

**AI로 가능성을 확장하는 Engineer,
서용득입니다.**

이력사항



서용득 Seo Yongdeuk

- 복잡한 문제를 분석하고 해결하는 과정에서 성취감을 느낍니다.
- AI 기술을 실제 서비스와 산업 현장에 적용할 수 있습니다.

M. 010-3046-8377
E. sod7050@gmail.com
W. github.com/yongchooon

Graduation	2017.03 - 2024.02	국립부경대학교 공학사 시스템경영안전공학부 기술데이터공학전공
	2024.03 - 현재	국립부경대학교 석사과정 산업및데이터공학과 산업데이터공학융합전공
Experience	2022.04 - 현재	국립부경대학교 산업AI연구실 연구원
	2023.08 - 2023.12	주식회사 토모큐브 AI팀 Intern
Project		
2023.08 - 2023.12	Segment Anything 기반 Auto-Annotator 개발 및 사내 배포 Segment Anything, OpenCV, Docker, FastAPI, MySQL, SQLAlchemy	
2024.04 - 2024.08	DALDA: LLM, Diffusion 활용 데이터 증강 프레임워크 개발 및 논문 게재 [ECCV Workshops 2저자 게재] Stable Diffusion, GPT-4, IP-Adapter, CLIP	
2024.11 - 2024.12	KIMCHI: 한국어 문화, 역사를 반영한 한국적 Diffusion 모델 개발 및 논문 게재 [한국CDE학회 동계학술대회 게재] Stable Diffusion, LLaVA, LLaMA, LoRA	
2025.03 - 2025.08	STELLAR: 다국어 및 실제 환경 Scene Text Editing 모델 개발 [AAAI Main Conference 심사 중] Stable Diffusion, ControlNet, PaddleOCR	
2025.09 – present	T-RADAR: 유사 상표 이미지 검색 및 AI Agent 기반 상표 등록 시뮬레이션 시스템 개발 [현재 진행 중] LangGraph, DINO, MetaCLIP2, BM25, PGVector, PostgreSQL, FastAPI, Docker	
Awards		
2024.08	제 1회 해커톤 경진대회 우수상 [Upstage Solar API 활용 서비스 개발] 국립부경대학교 성균관대학교 ICAN사업단	
2025.02	한국CDE학회 동계학술대회 우수발표상	

Patent	
2023	[출원, 공개] 인공지능 학습을 위한 화재이미지 생성 시스템 및 방법 공개번호 1020250021779
2024	[출원] 대규모 언어 모델의 프롬프트를 적용한 환산 모델을 통한 데이터 증강 방법 및 장치 출원번호 1020240137347

Skill	
Main	: 3회 이상 프로젝트 경험 보유
Sub	: 프로젝트 경험 보유
- 언어 : Python JavaScript	
- ML Library : PyTorch OpenCV	
OpenAI Diffusers PaddleOCR	
- 웹 프레임워크 / DB : FastAPI	
MySQL MongoDB PostgreSQL	
- DevOps : Git Docker Slack	

[(주)토모큐브 인턴십 중 개발] 세포 이미지 데이터셋 구축을 위한 Web 기반 고정밀·고효율 Auto-Annotator
백엔드 서버 구축, Segment Anything Model (SAM) 활용 기능 및 Fine-tuning 기능 구현 후 사내 배포

| 기간 2023.08 - 2023.12

| 기여도 70%

Problem

- 세포 이미지 데이터셋 구축의 어려움
 - 대량 수작업 → 시간/비용 소모

Solution

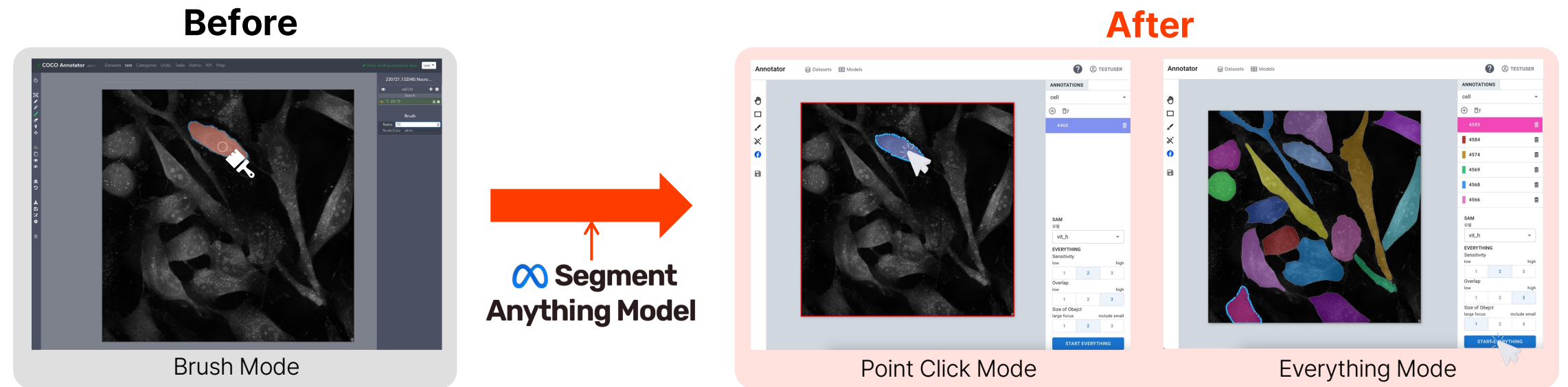
- SAM 기반 Auto Annotation 기능 구현
 - Point Click Mode
 - Everything Mode
 - ⇒ annotation 작업 속도 큰 폭으로 단축
- 사용자 어노테이션 데이터를 활용한 Fine-tuning 기능 추가
 - SAM prediction을 통한 준가공 데이터 수집 후 Editing
 - 가공 데이터로 모델을 LoRA Fine-tuning

