

전장 상황 인식을 위한 Arma 3 기반의 Scene Graph 생성 및 VQA 데이터 자동 생성에 관한 연구

김성현*¹, 박종현*², 김예찬*³, 권철희**¹, 조규태**², 진정훈**³, 김지아**⁴, 전문구***

Generating Training Data using Arma 3 for Battlefield Situation Recognition

Sunghoon Kim*¹, Jonghyun Park*², Yechan Kim*³, Cheolhee Kwon**¹, Kyutae Cho**²,
Junghun Jin**³, Jia Kim**⁴, and Moongu Jeon***

이 연구는 LIG NEX1 산학협력과제 지원으로 연구되었음

요 약

본 연구에서 우리는 전장 상황에 특화된 모델 훈련 데이터를 효율적으로 생성할 수 있는 프레임워크를 제안한다. 프레임워크는 먼저 Arma3 시뮬레이터를 활용하여 전쟁 시뮬레이션을 수행한 뒤, 이와 관련한 영상 및 메타 정보를 자동으로 수집한다. 다음으로 수집한 메타 정보를 바탕으로 장면 그래프 (Scene Graph)를 생성하고, 사전에 정의한 질문-답변 템플릿을 바탕으로 질문-응답 텍스트 레이블을 자동 생성하는 알고리즘을 개발하였다. 제안된 프레임워크로 생성한 데이터를 통해 대규모 비전-언어 모델인 LLaVA의 파인튜닝 가능성을 검증하였으며, 이를 통해 전장 상황 데이터 부족 문제를 해결하고 특화된 모델 개발을 촉진할 수 있음을 확인하였다.

Abstract

In this study, we propose a framework for effectively generating training data for AI models tailored to battlefield scenarios. Using the Arma 3 simulator, we conducted war simulations and developed an algorithm to automatically collect video data along with associated metadata. Based on the collected metadata, we generated scene graphs and designed an algorithm to automatically create question-answer text labels using predefined templates. Leveraging the synthetic training data produced by the proposed framework, we fine-tuned LLaVA, a large-scale vision-language model. Our results demonstrate that the proposed approach effectively addresses the issue of data scarcity and facilitates the development of models specialized for battlefield environments.

Keywords

visual question answering, scene graph, battlefield, large-scale vision-language model

* GIST 전기전자컴퓨터공학과

- ORCID¹: <https://orcid.org/0009-0001-5221-6036>

- ORCID²: <https://orcid.org/0009-0005-5404-0707>

- ORCID³: <https://orcid.org/0000-0002-2438-3590>

** LIG NEX1 AI 연구소 AI연구개발팀

- ORCID¹: <http://orcid.org/0000-0002-5811-1622>

- ORCID²: <http://orcid.org/0009-0008-4037-1486>

- ORCID³: <http://orcid.org/0009-0003-1737-9213>

- ORCID⁴: <http://orcid.org/0009-0003-9050-7023>

*** GIST 전기전자컴퓨터공학과 교수(교신저자)

- ORCID: <https://orcid.org/0000-0002-2775-7789>

· Received: Feb. 26, 2025, Revised: Mar. 27, 2025, Accepted: Mar. 31, 2025

· Corresponding Author: Moongu Jeon

Dept. of Electrical Engineering and Computer Science, GIST, Korea

Tel.: +82-62-715-2406, Email: mgjeon@gist.ac.kr

I. 서 론

현대 전쟁에서 AI의 활용 가능성은 날이 갈수록 증가하고 있다[1]. 이러한 변화는 드론 및 첨단 네트워크 통신 기술의 도입으로 더욱 가속화되고 있으며, 이는 군사 작전의 정밀성과 효율성을 향상시키는 동시에 인명 손실을 최소화하는 데 크게 기여할 수 있다. 특히 AI를 활용하여 복잡하고 역동적인 전장 환경에서 군사 물체를 탐지 및 추적하거나 정확한 상황 인식(Situational awareness)을 통해 신속하면서도 정확한 결정을 내리는 것이 임무의 성공을 좌우하는 핵심 요인으로 자리 잡고 있다.

