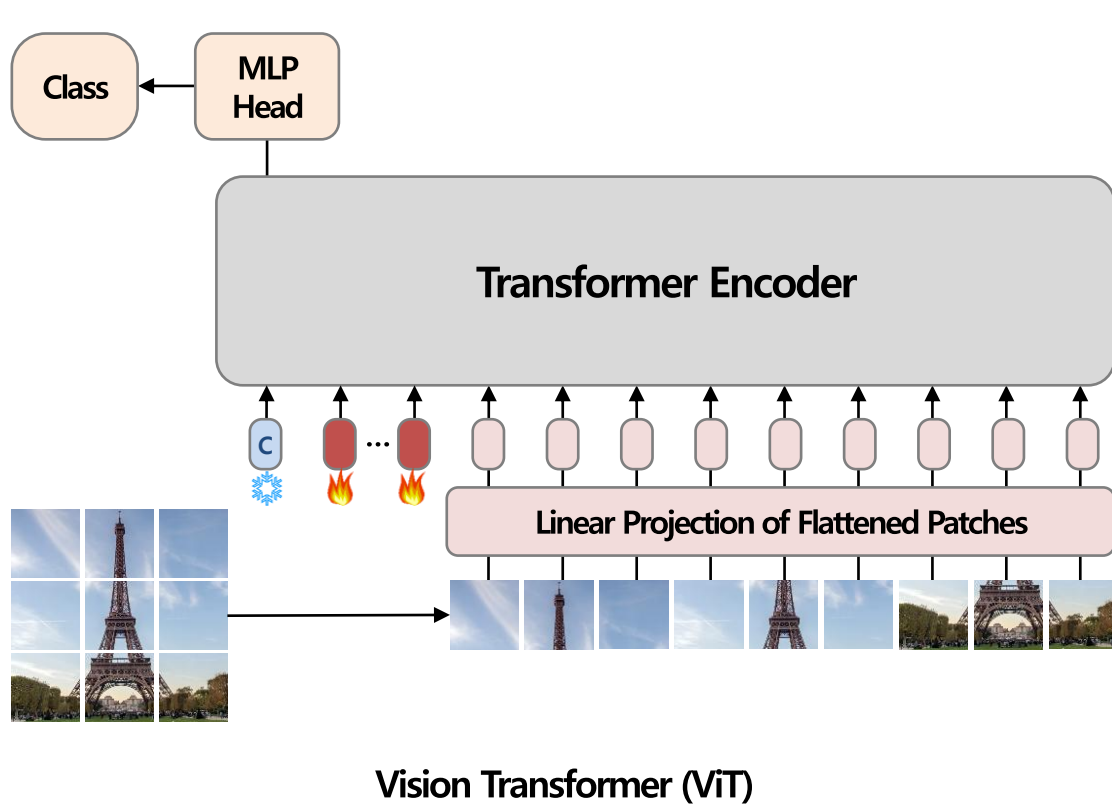


E²VPT

How E²VPT Works

❖ ① Effective Prompting: Key-Value Prompts

- 기존 Input Visual Prompt에 더해, self-attention의 Key-Value에 task-specific prompt를 직접 삽입

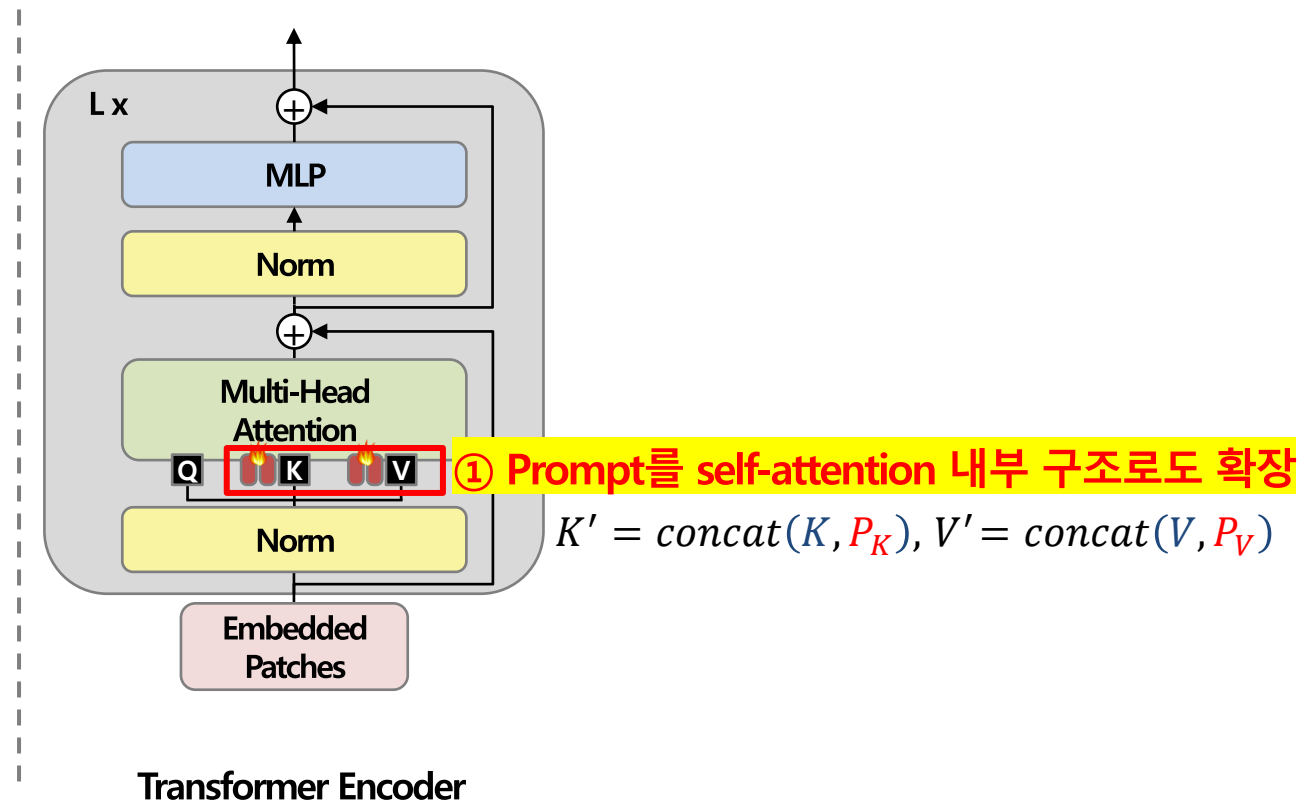
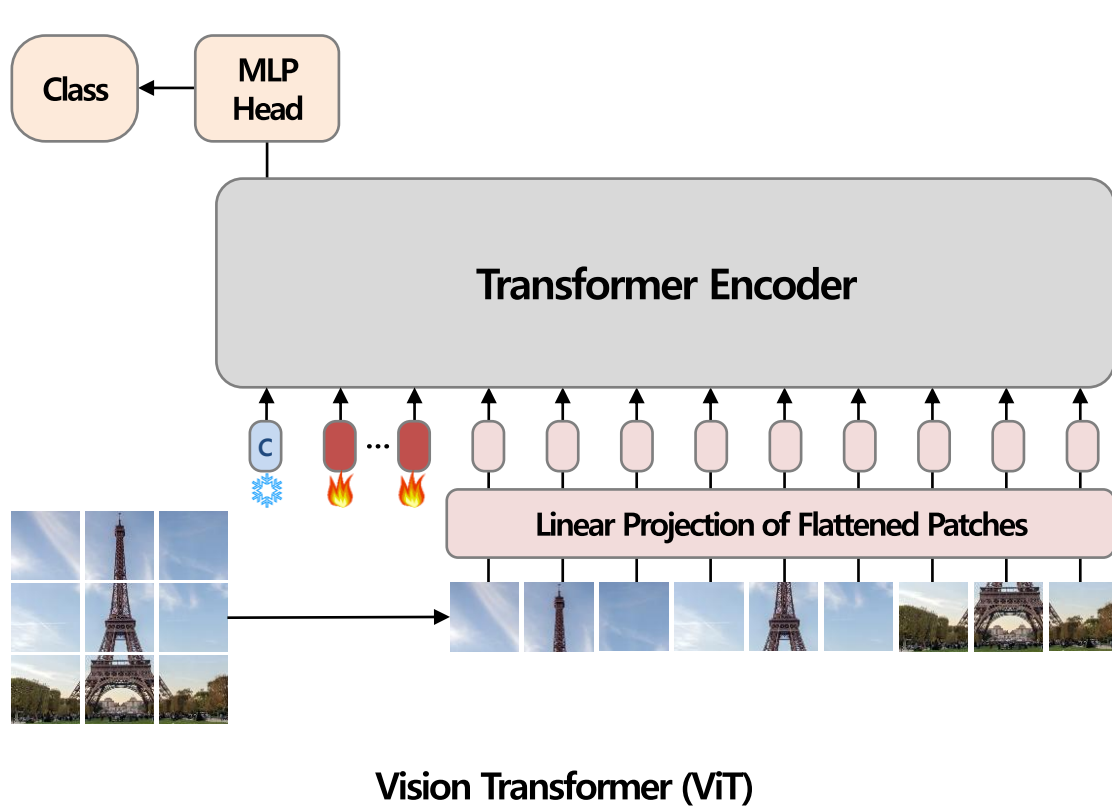