My Role

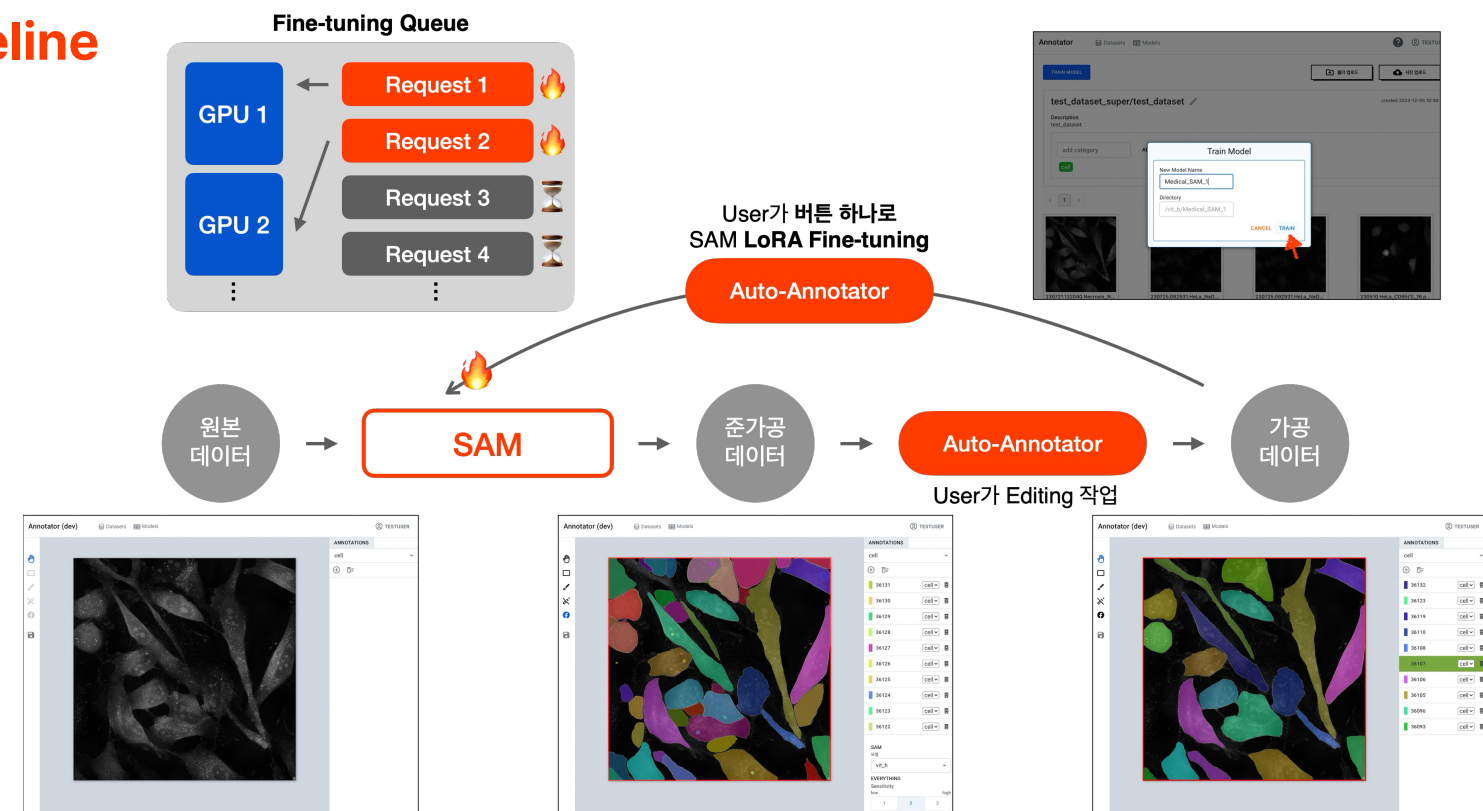
- Backend Server 구축 및 배포
 - FastAPI, Docker, MySQL 활용
- Fine-tuning 기능 및 학습 Queue 구현
 - Celery 활용: 한정적 GPU 자원을 효율적으로 사용

Results

- 수작업 대비 annotation 효율성 및 정확도 ↑ (Mean Dice Coefficient 약 2.7%p▲; 0.9107 → 0.9389)
- 한정된 GPU 환경에서 안정적 Queue 설계 ⇒ 자원 활용 최적화



Fine-tuning Pipeline



[석사과정 중 연구 및 논문 게재] Few-shot 학습 환경에서 합성 데이터의 다양성과 일관성을 균형있게 제어하는 데이터 증강 프레임워크

| 기간 2024.04 - 2024.08

| 기여도 50%

Problem

- Few-shot 학습 환경에서 합성 데이터로 증강 시
 - 다양성 증가, But target 분포와 멀어질 위험
 - ⇒ 성능 저하 발생

Solution

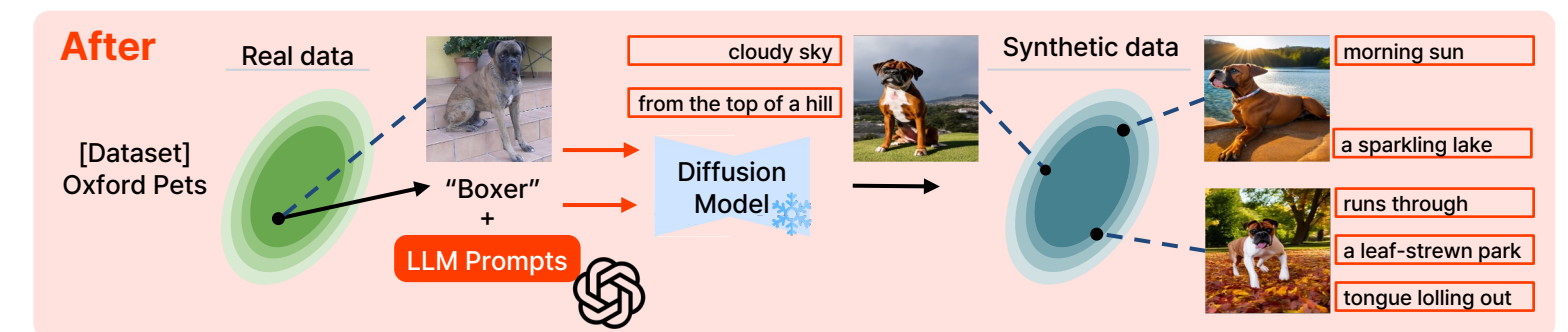
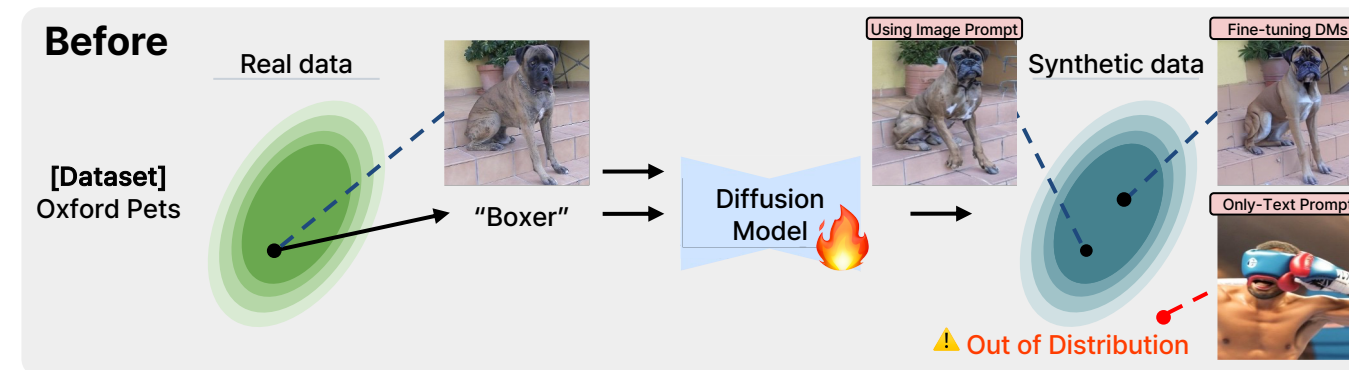
- LLM (GPT-4) 활용: **class의 의미적 특징을 확장**한 text prompt 생성
- Stable Diffusion + IP-Adapter 기반 이미지 생성
 - 클래스 일관성 유지 ⇒ **target 분포 내 생성 유도**
- Adaptive Guidance Scaling 설계
 - CLIPScore ↓ ⇒ weight λ ↑ ⇒ diversity ↓
 - CLIPScore ↑ ⇒ weight λ ↓ ⇒ diversity ↑
 - 동적으로 **학습 데이터 다양성을 균형있게 제어** ⇒ 성능 ↑