현재도 탐지(Detection), 추적(Tracking), 상황 인식을 위한 AI 모델들이 존재하지만, 전장과 같은 특수한 환경에서 활용될 때 성능이 저하되는 문제가 존재한다. 이는 주로 모델이 훈련된 데이터셋과 군사적 응용 환경 간의 괴리에서 비롯된다. 예를 들어, 기존의 일반적인 모델 학습 데이터셋은 주로 일상적인 장면이나 객체를 대상으로 구성되어 있으며, 군사적 요소(예: 병력 이동, 전투 상황)를 포함하지 않는다. 이러한 데이터 부족 문제는 다양한 모델들의 군사적 활용 가능성을 제한하는 주요 원인으로 작용한다.

본 연구에서는 군사 시뮬레이션 게임 Arma 3[2]를 기반으로 한 데이터 생성 프레임워크를 설계하였다. Arma 3를 채택한 이유는 Arma 3는 현실성 높은 군사 기물들을 지원하여 현실적인 가상 데이터 생성이 가능하며 스크립트 언어 코딩을 통해 데이터 생성을 자동화할 수 있으며 동시에 필요한 데이터셋 생성에 필요한 메타 데이터도 추출 가능하기 때문이다.

이 프레임워크는 군사적 시뮬레이션 환경에서 현실적인 군사 데이터를 자동으로 생성한다. 드론 촬영 영상, 객체 메타데이터, 전장 상황에서 발생한 이벤트, 질문-응답(QA)데이터를 다양한 전투 양상마다 생성할 수 있다.

본 연구를 통해 생성한 프레임워크를 이용하면 AI 모델 학습에 필요한 그림 1과 같은 군사 데이터를 취득 가능하다. 또한 실제 군사 데이터를 취득할 수 없는 민간 연구원도 전장 데이터를 연구에 활용

가능하다. 이에 따라 본 연구는 군사 목적 모델 훈련 및 전장 분석과 같은 다양한 군사적 응용 분야에서 강력한 도구로 활용될 잠재력을 지니고 있다. 단순한 전장 상황 영상이 아닌 그림 1에서 확인되는 것과 같이 기물들의 위치에 대한 메타데이터를 얻을 수 있기 때문에 detection 모델 연구가 가능하며 위치 데이터도 연속적으로 이용하여 tracking연구도 가능하다. 그림 2와 같은 질문-응답(QA) 데이터도 생성되어 VQA 모델 연구 또한 가능하다.

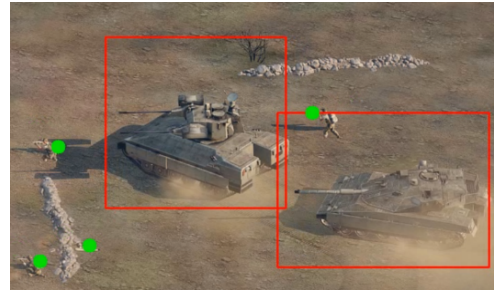


그림 1. 주석이 포함된 전투 장면 생성 예시
Fig. 1. Annotated battle scene generation example

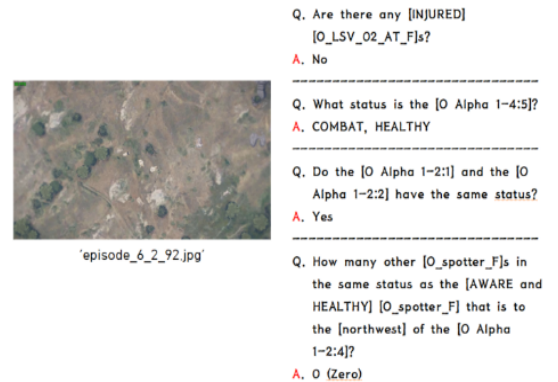


그림 2. 생성된 영상의 특정 frame-QA 레이블 예시
Fig. 2. Example of QA labels for a specific frame in the generated video

II. 관련 연구

2.1 장면 그래프

VQA(Visual Question Answering)을 AI가 진행하기 위해서는 이미지 또는 영상의 정보를 적합한 방식으로 저장해야 한다.