E²VPT

How E²VPT Works

❖ ① Effective Prompting: Key-Value Prompts

- 기존 Input Visual Prompt에 더해, self-attention의 Key-Value에 task-specific prompt를 직접 삽입

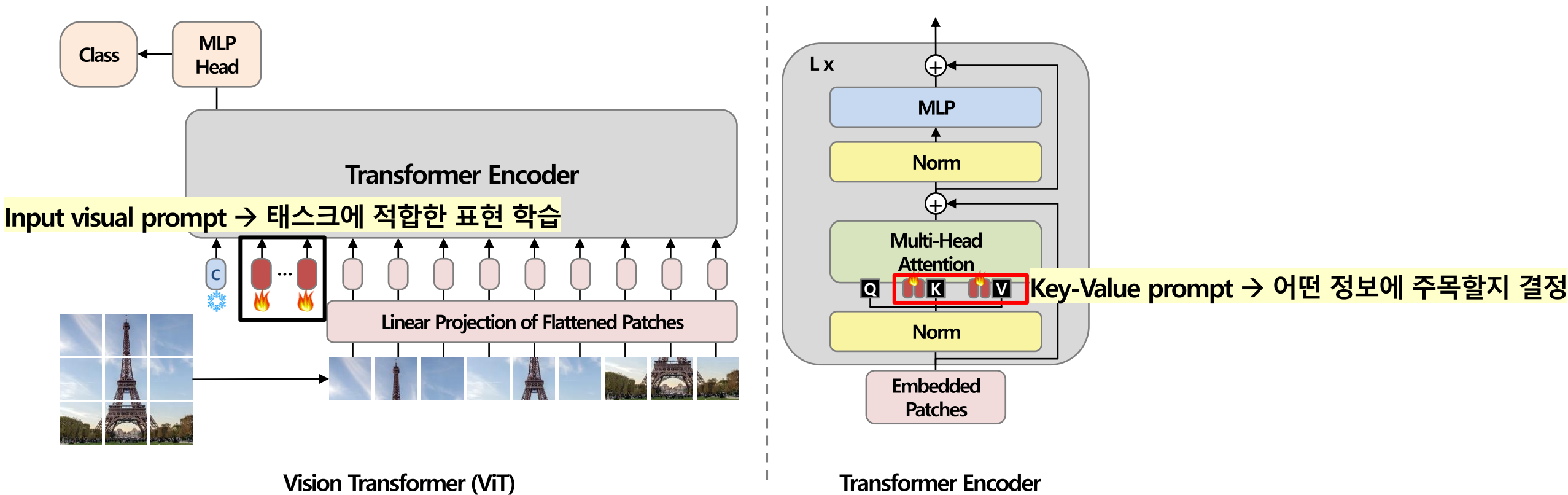


E²VPT

How E²VPT Works

❖ ① Effective Prompting: Key-Value Prompts

- 기존 Input Visual Prompt에 더해, self-attention의 Key-Value에 task-specific prompt를 직접 삽입

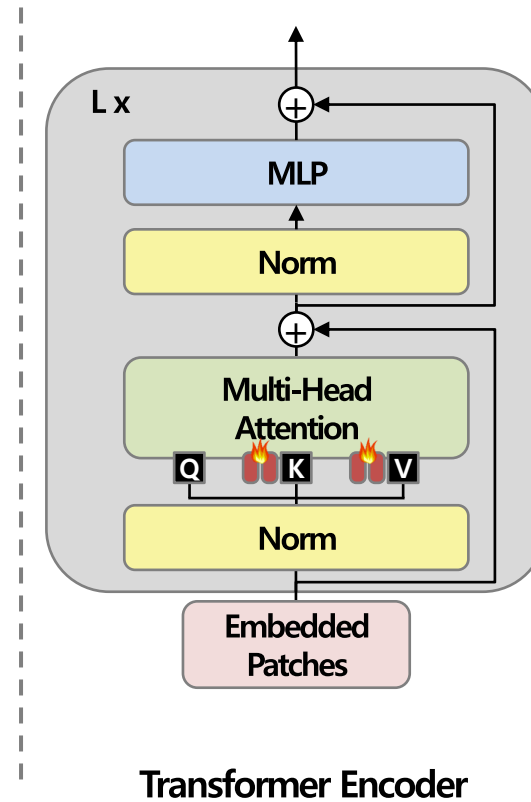
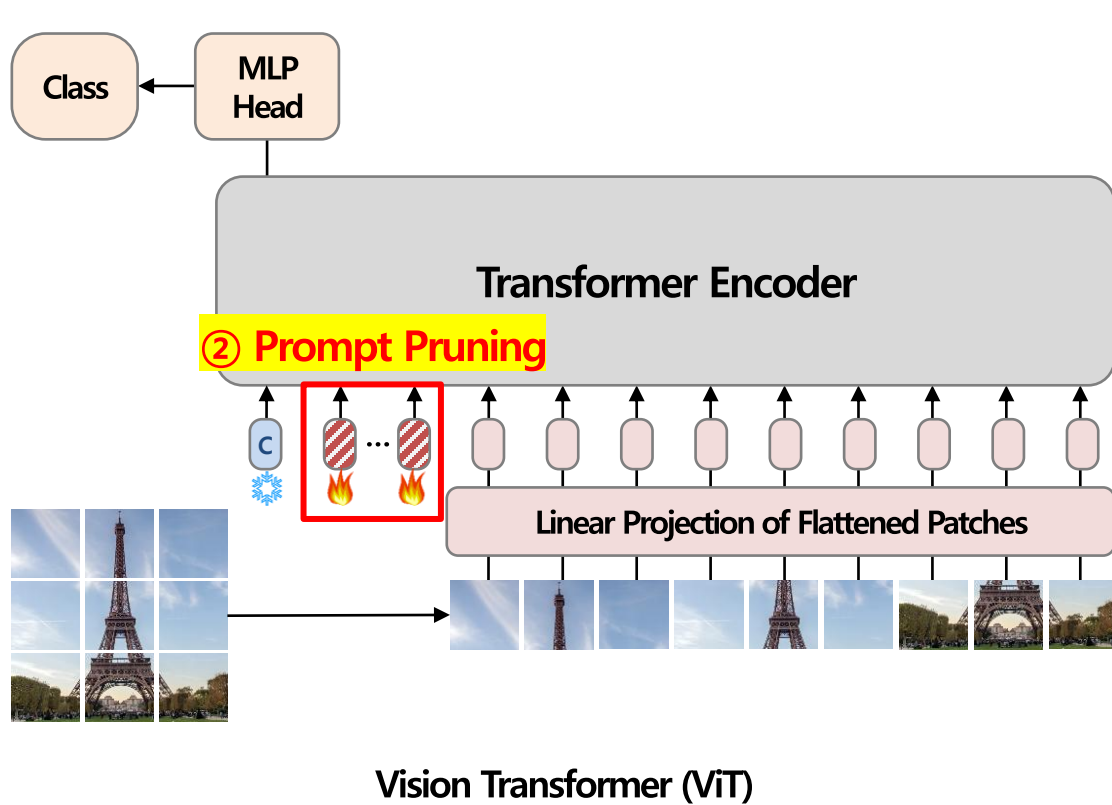


E²VPT

How E²VPT Works

❖ ② Efficient Prompting: Prompt Pruning

- 중요도가 낮은 visual prompt의 token 또는 segment를 제거하여 성능 저하 없이 파라미터 효율성 향상



E²VPT

How E²VPT Works

❖ ② Efficient Prompting: Prompt Pruning

- **Lottery Ticket Hypothesis (ICLR 2019):** 큰 신경망 안에 전체 네트워크와 같은 성능을 낼 수 있는 작은 부분 네트워크가 존재한다는 가설