My Role

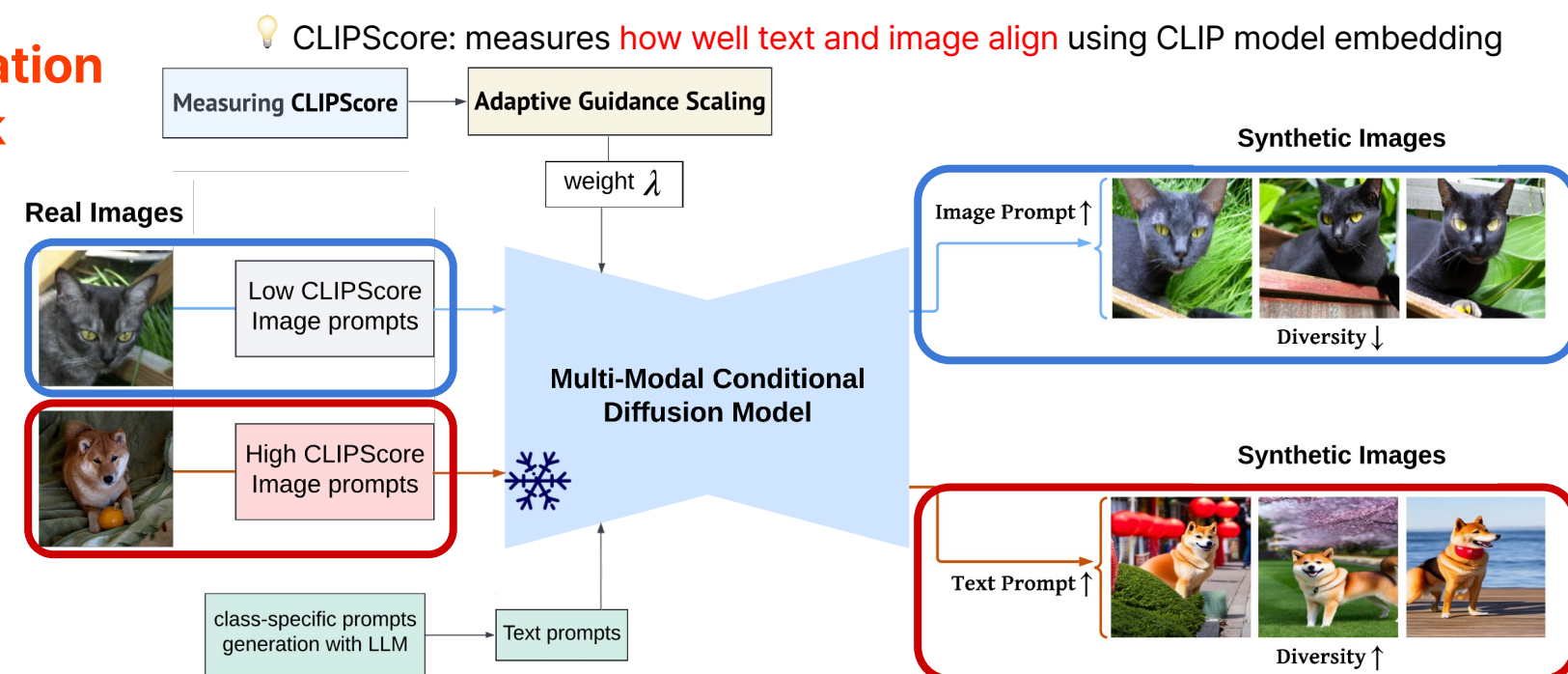
- 합성 데이터 증강 학습 파이프라인 설계 및 구현
- 실험, 결과 분석 및 논문 작성

Results

- 3개 벤치마크 데이터셋에서 1-shot classification accuracy가 SOTA 모델 대비 평균 9.07%p 상승
- ECCV 2024 Workshops 논문 2저자 게재 및 Oral Session



Data Augmentation Framework



[석사과정 중 연구 및 논문 게재] 한국어 텍스트를 이해하고 한국적 이미지를 자연스럽게 생성하는 Diffusion 모델 개발

Problem

- 기존 T2I 생성 모델의 한계
 - 영어 중심 학습 $\Rightarrow$ 한국어 기반 생성 성능 저하
 - 한국 고유의 문화적 맥락 반영 불가

Solution

- 한국어 프롬프트 생성 후 추가 학습
 - LLaVA + LLaMA 기반
 - 한국적 이미지에 대한 묘사
- 한국어 - 한국적 이미지 데이터셋 활용 학습
 - LoRA Tuning : Ko-CLIP Text Encoder, Stable Diffusion