이 과정에서 현재 채택되고 있는 방식 중 하나는 장면 그래프(Scene graph)로 이는 J. Johnson et al.[3]이 제안한 개념으로, 특정 장면 내 객체들에 대한 시각적 정보를 담고 있는 그래프이다. 이 그래프에서 노드는 객체의 바운딩 박스(bounding boxes)와 해당 객체의 속성에 대응하며, 엣지(edges)는 객체 간의 쌍(pair-wise) 관계를 나타낸다. 객체(예: “남자(man)”, “소화전(fire hydrant)”, “반바지(shorts)”, 객체의 속성(예: “소화전은 노란색이다(fire hydrant is yellow)”), 그리고 객체 쌍 간의 관계(예: “남자가 소화전을 뛰어넘고 있다(man jumping over fire hydrant)”)를 그래프에 나타냄으로써 객체 간의 의미 관계를 나타낼 수 있다[4].

다양한 Feature Vector, Visual Embedding과 같은 저장 방식이 아닌 Scene Graph를 채택한 이유는 Scene Graph의 경우 객체 간의 관계를 명시적으로 표현하기 때문에 질의응답 데이터 생성에 있어 유리하기 때문이다.

초기 장면 그래프 생성은 객체를 먼저 탐색한 후 객체 사이의 관계를 탐색하는 방식이 이용되었으나, 높은 검출 성능을 보여주나 복잡한 연산으로 인해 현재 장면 그래프 생성에는 DETR과 같은 모델을 활용하여 end-to-end 기반의 one-stage 방식이 활용되고 있다[5].

비디오에서의 장면 그래프 생성은 이미지에서의 생성과 차이점이 존재한다. 정적 이미지와 다르게 시간적 차원을 포함하고 있어서 비디오 내 객체의 관계는 시공간 그래프(Spatio-temporal graph) 구조를 가지게 된다. 2D 장면 그래프와 차이점은 객체가

고정된 바운딩박스가 아닌 비디오에서의 궤적으로 나타나게 되고, 관계도 특정 시간 동안만 유지가 된다[6].

모델 기반의 장면 그래프 생성은 이미지 정보만을 활용하기 때문에 모델의 성능에 따라 정확도가 달라지며 이로 인한 정보의 손실이 발생할 수 있다. 본 연구에서는 모델을 이용하여 장면 그래프를 생성하는 것이 아닌 시뮬레이션 환경에서 메타 정보를 수집하여 정확한 장면 그래프를 생성하였다.

2.2 시각적 질의 응답 템플릿

T. Qian et al.[7]은 자율주행을 위한 시각적 질의 응답 연구를 진행하였다. 자율주행의 경우 3차원 정보를 분석하는 것이 중요하기 때문에 3차원 장면 정보에 대해 질의 응답을 진행하는 3D Visual Question Answering(3D-QA)에 집중하였다. 3D-QA를 진행하기 위해서 해당 논문은 66개의 질문 템플릿을 제작하였다.

본 논문에서는 T. Qian이 제시한 3D-QA를 위한 질문 템플릿을 참고하여 전장 상황의 3D-QA에 최적화된 질문 템플릿을 제작하였다.

III. Arma3 기반 전쟁 시뮬레이션 영상 및 메타데이터 자동 생성 알고리즘

전장 상황 데이터 생성 도구의 전체 구조는 그림 3과 같다.