Published as a conference paper at ICLR 2019

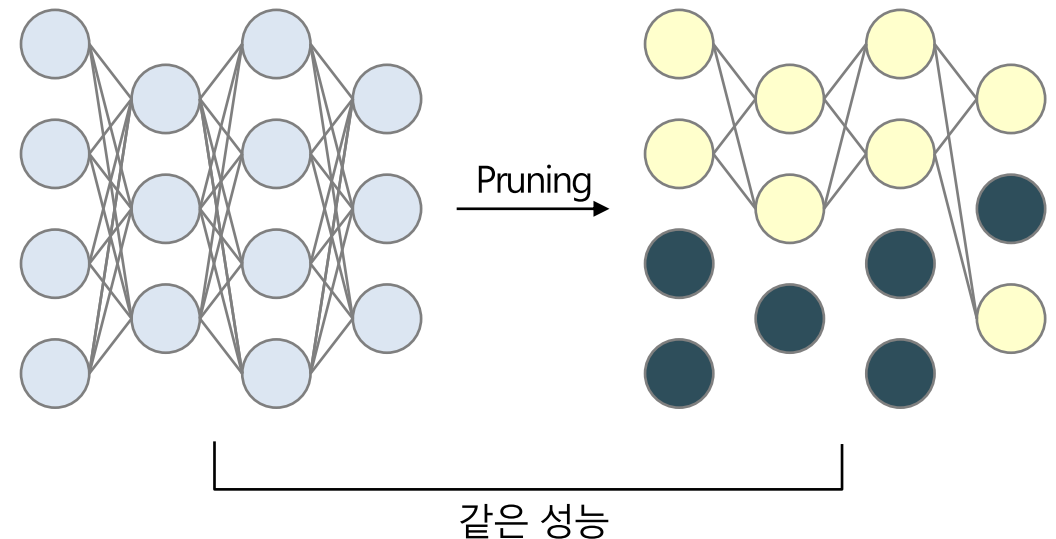
THE LOTTERY TICKET HYPOTHESIS: FINDING SPARSE, TRAINABLE NEURAL NETWORKS

Jonathan Frankle
MIT CSAIL
jfrankle@csail.mit.edu

Michael Carbin
MIT CSAIL
mcarbin@csail.mit.edu

ABSTRACT

Neural network pruning techniques can reduce the parameter counts of trained networks by over 90%, decreasing storage requirements and improving computational performance of inference without compromising accuracy. However, contemporary experience is that the sparse architectures produced by pruning are difficult to train from the start, which would similarly improve training performance.

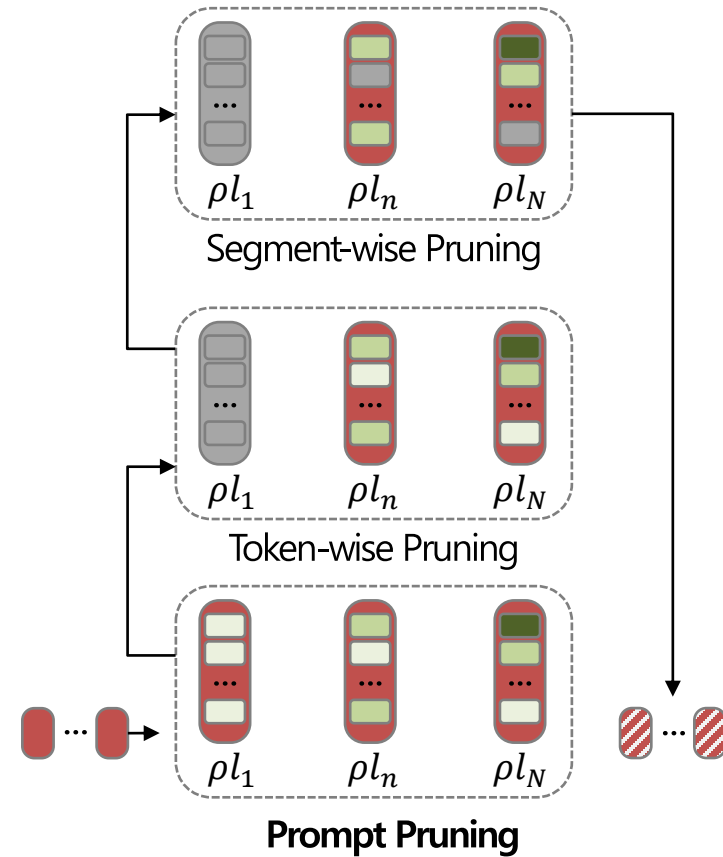
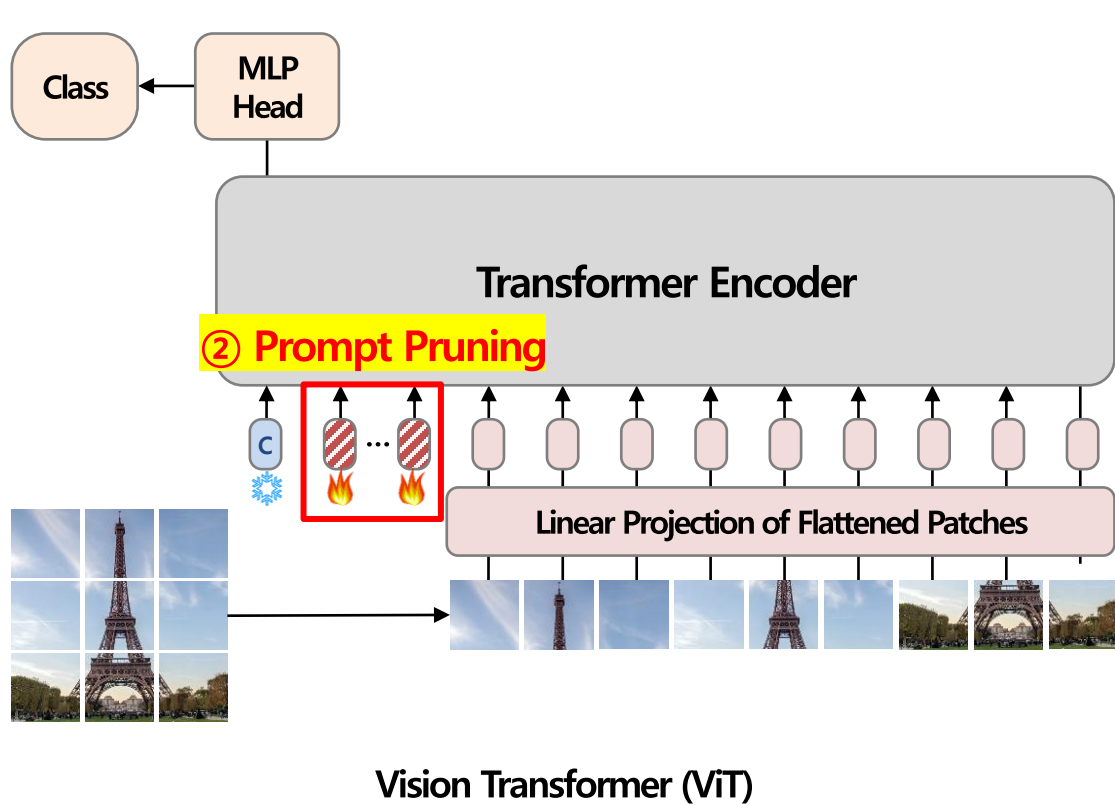