My Role

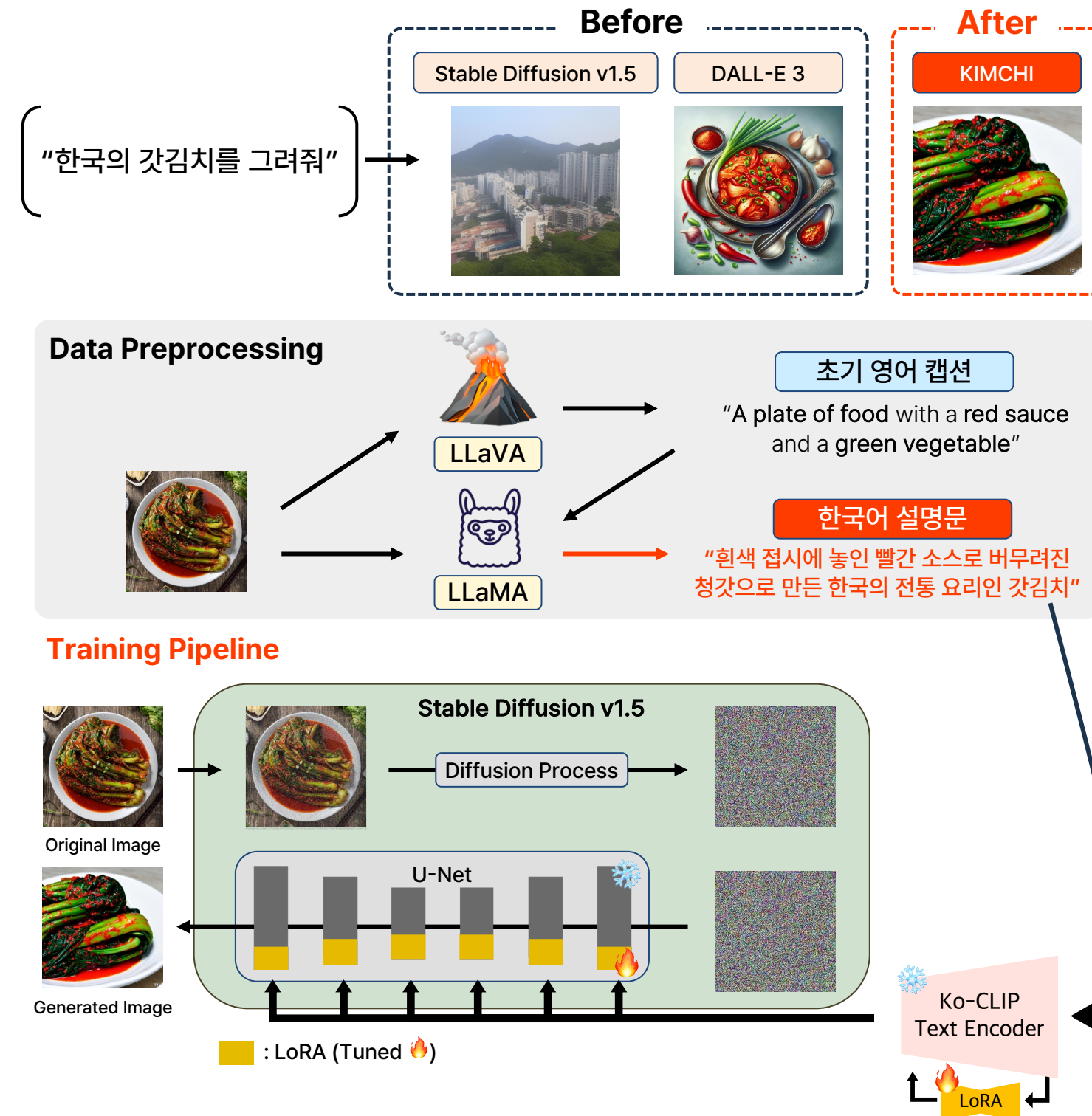
- 데이터셋 정제·구축 - 학습 파이프라인 설계 및 구현
- 실험·평가 및 결과 분석, 논문 작성

Results

- Human Evaluation 결과, **최고 성능 달성**
- 최신 생성 모델과 비교하여, 한국적 음식, 물품, 문화재, 풍경 등 **다양한 문화적 요소를 반영한 이미지 생성 가능**

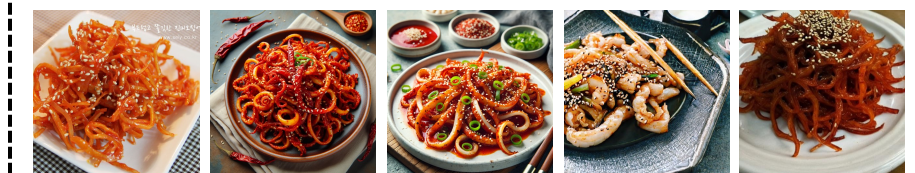
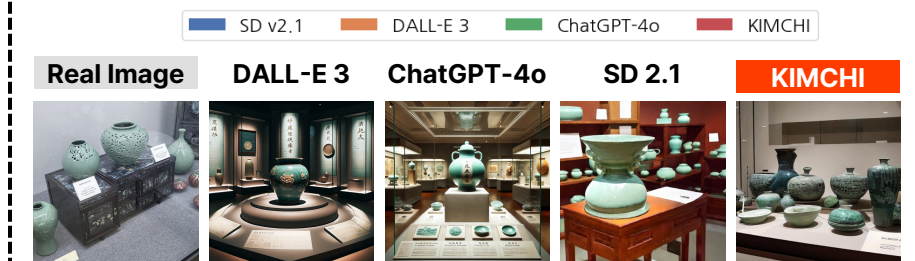
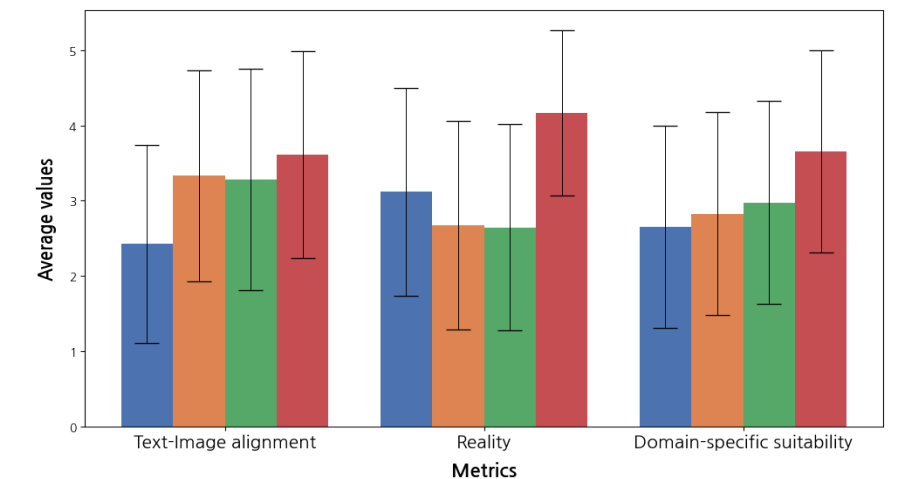
1. Text-image alignment 약 6%▲, Reality 약 21%▲, Domain-specific suitability 약 14%▲