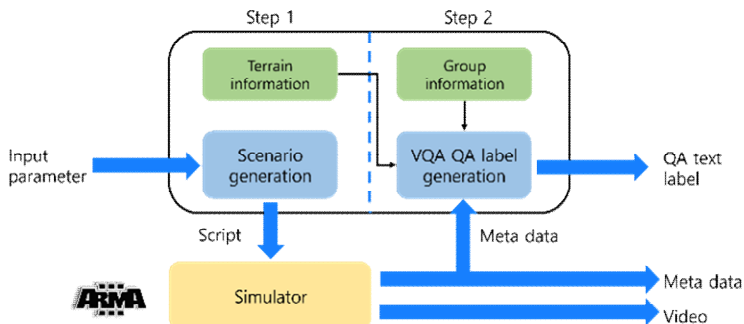


그림 3. 데이터 생성 장치의 구조, Step 1은 전장 영상 생성 및 메타데이터 추출 모듈, Step 2는 메타데이터 및 장면 그래프 기반 QA 텍스트 자동 생성 모듈

Fig. 3. Structure of the data generation device: Step 1 – Battlefield video generation and metadata extraction module, Step 2 – Automatic QA text generation module based on metadata and scene graph

이 도구는 전장 영상 생성 및 메타데이터 추출 모듈, 메타데이터 및 장면 그래프 기반 QA(Question Answering) 텍스트 자동 생성 모듈의 두 주요 구성 요소로 나뉜다.

3.1 전장 영상 생성 및 메타데이터 추출 모듈

전장 영상 생성 및 메타데이터 추출 모듈은 지형 정보 메타데이터 생성 알고리즘과 전장 영상 생성 및 전장 메타데이터 추출 알고리즘으로 구성되어 있다.

3.1.1 지형 정보 메타데이터 생성 알고리즘

지형 정보 메타데이터는 전장 상황에서 실시간으로 추출되는 메타정보가 아니기 때문에 전장 시뮬레이션과 독립적으로 추출된다.

지도에서 다양한 지형 정보를 추출한다. 각 좌표의 고도 데이터를 분석하여 고도 지도를 생성하고, 생성된 지도는 SVG 파일 형식으로 저장하였다.

SVG 지도의 색상을 사용해 지형 유형을 세분화하여 식별할 수 있다. 그림 4는 색상을 분리해 추출한 지형 정보 중 하나인 Urban Area의 일부분이다.

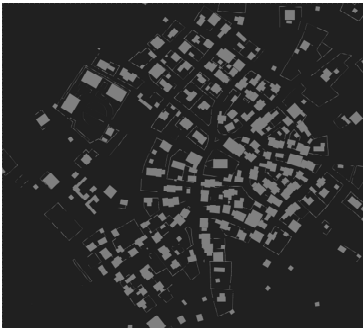


그림 4. 추출한 Urban Area 지형 정보
Fig. 4. Extracted Urban Area terrain information

SVG 파일 분석으로는 알 수 없는 영역들이 존재한다. 숲 또는 도시 영역을 검출하기 위해 나무 위치 데이터, 또는 건물 데이터를 활용하여 특정 반경 내에 나무 또는 건물의 개수에 따라 해당 지역을 숲 또는 도시로 분류했다.

3.1.2 전장 영상 생성 및 전장 메타데이터 추출 알고리즘

영상 생성 및 메타데이터 추출 모듈의 동작 알고리즘은 다음과 같다:

1. 그룹 생성

우호 및 적군 진영(WEST 및 EAST, Arma3 용어 기준)을 사전에 정의된 거리에서 생성하여 두 그룹 간의 초기 분리를 설정한다. 그룹 정보에 대한 메타데이터를 저장한다.

2. 병력 이동 시뮬레이션

양측 진영이 서로를 향해 이동하도록 프로그래밍하여 현실적인 전장 교전을 시뮬레이션한다.