E²VPT

How E²VPT Works

❖ ② Efficient Prompting: Prompt Pruning

- 중요도가 낮은 visual prompt의 token 또는 segment를 제거하여 성능 저하 없이 파라미터 효율성 향상

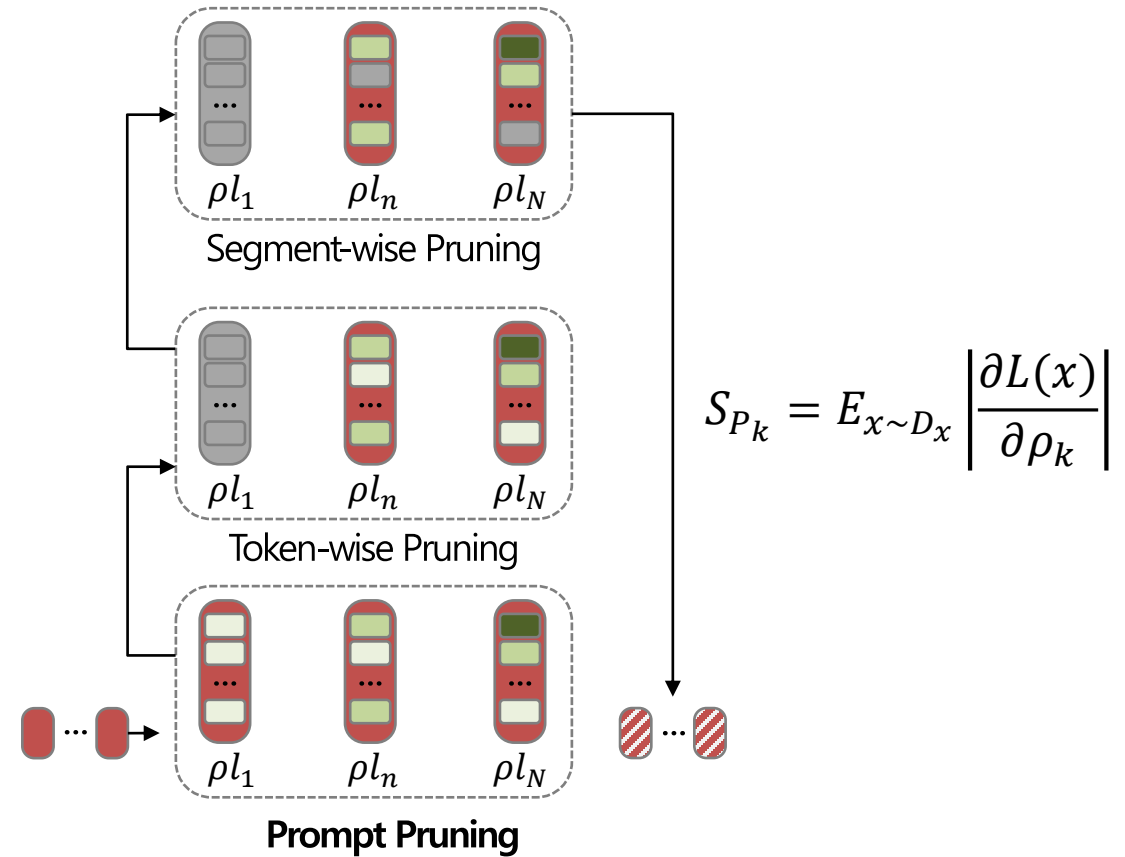
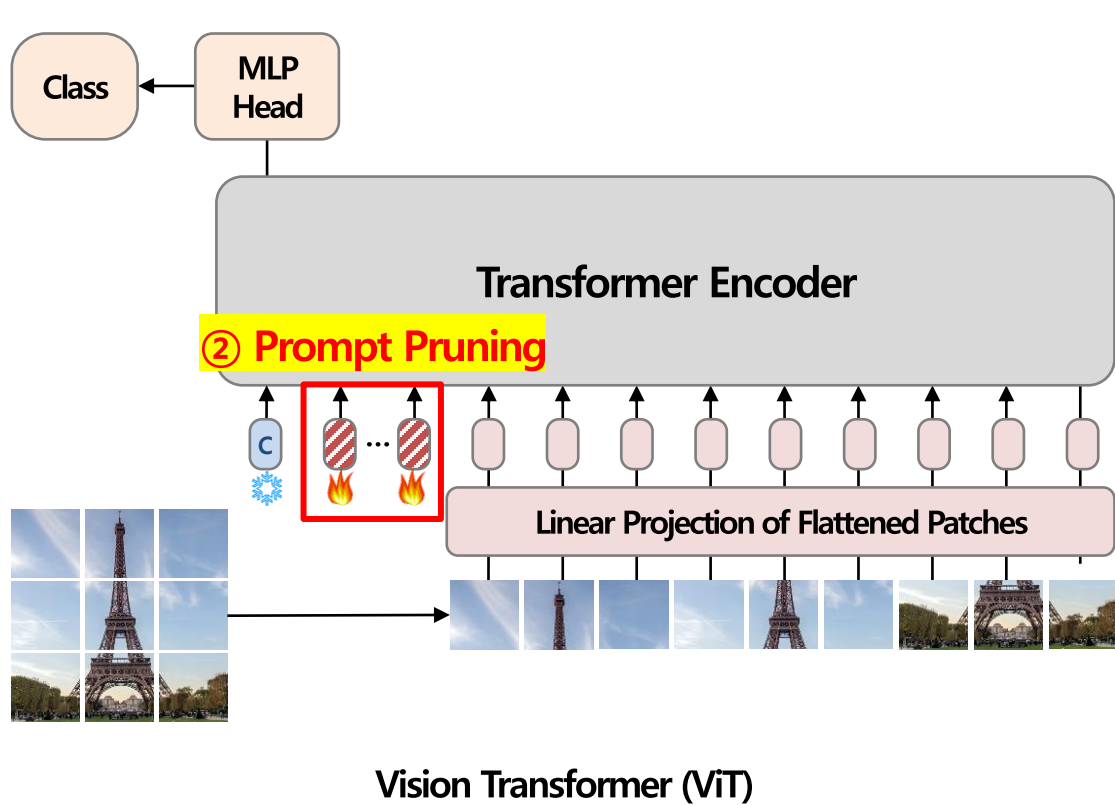


E²VPT

How E²VPT Works

❖ ② Efficient Prompting: Prompt Pruning

- 중요도가 낮은 visual prompt의 token 또는 segment를 제거하여 성능 저하 없이 파라미터 효율성 향상

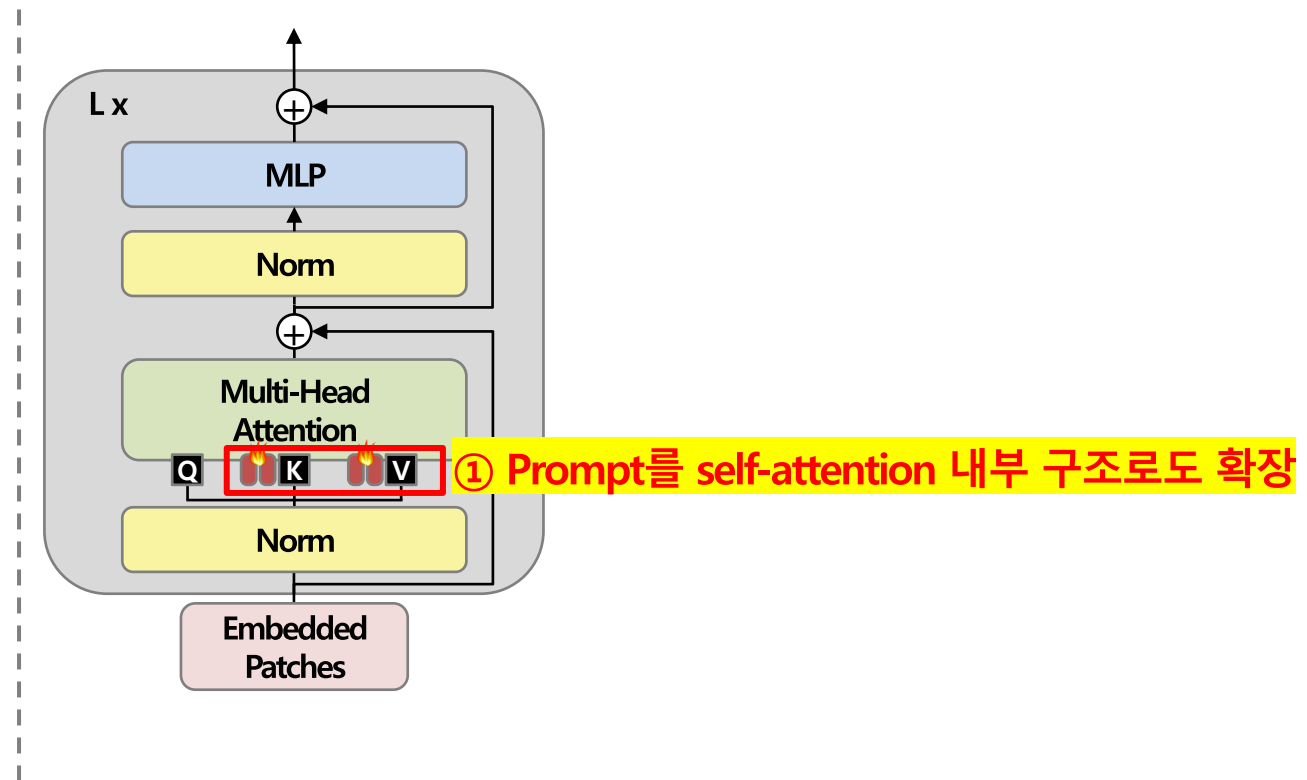
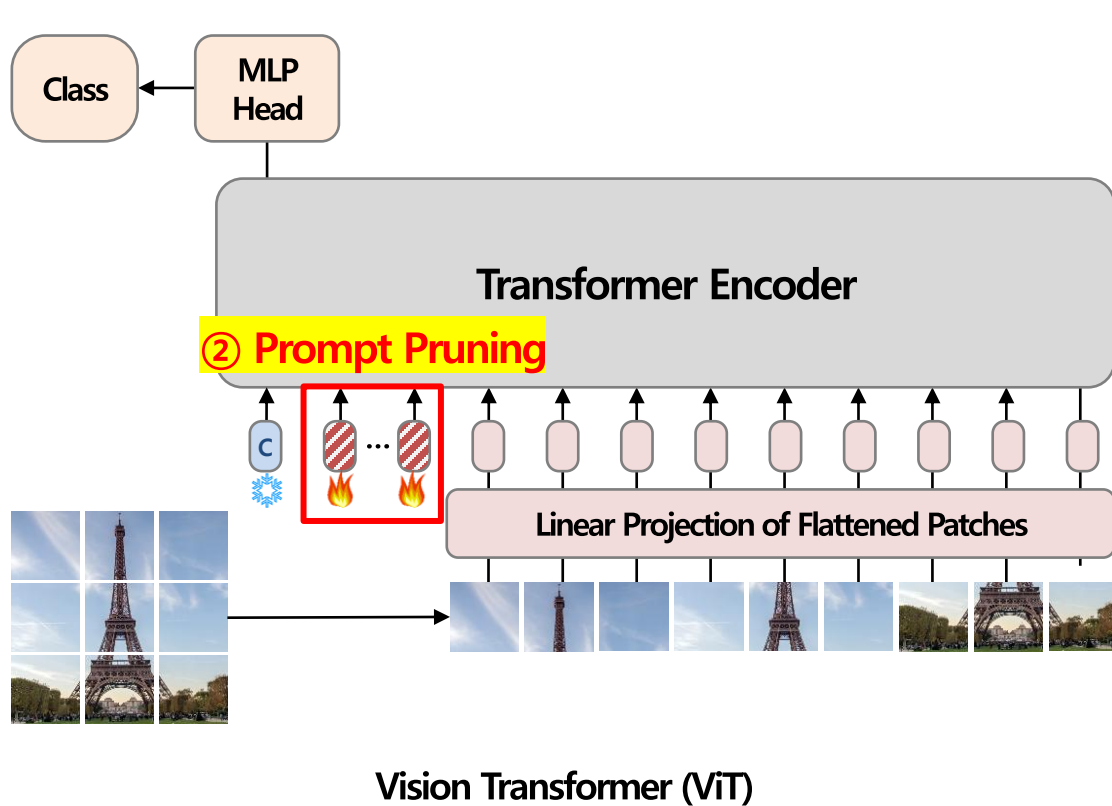