2. 한국CDE학회 2025 동계학술대회 논문 발표 및 우수발표상 수상



| 기간 2024.11 - 2024.12
| 기여도 50%

Results



[석사과정 중 연구 및 논문 심사중] 저자원 언어와 real-world 이미지 중심의 Scene Text Editing(STE) 모델 개발

ControlNet 기반 파이프라인과 2단계 학습 전략 제안, 학습 데이터셋 직접 구축, Text Appearance Similarity 평가를 위한 새로운 지표 TAS 제안

| 기간 2025.03 - 2025.08

| 기여도 100%

Problem

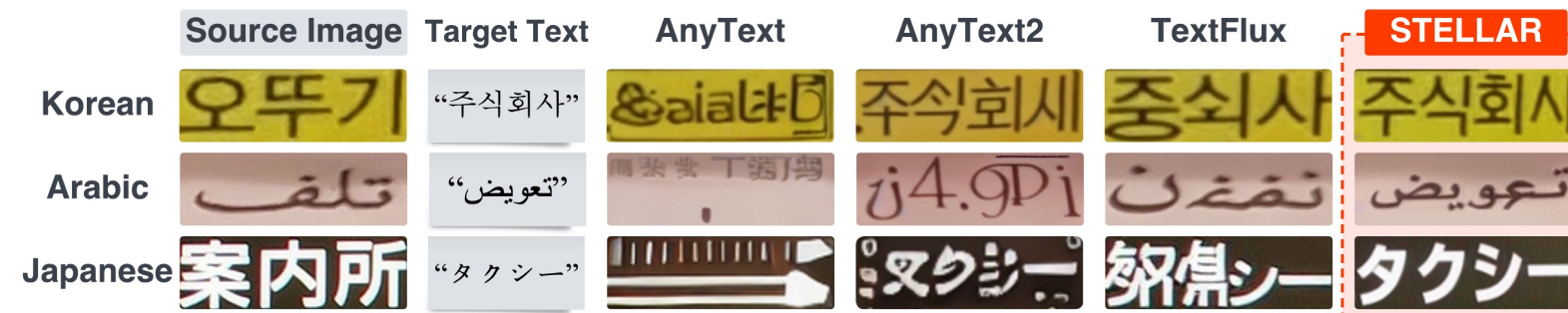
- 기존 STE 모델은 대부분 영어, 중국어 위주 학습
 - 한국어, 아랍어, 일본어 등의 다국어 모델 부재
- 합성 데이터 의존 학습 다수
 - Real-world 데이터에 대한 생성 성능 저하
- 모델 평가 지표의 한계 존재

Solution

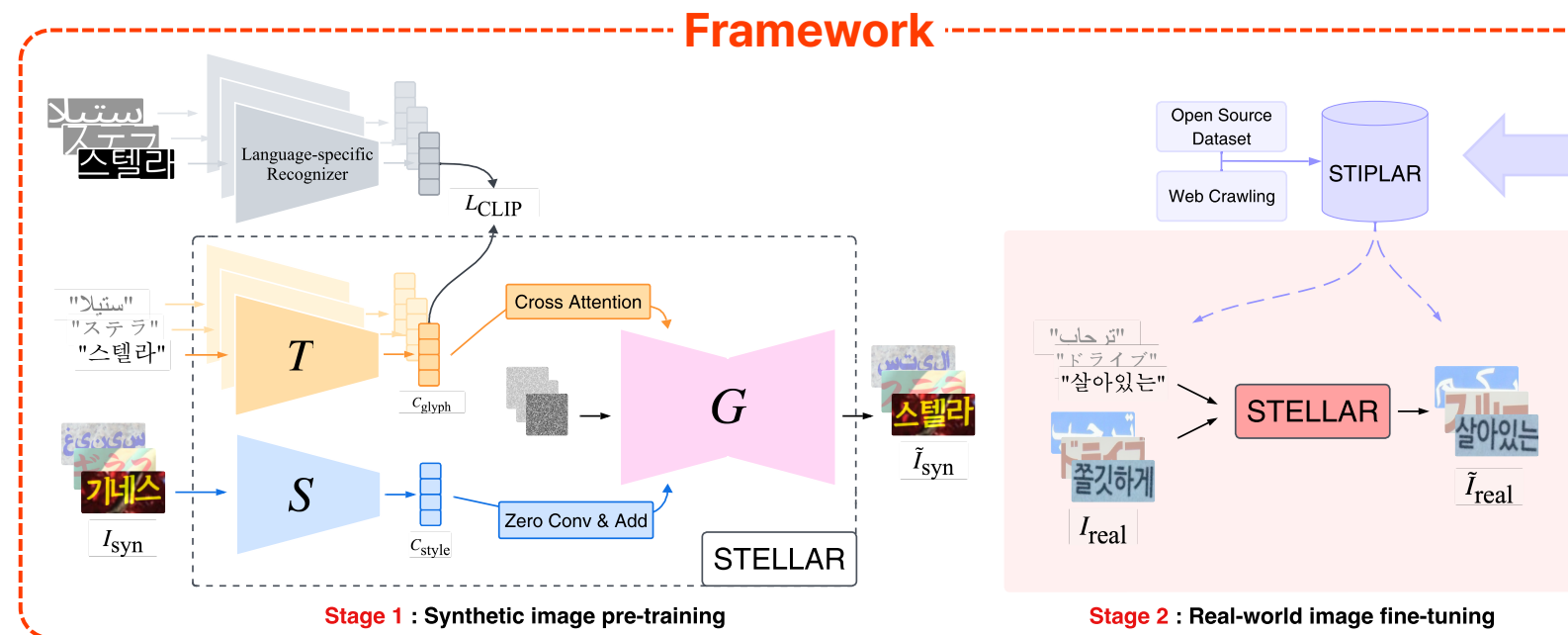
- 다국어 OCR Recognizer 기반으로 전용 Text Encoder 사전학습
- Text Style Encoder 학습 : 글꼴, 색상, 배경을 분리
- ControlNet 2단계 학습 전략 제안(합성 → 실제)
 - ⇒ 효과적인 domain adaptation 달성
- 다국어 실제 이미지 쌍 데이터셋 STIPLAR 구축 및 공개
 - 오픈 소스 데이터셋 + 웹 크롤링 활용 데이터 수집
 - PaddleOCR, LLM(GPT-4) 기반 필터링
 - ⇒ 2단계 학습에 사용하여 실제 이미지 편집 성능 향상
- 텍스트 스타일 보존도 평가 지표 TAS 제안
 - 글꼴, 폰트, 배경을 분리하여 개별적 유사도 평가
 - ⇒ Metric의 신뢰성, 해석력 증가