3. 전투 장면 기록

Arma3의 기본 제공 드론 카메라를 활용하여 전투 장면을 비디오로 녹화한다. 이 과정에서 객체 클래스(병사, 차량 등)와 체력 상태 등의 메타데이터를 1초마다 타임스탬프와 함께 추출하고 ini 파일 형식으로 저장한다.

데이터 생성에는 진영 수, 게임 지도, 초기 좌표, 카메라 고도, 시야각(FOV), 초기 사격 범위, 진영 간 거리, 비디오 클립 수, 각 비디오 클립의 길이, 시나리오 수를 입력 매개변수로 사용할 수 있다.



["v-event-fired", [2024,11,13,23,49,32,931],
["b Alpha 1-4:3", "SmokeLauncher",
"SmokeLauncherAmmo"], Alpha 1-4:3]

그림 5. 연막이 터진 전장 상황과 해당 이벤트 메타데이터

Fig. 5. Battlefield scenario with smoke explosion and corresponding event metadata

메타데이터는 객체 클래스, 체력 상태, 위치 정보, 발생한 이벤트 등을 담고 있다. 이벤트 메타데이터에는 `enemydetected`, `unitfired`, `vehiclefired`, `damaged`, `killed` 등이 존재한다. 이벤트 메타데이터 항목은 발생 위치, 대상 유닛, 그리고 각 이벤트와 관련된 세부 정보를 반환하며 CSV 파일로 저장된다. 그림 5는 연막이 터진 전장 상황의 이벤트 메타데이터의 예시이다. 해당 상황의 이벤트 메타데이터는 순서대로 일어난 사건, 시간, 연막을 발사한 기물, 발사한 장치, 발사한 탄약, 연막 피격 유닛이다.

3.2 메타데이터 및 장면 그래프 기반 QA 텍스트 자동 생성 모듈

VQA 모델 학습을 위해 비디오와 해당하는 QA 레이블이 필수적이다. 본 연구에서는 QA 레이블을 자동으로 생성하는 알고리즘을 개발하여, 주작업으로 주석을 달아야 하는 비효율성을 제거하였다. 이 알고리즘은 시간 절약과 비용 효율성을 제공하며, 데이터 주석 과정의 일관성 또한 보장한다. 이 알고리즘은 장면 그래프 생성 모듈과 질문-응답(QA) 텍스트 레이블 생성 모듈로 이루어져 있다. 장면 그래프를 생성하는 이유는 QA 텍스트 레이블 생성을 자동화하기 위함이다.

3.2.1 장면 그래프 생성 모듈

전장 상황 영상에 대한 QA는 이전 상황에 대한 정보를 포함하고 있어야 하므로 시간 정보를 포함하고 있는 장면 그래프가 요구되었다. 그림 6은 장면 그래프의 구조이다.

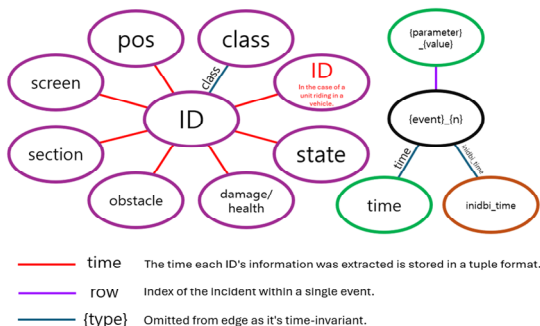


그림 6. 장면 그래프의 구조

Fig. 6. Structure of a generated scene graph

장면 그래프 생성 방식은 다음과 같다.

1. 객체 메타정보를 객체 csv로, 이벤트 메타정보를 이벤트 csv로 저장한다.

2. 유닛들의 위치 정보를 토대로 terrain을 선정하고 그래프에 node로 terrain을 추가한다

3. 객체 csv의 한 row에 저장되어 있는 attribute들을 ID에 연결한다. 이때 edge에는 객체 메타정보가 추출한 시간을 tuple 형식으로 입력한다. class attribute의 경우 시간에 따라 변하는 값이 아니기 때문에 edge에 시간을 입력하지 않는다.