E²VPT

How E²VPT Works

❖ ② Efficient Prompting: Prompt Pruning

- 중요도가 낮은 visual prompt의 token 또는 segment를 제거하여 성능 저하 없이 파라미터 효율성 향상



E²VPT

Experiments

❖ E²VPT Experiment #1: Parameter Efficiency & Performance

- **비교 대상:** Full fine-tuning / Partial tuning / Extra module
- **학습 파라미터:** VPT **0.73%** → E²VPT **1.18%**
- **성능:** E²VPT는 24개 task 중 **VPT 대비 21개, Full Fine-Tuning 대비 19개** task에서 높은 성능 → 더 적은 학습 파라미터로 더 높은 성능 달성

ViT-Base/16 [12] (85.8M)	Tuned/ Total	Scope		Extra params	FGVC [34] [5]	VTAB-1k [96] [19]		
		Input	Backbone			Natural [7]	Specialized [4]	Structured [8]
Full [CVPR22][32]	100.00%		✓		88.54%	75.88%	83.36%	47.64%
Linear [CVPR22][32]	0.08%				79.32% [0]	68.93% [1]	77.16% [1]	26.84% [0]
Partial-1 [NeurIPS14][91]	8.34%				82.63% [0]	69.44% [2]	78.53% [0]	34.17% [0]
MLP-3 [CVPR20][9]	1.44%			✓	79.80% [0]	67.80% [2]	72.83% [0]	30.62% [0]
Sidetune [ECCV20][98]	10.08%		✓	✓	78.35% [0]	58.21% [0]	68.12% [0]	23.41% [0]
Bias [NeurIPS17][70]	0.80%		✓		88.41% [3]	73.30% [3]	78.25% [0]	44.09% [2]
Adapter [NeurIPS20][6]	1.02%		✓	✓	85.66% [2]	70.39% [4]	77.11% [0]	33.43% [0]
VPT [ECCV22][34]	0.73%	✓		✓	89.11% [4]	78.48% [6]	82.43% [2]	54.98% [8]
Ours	0.39%	✓	✓	✓	89.22% [4] {4}	80.01% [6] {5}	84.43% [3] {4}	57.39% [8] {7}

E²VPT

Experiments

❖ E²VPT Experiment #2: Ablation on Key Design Components

- 실험 목적: E²VPT의 세 구성 요소별 기여도 분석
- 구성 요소: Visual Prompt / Key-Value Prompt / Pruning & Rewinding
- 평가 태스크: VTAB-1k Natural **SVHN**, FGVC **NABirds**

Fine-tuning Techniques			VTAB-1k Natural SVHN [62]			FGVC NABirds [77]		
Visual Prompts	Key-Value Prompts	Pruning & Rewinding	Pruning	Tuned / Total	Accuracy	Pruning	Tuned / Total	Accuracy
✓			0.0%	0.54%	78.1%	0.0%	1.02%	84.2%
✓	✓		0.0%	0.55%	83.8%	0.0%	1.05%	84.5%
✓		✓	56.3%	0.42%	79.0%	34.4%	0.63%	84.2%
✓	✓	✓	62.5%	0.43%	85.3%	40.0%	0.65%	84.6%

E²VPT

Experiments

❖ E²VPT Experiment #2: Ablation on Key Design Components

- 실험 목적: E²VPT의 세 구성 요소별 기여도 분석
- 구성 요소: Visual Prompt / Key-Value Prompt / Pruning & Rewinding
- 평가 태스크: VTAB-1k Natural **SVHN**, FGVC **NABirds**

Fine-tuning Techniques			VTAB-1k Natural SVHN [62]			FGVC NABirds [77]		
Visual Prompts	Key-Value Prompts	Pruning & Rewinding	Pruning	Tuned / Total	Accuracy	Pruning	Tuned / Total	Accuracy
✓			0.0%	0.54%	78.1%	0.0%	1.02%	84.2%
✓	✓		0.0%	0.55%	83.8%	0.0%	1.05%	84.5%
✓		✓	56.3%	0.42%	79.0%	34.4%	0.63%	84.2%
✓	✓	✓	62.5%	0.43%	85.3%	40.0%	0.65%	84.6%

Key-Value Prompt 추가 시, 파라미터 증가는 거의 없지만 성능 향상

E²VPT

Experiments

❖ E²VPT Experiment #2: Ablation on Key Design Components

- 실험 목적: E²VPT의 세 구성 요소별 기여도 분석
- 구성 요소: Visual Prompt / Key-Value Prompt / Pruning & Rewinding
- 평가 태스크: VTAB-1k Natural **SVHN**, FGVC **NABirds**

Fine-tuning Techniques			VTAB-1k Natural SVHN [62]			FGVC NABirds [77]		
Visual Prompts	Key-Value Prompts	Pruning & Rewinding	Pruning	Tuned / Total	Accuracy	Pruning	Tuned / Total	Accuracy
✓			0.0%	0.54%	78.1%	0.0%	1.02%	84.2%
✓	✓		0.0%	0.55%	83.8%	0.0%	1.05%	84.5%
✓		✓	56.3%	0.42%	79.0%	34.4%	0.63%	84.2%
✓	✓	✓	62.5%	0.43%	85.3%	40.0%	0.65%	84.6%

Pruning & Rewinding 적용 시, 많은 프롬프트를 제거하고도 성능 유지

E²VPT

Experiments

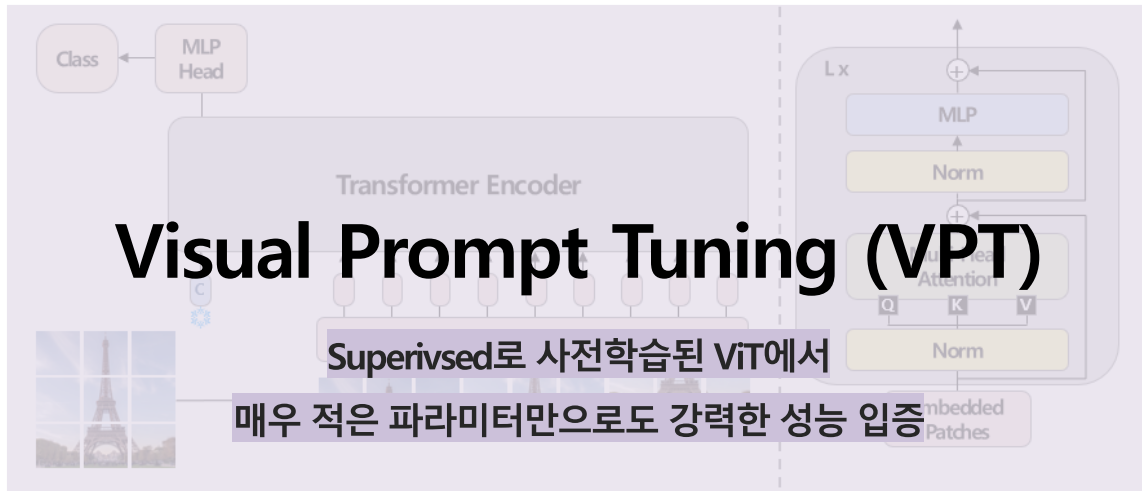
❖ E²VPT Experiment #2: Ablation on Key Design Components