Results

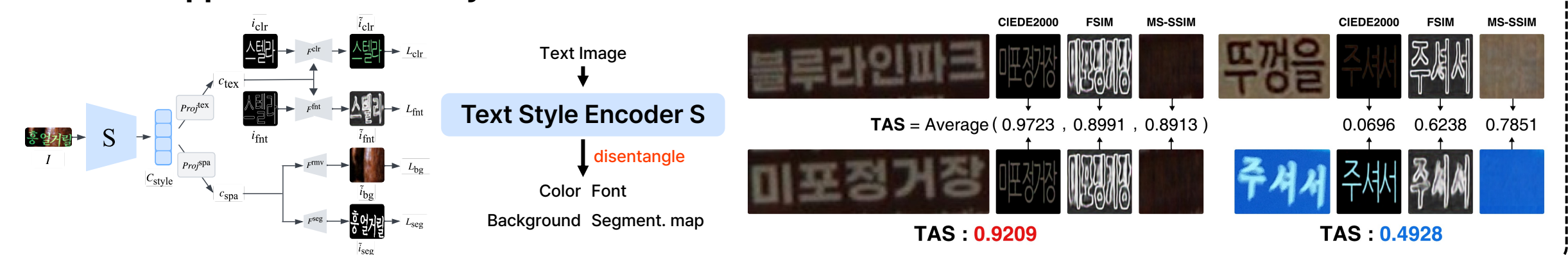
- Flux 기반 SOTA 모델에 비해 3개 언어에 대해 평균적으로 TAS 약 1.8%p▲, Recognition Accuracy 약 40.4%p▲
- AAAI 2026 Main Conference 제출 후 심사 중



Dataset Examples



TAS: Text Appearance Similarity



[석사과정 중 / 개발 진행 중] 상표 이미지와 상표명 기반 유사 상표에 대한 다중 소스 검색 시스템
AI Agent 기반 상표 등록 과정에서의 심사관과 출원인 시뮬레이션 및 보고서 제공

| 기간 2025.09 – present
| 기여도 70%

Problem

- 기존에는 상표 등록 전 **번거로운 반복 작업 존재**
 - 주요부 비엔나 코드 활용 상표 검색 수작업
- 상표 등록 및 거절은 **복합적이고 주관적인** 판단에 근거함
 - 거절 사유 및 등록 가능성을 **예측하기 어려움**

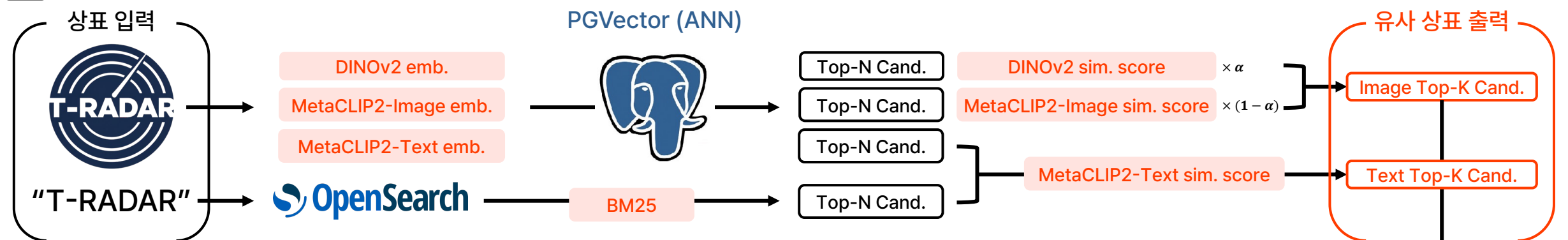
Solution

- 이미지 / 텍스트에 대한 **ANN** 기반 유사 상표 검색
 - 상표 데이터셋 수집·정제 후 **PGVector** 기반 검색 DB 구축
 - 상표 공보 / 의견제출통지서 / 거절결정서
 - 다중 소스 검색 및 통합 랭킹**: DINOv2, MetaCLIP2 embedding 유사도 및 BM25 활용
- AI Agent (LangGraph)** 기반 상표 등록 시뮬레이션
 - 심사관과 출원인 Agent의 거절 및 반박 과정**을 리포터 Agent가 정리 후 출력
 - 평가자 Agent가 상표 침해 위험도 및 등록 가능성 score 제공