4. 이벤트 csv를 순회하며 3번에서 입력한 메타정보 이전에 일어난 이벤트들을 그래프에 추가한다. 이때 하나의 이벤트로 여러 사건들이 동시에 발생할 수 있다. 따라서 event에서 발생한 첫 번째 사건을 row 1, 두 번째를 row 2와 같은 방식으로 labeling하고 이벤트 node와 연결할 때 edge에 row number를 기재하여 어떤 사건의 parameter인지 확인할 수 있도록 한다.

5. 이벤트가 발생한 시간인 time, 객체 메타정보를 확인할 수 있는 inidbi_time의 경우 사건마다 다른 값이 아니기 때문에 edge에 row number를 기재하지 않고 event node와 연결한다.

시간 정보를 포함하기 위해 전장 상황 영상에서 식별된 객체는 그래프에 추가되거나 기존 정보가 업데이트되고, 식별되지 않은 객체는 가장 최근 상태를 유지한다. 이와 같은 시간 축 통합은 전장 데이터의 변화와 동적 관계를 효과적으로 표현할 수 있다.

3.2.2 질문-응답(QA) 텍스트 레이블 생성 모듈

질문 템플릿 생성에는 [7]을 참조하였다. 질문에는 {}로 이루어진 Placeholder가 존재한다. {}사이에는 질문에 들어가야 할 정보의 속성이 적혀 있고 사전에 생성해 놓은 정보에서 Placeholder의 속성과 일치하는 정보가 랜덤으로 선택되어 {}를 대체하는 방식으로 질문이 생성된다. 이 과정에서 사용하는 정보에는 3.1.1에서 생성한 지형 정보 메타데이터, 3.1.2에서 생성한 전장 메타데이터, 3.2.1에서 생성한 장면그래프에 포함된 노드들의 내용이다.

{ }가 포함된 질문 템플릿마다 쌍을 이루는 답변 파싱 코드를 작성하였다. 해당 코드들은 { }가 랜덤으로 채워진 질문에도 자동으로 답변을 생성할 수 있다. 표 1은 질문과 답변 코드 쌍의 예시이다.

표 1. 질문 템플릿-답변 코드 쌍 예시
Table. 1. Example of question-answer code pair

Question	Answer 파싱 코드 (* Python3 언어)
What is the state of a unit at the {SE1} of the screen {n} seconds before?	<pre> time = self.find_closest_time(self.time[-1], self.values['n']) objects = self.ID_having_attribute(self.values[' SE1']).get(time, []) result=[] for object in objects: result.append(self.attribute_of_ID('st ate', object).get(time)) return result </pre>
{ } is a placeholder that is replaced with information extracted from the scene graph to auto-generate questions. {i}, {Ci}, {Di}, {Si}, {Hi}, {Di}, {Ci}, {SEi}, {i} (i is natural numbers): represent randomly assigned object ID, class, direction, state, health, damage, class, screen position, etc. These random values enable the generation of diverse questions. By generating appropriate answers to questions created with these randomized values, various question-answer pairs can be produced programmatically.	

템플릿 기반 접근법은 효율적이고 일관된 질문 생성을 가능하게 하며, 수작업 주석 없이 대규모 데이터셋을 생성할 수 있는 방법을 제공한다. 이러한 방법은 전장 시각적 질문 응답의 복잡성을 해결하는 데 기여하며, 제안된 벤치마크를 실제 군사 환경에 보다 적합하게 만든다.

IV. 결 론

본 연구를 통해 AI 모델 학습을 위한 현실적인 가상 전투 영상, 메타데이터, QA 데이터를 생성하는 프레임워크를 제안하였다. 이 프레임워크를 이용해 생성된 데이터셋은 다양한 AI 분야에서 활용될 것으로 기대된다. 위치 정보 메타데이터를 활용하여 군사 기물 detection 및 tracking 연구도 가능할 것이

다. QA 데이터를 활용하여 전장 상황에 특화된 VQA 모델 학습 또한 가능하다.