- 실험 목적: E²VPT의 세 구성 요소별 기여도 분석
- 구성 요소: Visual Prompt / Key-Value Prompt / Pruning & Rewinding
- 평가 태스크: VTAB-1k Natural **SVHN**, FGVC **NABirds**

Fine-tuning Techniques			VTAB-1k Natural SVHN [62]			FGVC NABirds [77]		
Visual Prompts	Key-Value Prompts	Pruning & Rewinding	Pruning	Tuned / Total	Accuracy	Pruning	Tuned / Total	Accuracy
✓			0.0%	0.54%	78.1%	0.0%	1.02%	84.2%
✓	✓		0.0%	0.55%	83.8%	0.0%	1.05%	84.5%
✓		✓	56.3%	0.42%	79.0%	34.4%	0.63%	84.2%
✓	✓	✓	62.5%	0.43%	85.3%	40.0%	0.65%	84.6%

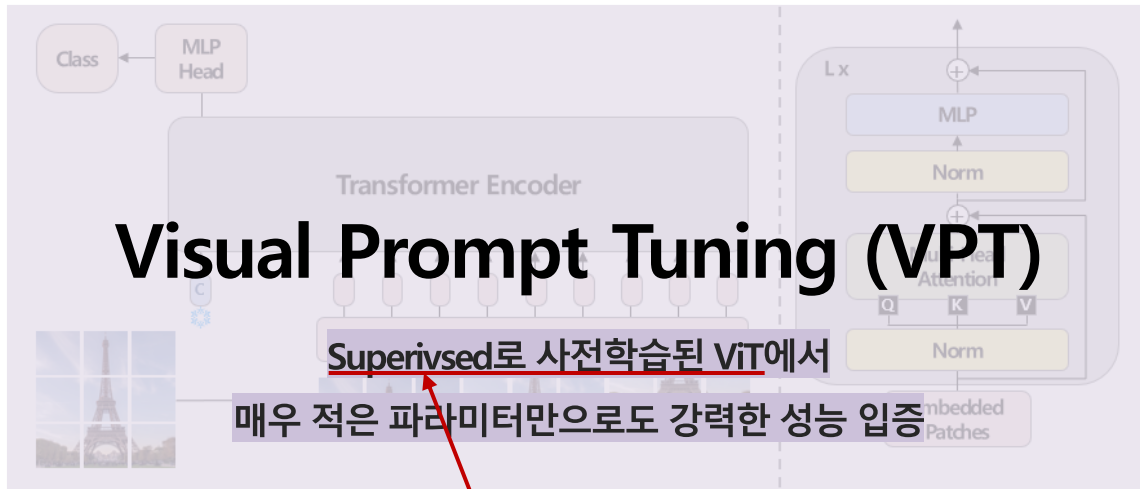
세 요소를 모두 적용 시, 최고 성능 달성

Gated Prompt Tuning



➔ VPT의 핵심 설계 과제:
Prompt를 어디에, 얼마나 개입시킬 것인가?

Gated Prompt Tuning



→ 여전히 self-supervised로 사전학습된 ViT에서는
VPT가 잘 작동하지 않는다

➔ VPT의 핵심 설계 과제:
Prompt를 어디에, 얼마나 개입시킬 것인가?

Gated Prompt Tuning

Self-supervised Learning

❖ Self-supervised Learning이란?

- 정의: 정답 레이블 없이 데이터 자체에서 학습 신호를 생성하여 representation을 학습하는 방식
- 대표적인 방법
 - MAE**: 이미지 일부 patch를 가린 뒤 가려진 부분을 복원하도록 학습 → 이미지의 **구조 · 형태와 같은 시각적 정보**를 학습
 - MoCo v3**: 동일 이미지의 서로 다른 augmented view는 가깝게 표현되도록 학습 → augmentation에 **불변한 representation** 학습
- 즉, self-supervised ViT는 layer별 representation 특성이 다르지만, VPT-Deep는 모든 layer에 prompt를 일괄 삽입 → 성능 저하

		MAE				MoCo v3			
ViT-B/16 (85.8M)		Total params	Natural	Specialized	Structured	Total params	Natural	Specialized	Structured
Total # of tasks			7	4	8		7	4	8
(a)	FULL	19.01×	59.29	79.68	53.82	19.01×	71.95	84.72	51.98
(b)	LINEAR	1.01×	18.87 (0)	53.72 (0)	23.70 (0)	1.01×	67.46 (4)	81.08 (0)	30.33 (0)
	PARTIAL-1	2.58×	58.44 (5)	78.28 (1)	47.64 (1)	2.58×	72.31 (5)	84.58 (2)	47.89 (1)
(c)	BIAS	1.03×	54.55 (1)	75.68 (1)	47.70 (0)	1.03×	72.89 (3)	81.14 (0)	53.43 (4)
	ADAPTER	1.17×	54.90 (3)	75.19 (1)	38.98 (0)	1.22×	74.19 (4)	82.66 (1)	47.69 (2)
(ours)	VPT-SHALLOW	1.01×	39.96 (1)	69.65 (0)	27.50 (0)	1.01×	67.34 (3)	82.26 (0)	37.55 (0)
	VPT-DEEP	1.04×	36.02 (0)	60.61 (1)	26.57 (0)	1.01×	70.27 (4)	83.04 (0)	42.38 (0)

Gated Prompt Tuning



➔ VPT의 핵심 설계 과제:
Prompt를 어디에, 얼마나 개입시킬 것인가?

prompt의 layer별 개입 정도를 학습:

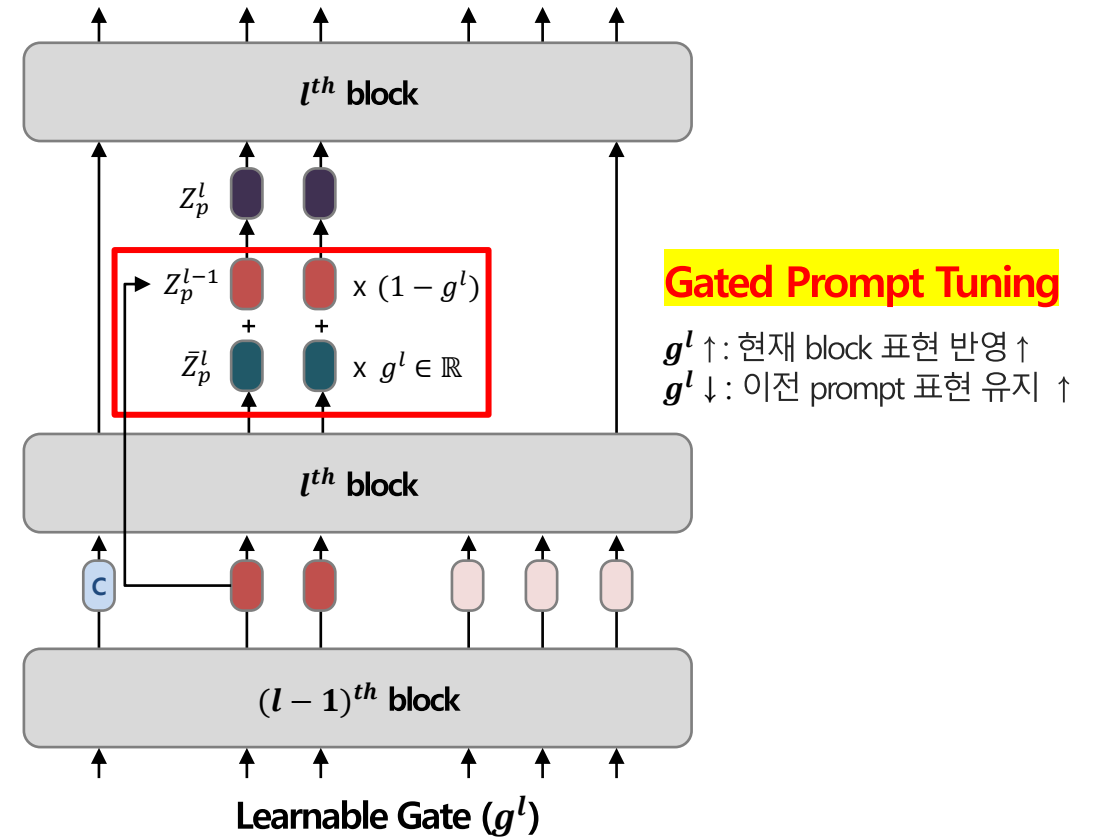
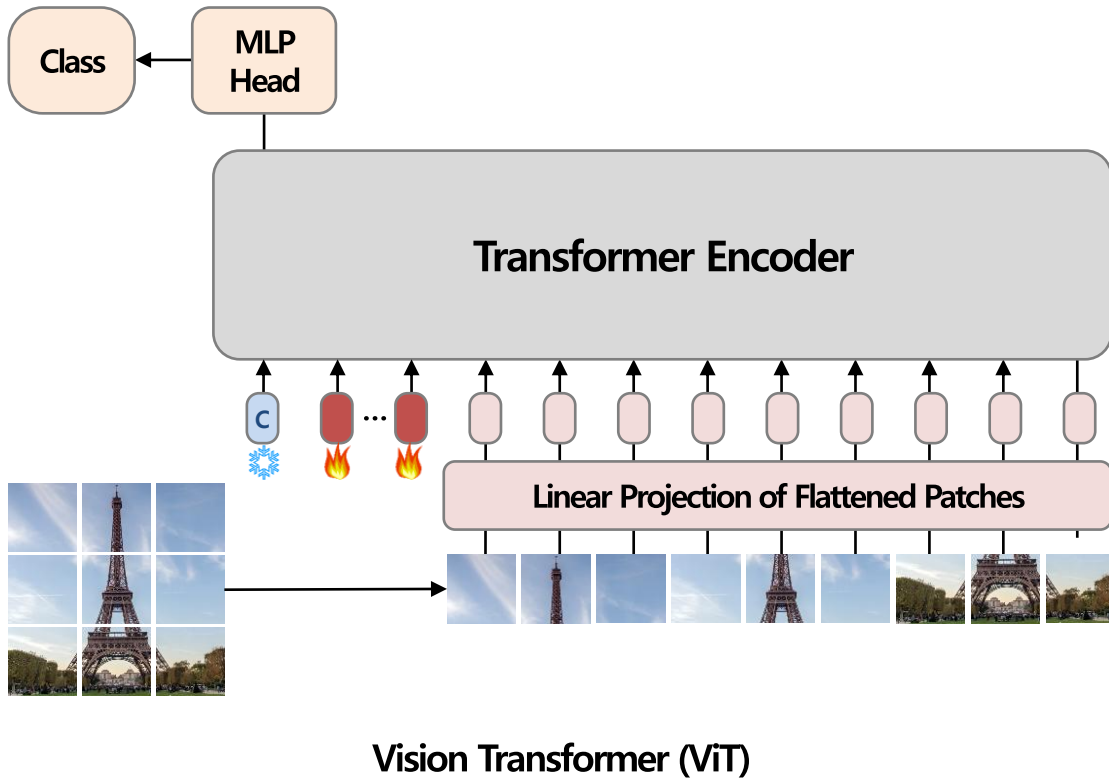


Gated Prompt Tuning

How Gated Prompt Tuning Works

❖ Learnable Gate: Block별 Prompt 개입 정도 조절

- 각 ViT block에 학습 가능한 gate g^l 를 적용하여 prompt representation의 반영 정도를 조절
- VPT-Shallow와 VPT-Deep의 장점을 모두 취함



Gated Prompt Tuning

Experiments

❖ Gated Prompt Tuning Experiment #1: Classification results on FGVC

- 비교 대상: VPT-Shallow / VPT-Deep
- 평가 환경: MAE / MoCo v3로 self-supervised pre-trained ViT-B/16
- 성능: MAE 73.39%, MoCo v3 83.00% → VPT-Shallow 대비 큰 폭 향상, VPT-Deep 대비 우수하거나 대등한 성능

SSL	METHOD	TOTAL PARAMS	CUB	FLOWERS	CARS	DOGS	NABIRDS	AVG
MAE	VPT-SHALLOW	1.02×	42.15	69.15	43.38	77.07	57.43	57.84
	VPT-DEEP	1.02×	68.33	80.05	67.67	78.83	65.22	72.02
	OURS	1.02×	70.56	78.55	71.70	78.9	67.26	73.39
MoCo v3	VPT-SHALLOW	1.02×	79.05	90.47	71.91	81.97	72.92	79.26
	VPT-DEEP	1.02×	82.67	94.41	79.18	83.33	75.99	83.12
	OURS	1.02×	82.86	93.71	79.02	83.37	76.02	83.00

Gated Prompt Tuning

Experiments

❖ Gated Prompt Tuning Experiment #2: Classification results on VTAB-1K

- 비교 대상: VPT-Shallow / VPT-Deep
- 평가 환경: MAE / MoCo v3로 self-supervised pre-trained ViT-B/16
- 성능: MAE 49.22%, MoCo v3 65.80% → 두 환경 모두 VPT-Shallow / VPT-Deep 대비 높은 성능

SSL	METHOD	TOTAL PARAMS	NATURAL (7)	SPECIALIZED (4)	STRUCTURED (8)	AVG
MAE	VPT-SHALLOW	1.01×	39.96 [†]	69.65 [†]	27.50 [†]	40.96 [†]
	VPT-DEEP	1.04×	36.02 [†]	60.61 [†]	26.57 [†]	37.22 [†]
	OURS	1.01×	47.61	76.86	36.80	49.22
MoCo v3	VPT-SHALLOW	1.01×	67.34 [†]	82.26 [†]	37.55 [†]	57.94 [†]
	VPT-DEEP	1.01×	70.27 [†]	83.04 [†]	42.38 [†]	61.22 [†]
	OURS	1.01×	74.84	83.38	49.10	65.80

Gated Prompt Tuning

Experiments

❖ Gated Prompt Tuning Experiment #3: Semantic Segmentation on ADE20K

- **평가 환경:** SETR-PUP / MAE, MoCo v3로 self-supervised pre-trained ViT-B/16
- **평가 지표:** mIOU ↑ (SS: Single-Scale / MS: Multi-Scale inference)
- **성능:** MAE 38.44 / 39.81, MoCo v3 36.81 / 38.55 → VPT-Deep보다 적은 prompt로 더 높은 성능

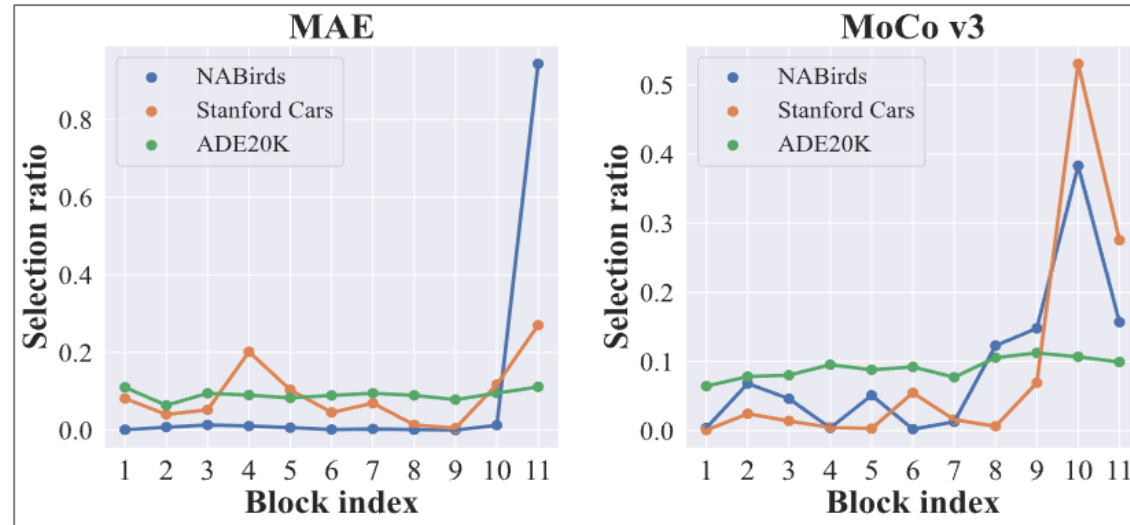
SSL	METHOD	# OF PTs	mIOU (SS)	mIOU (MS)
MAE	VPT-SHALLOW	100	34.20	35.23
	VPT-DEEP	10($\times 12$)	37.76	38.80
	OURS	100	38.44	39.81
MoCo v3	VPT-SHALLOW	100	34.55	36.18
	VPT-DEEP	10($\times 12$)	35.50	37.15
	OURS	100	36.81	38.55