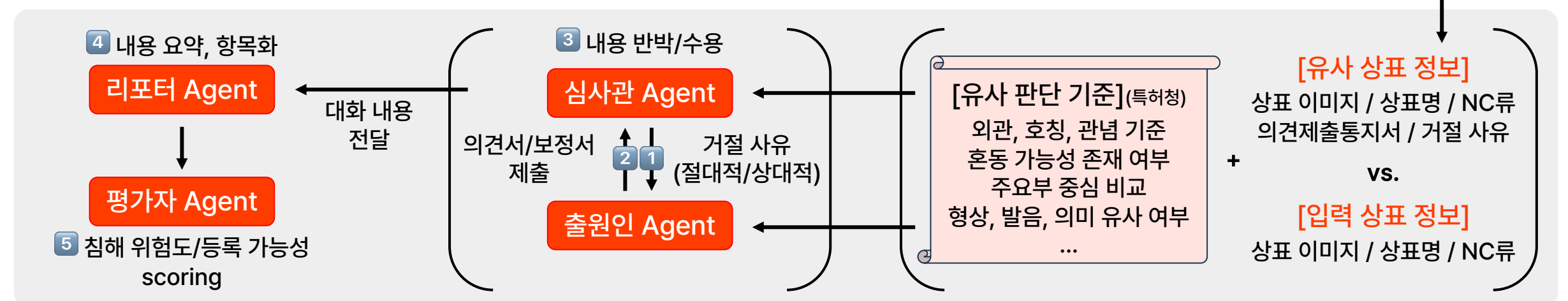
Objective

- 의견제출통지서에 기재된 선등록상표를 유사 상표 정답으로 한 테스트셋에서 $\text{Recall@20} \geq 0.50$
- 특허 등록 및 실제 서비스화

1 유사 상표 이미지 검색



2 상표 등록 시뮬레이션



3 사용자 출력 예시

	유사 상표	상태	NC류	이미지 유사도	텍스트 유사도	시뮬레이션 결과 요약	침해 위험도	등록 가능성
Ex 1)	T "TRA-DER"	등록	비인접	0.4	0.9	거절 사유 예측 반박 의견 보정 가능한 내용	50%	25%
Ex 2)	R "TRADER"	거절	인접	0.8	0.9		80%	

[학사과정 중 연구 및 특허 등록, 공개] 실제 수집이 어려운 화재 이미지 합성을 통해 학습 데이터 부족 문제를 해결
LoRA, ControlNet 학습을 기반으로 화재 이미지 생성을 위한 Diffusion Model 개발

| 기간 2023.03 - 2023.05
| 기여도 50%

Problem

- CCTV 화재 감지 모델 개발에 필요한 학습 데이터가 부족함
- 실제 화재 이미지는 안전, 비용, 환경적 제약으로 수집이 어려움
- 기존에 수집된 화재 이미지 데이터들은 인위적 수집
 - ⇒ 중복/편향 다수 존재

Solution

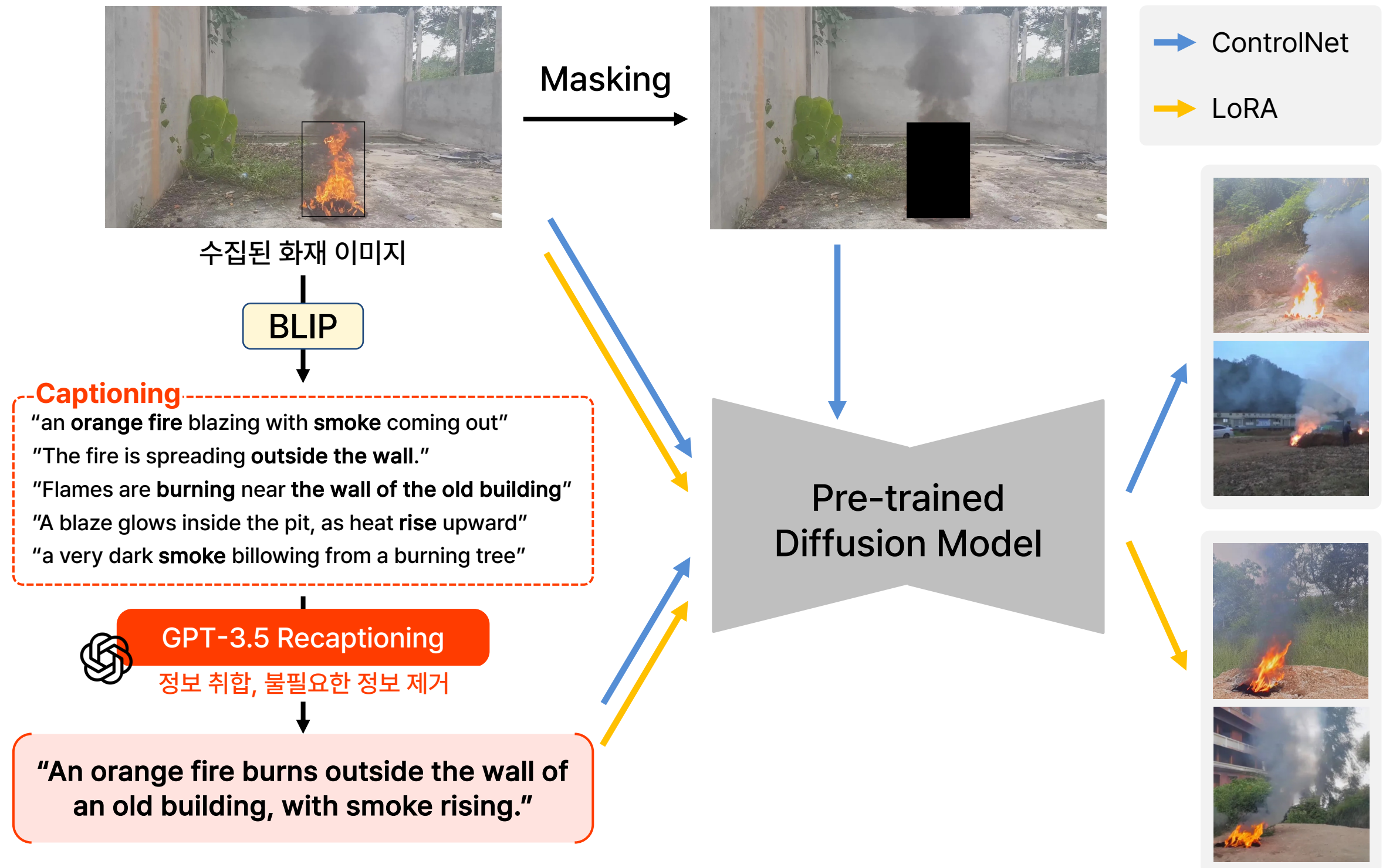
- Diffusion Model 기반 사실적인 합성 화재 이미지 생성 시스템 개발
- Dataset : AIHub에서 CCTV 화재 이미지 수집
- VLM, LLM 활용 이미지 Captioning & Preprocessing
 - BLIP 활용 5개 캡션 추출 → GPT 3.5가 1개 프롬프트로 요약
- Solution 1 : 화재 영역 Masking 후 ControlNet 학습
- Solution 2 : Diffusion Model을 LoRA 학습
- ⇒ 다양한 배경, 불꽃 형태, 크기를 반영한 화재 합성 데이터 생성 가능