Gated Prompt Tuning

Experiments

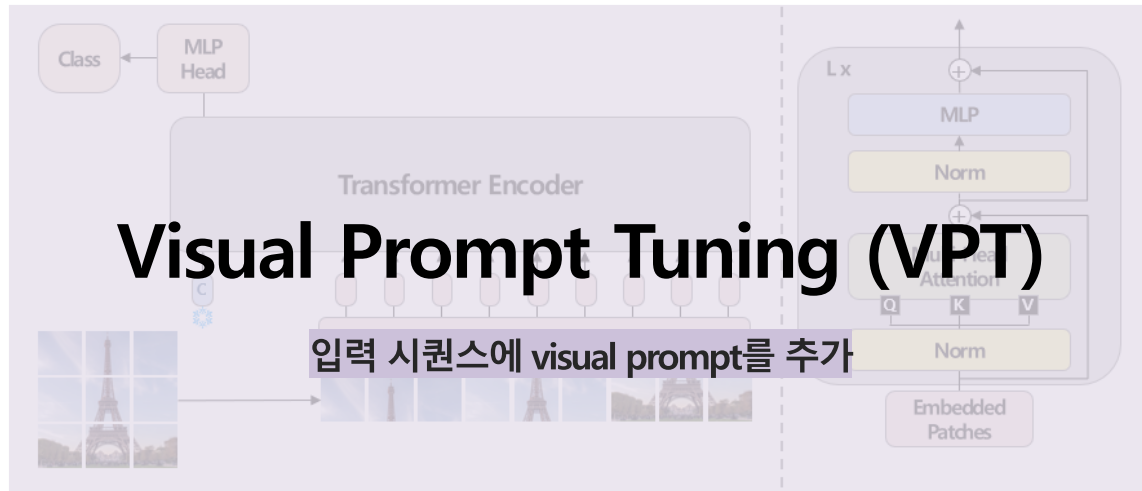
❖ Gated Prompt Tuning Experiment #4: Analysis on Learned Gates

- **Selection ratio:** 각 block이 최종 prompt representation에 미치는 영향
- **Task별 차이:** MAE에서 NABirds는 11번째 block, Stanford Cars는 4·11번째 block에 집중
- **Backbone별 차이:** 동일 task에서도 MAE와 MoCo v3가 서로 다른 block을 선택
- **ADE20K:** 여러 block을 비교적 고르게 활용 → segmentation에는 multi-level feature가 중요



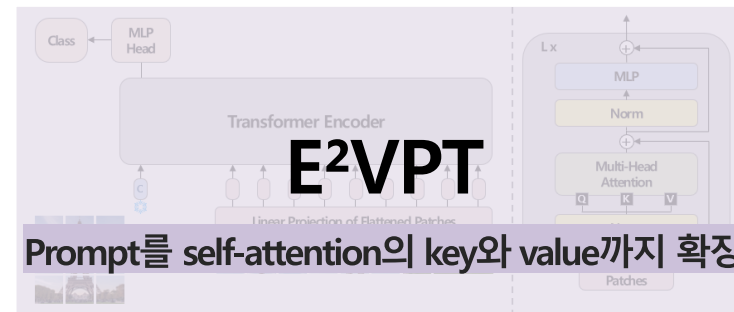
Task와 SSL backbone에 따라 필요한 block을 선택적으로 반영

Conclusion

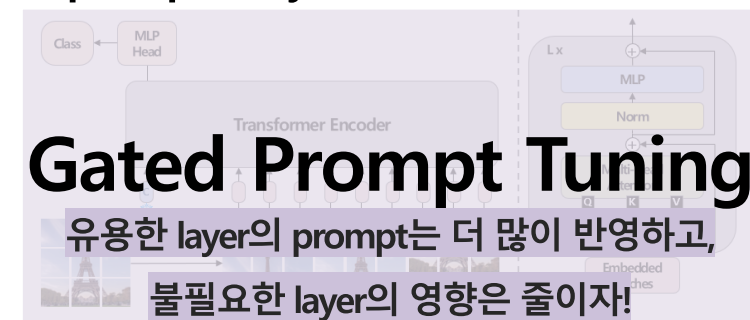


➔ VPT의 핵심 설계 과제:
Prompt를 어디에, 얼마나 개입시킬 것인가?

self-attention 구조 내부로도 확장:

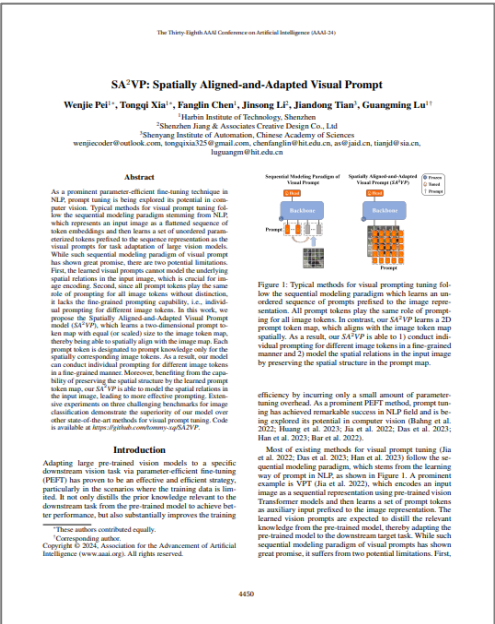


prompt의 layer별 개입 정도를 학습:



Conclusion

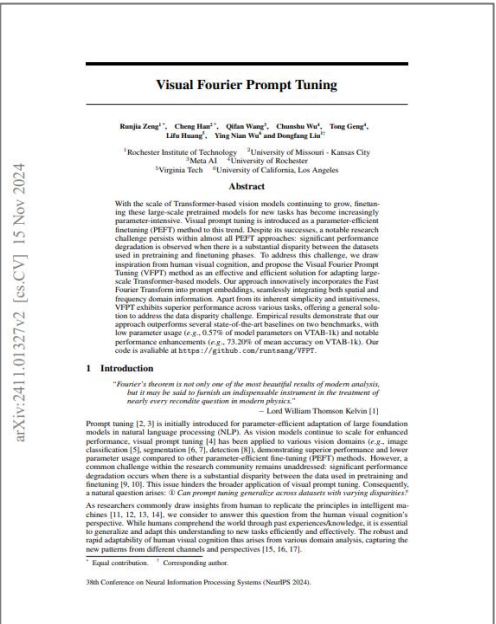
VPT 계열 연구 방향



SA²VP: Spatially Aligned-and-Adapted Visual Prompt

(Pei et al., AAAI 2024)

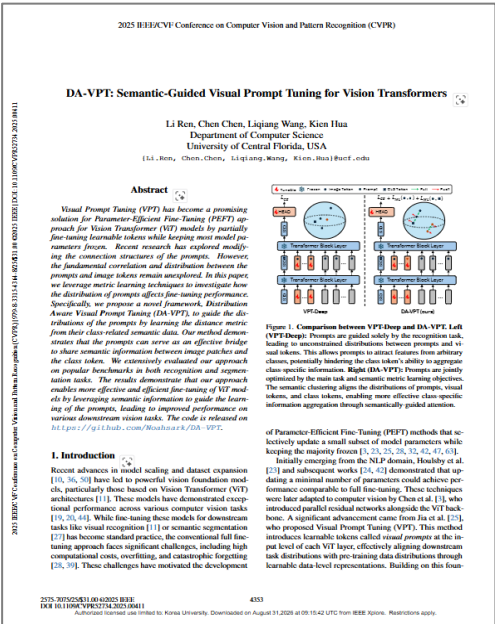
→ prompt의 공간 구조를 2차원으로 정교화



Visual Fourier Prompt Tuning

(Zeng et al., NeurIPS 2024)

→ prompt의 표현 영역을 주파수 도메인으로 확장



DA-VPT: Semantic-Guided Visual Prompt Tuning for Vision Transformers

(Ren et al., CVPR 2025)

→ prompt 학습에 의미적 제약을 부여

Thank You

References

- [1] Jia, M., Tang, L., Chen, B. C., Cardie, C., Belongie, S., Hariharan, B., & Lim, S. N. (2022). Visual prompt tuning. In European Conference on Computer Vision (pp. 709-727).
- [2] Han, C., Wang, Q., Cui, Y., Cao, Z., Wang, W., Qi, S., & Liu, D. (2023). E²VPT: An effective and efficient approach for visual prompt tuning. In Proceedings of the IEEE/CVF International Conference on Computer Vision (pp. 17491-17502).
- [3] Yoo, S., Kim, E., Jung, D., Lee, J., & Yoon, S. (2023). Improving visual prompt tuning for self-supervised vision transformers. In International Conference on Machine Learning (pp. 39848-39861).