My Role

- 화재 이미지 데이터 수집 및 전처리 (Masking, Captioning)
- LoRA 기반 Diffusion Model 학습 파이프라인 설계 및 구현

Results

특허 출원 및 공개 [인공지능 학습을 위한 화재이미지 생성 시스템 및 방법] (공개번호 1020250021779)



[석사과정 중 해커톤 참여 및 우수상 수상] STT 기반 운동 내역 자동 기록 및 추천 앱의 백엔드 개발
Whisper API와 Upstage Solar API 기반 [음성 → 텍스트 → 데이터 파싱 → 저장 → 추천] 기능 구현

| 기간 2024.08 - 2024.08
| 기여도 70%

Problem

- 운동 내역을 수작업으로 기록 시 사용자 불편 존재
 - 기능 불편 : 숫자를 여러 번 입력해야 함
 - 오랜 시간 소요
 - 단순히 기록에 그치고, 부가적인 정보는 얻을 수 없음

Solution

- Whisper API 활용 STT : 음성으로 운동 내역을 녹음 → 텍스트 추출
- Upstage Solar API (LLM) 활용
 - 텍스트 → 데이터 파싱 (운동명, 무게, 단위, 횟수, 세트 수 등)
 - 지난 운동 기록을 기반으로 다음 일자의 운동 내역 추천
- FastAPI, MySQL 기반 백엔드 서버 구현
- Flutter 기반 프론트엔드 서버 구현

My Role

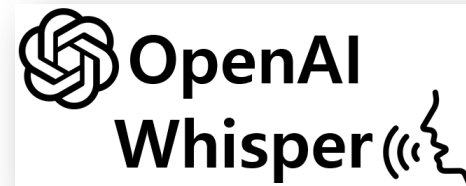
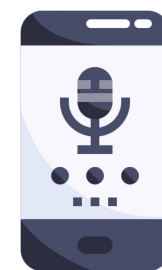
- FastAPI, MySQL 기반 백엔드 서버 구현
- Whisper API, Upstage Solar API 기반 기능 구현
- 프롬프트 엔지니어링 : 데이터 파싱 및 운동 내역 추천

Results

제 1회 LLM 활용 서비스 개발 해커톤 경진대회에서 우수상 수상

주최:국립부경대학교 성균관대학교 ICAN 사업단, 후원: Upstage

Pipeline: 운동 내역 자동 기록



"오늘 8월 13일 화요일.. 등 운동 랫풀다운 30키로 15회 한 세트.. 35키로 15회 한 세트.. 아 그리고 뭐였더라.. 아! 40키로 15회 한 세트하고 40키로 12회 한 세트..."

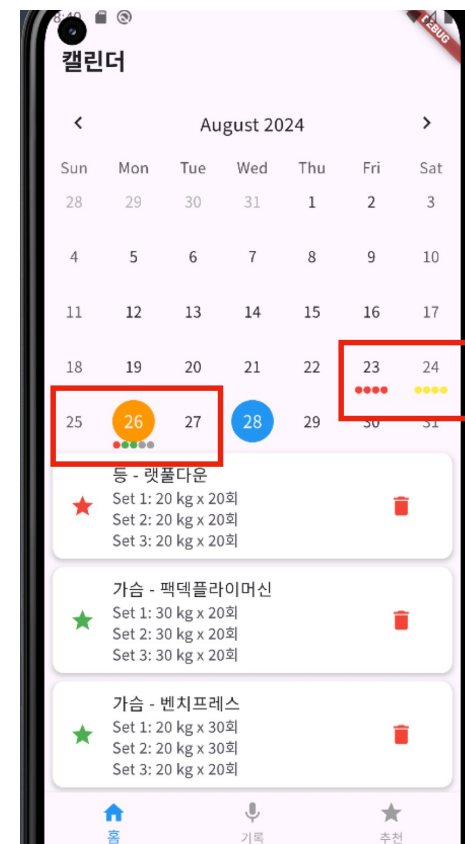


```
"등": {
  "랫풀다운": [
    [30, "kg", 15], [35, "kg", 15],
    [40, "kg", 15], [40, "kg", 12]
  ]
}
```

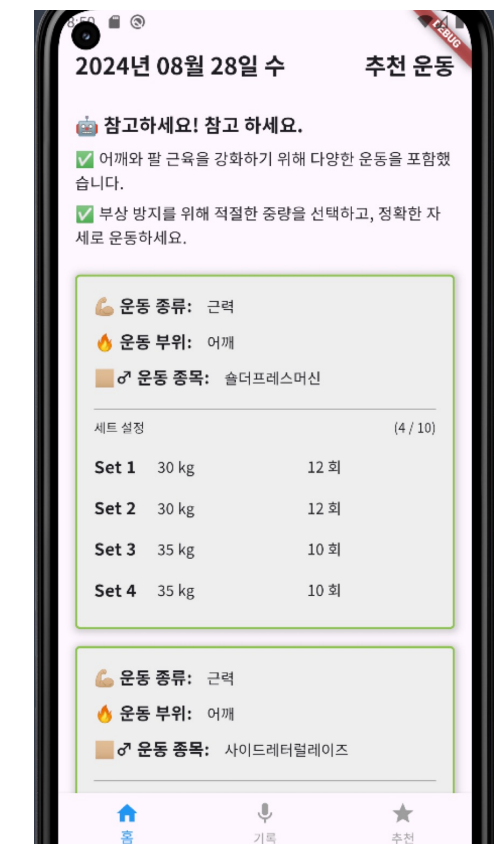
2024년 08월 13일 화

운동 종류:	근력
운동 부위:	등
운동 종목:	랫풀다운
세트 설정	(4 / 10)
Set 1	30 kg 15 회
Set 2	35 kg 15 회
Set 3	40 kg 15 회
Set 4	40 kg 12 회

운동 추천 기능



- 지난 5일 간의 운동 내용을 기반으로 오늘의 운동 내용을 추천해야 함
- 휴식 유무를 고려해야 함
- 중량 및 횟수에 대한 점진적 과부하를 반영하여 추천해야 함
- 추천에 대한 근거 및 유의사항을 제공해야 함
- Output 형식은 python list 형식이어야 함



THANK YOU

끝까지 검토해 주셔서 감사합니다.