V. 한계점 및 추후 연구 방향

추가적인 데이터를 제공하는 것은 AI 모델들의 학습 성능 향상에 기여할 수 있다. 대표적인 예로, 병사들의 위치를 시각적으로 표현한 미니맵이 있다. 미니맵은 병사들 간의 위치 관계를 명확히 나타내어 모델의 공간적 이해를 강화하고, 이를 통해 상황 인식을 개선하며 더 정확한 응답을 생성할 수 있는 기반을 마련할 수 있다.

현재 생성되는 답변은 단답형이거나 정형화된 템플릿의 빈칸을 채우는 방식에 그치는 경우가 많아, 다양한 문장 구조를 학습하는 데 한계가 있다. 이러한 문제를 해결하기 위해, LLM을 활용하여 단답형 응답을 보다 자연스럽고 풍부한 문장으로 확장해주는 모델을 결합하는 방식을 연구하고자 한다.

추후 연구에서는 본 연구에서 제안한 프레임워크를 바탕으로 데이터셋을 생성하고 전장상황에 특화된 VQA 모델을 연구할 계획이다. VQA 모델을 사용하여 추가 테스트를 수행하고, 상황인식 모델을 개발할 예정이다.

LLaVA[8] 모델을 해당 데이터셋을 이용하여 파인튜닝을 진행하여 데이터셋의 실효성을 입증할 것이다. 이후 LLaVA 모델에 RAG를 부착시키는 방법 등을 활용하여 더 정확한 답변을 생성할 수 있도록 VQA 모델을 고도화 나아가고, 모델 성능을 객관적으로 평가할 수 있는 구체적인 방법을 설계하는 것을 최종 목표로 한다.

References

- [1] East Asia Institute (EAI), "Structure and Future Prospects of the New Cold War on the Korean Peninsula", https://www.eai.or.kr/new/ko/pub/view.asp?intSeq=22707&board=kor_special. [accessed: Feb. 11, 2025]
- [2] Bohemia Interactive, Arma 3, <https://arma3.com/>. [accessed: Feb. 11, 2025]

- [3] J. Johnson, R. Krishna, M. Stark, L. J. Li, D. A. Shamma, M. S. Bernstein, and F. F. Li, "Image retrieval using scene graphs", 2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Boston, MA, USA, pp. 3668-3678, Jun. 2015.
- [4] G. Zhu, L. Zhang, Y. Jiang, Y. Dang, H. Hou, P. Shen, M. Feng, X. Zhao, Q. Miao, S. A. A. Shah, and M. Bennamoun, "Scene Graph Generation: A Comprehensive Survey", arXiv preprint arXiv:2201.00443, Jan. 2022. <https://doi.org/10.48550/arXiv.2201.00443>.
- [5] J. Im, J. Nam, N. Park, H. Lee, and S. Park, "EGTR: Extracting Graph from Transformer for Scene Graph Generation", arXiv preprint arXiv:2404.02072, Apr. 2024. <https://doi.org/10.48550/arXiv.2404.02072>.
- [6] C. Liu, Y. Jin, K. Xu, G. Gong, and Y. Mu, "Beyond short-term snippet: Video relation detection with spatio-temporal global context", Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. (CVPR), Seattle, WA, USA, pp. 10840-10849, Jun. 2020.
- [7] T. Qian, J. Chen, L. Zhuo, Y. Jiao, and Y.-G. Jiang, "NuScenes-QA: A Multi-Modal Visual Question Answering Benchmark for Autonomous Driving Scenario", Proc. AAAI Conf. Artif. Intell., Vancouver, Canada, Vol. 38, No. 2, pp. 28253, Feb. 2024.
- [8] H. Liu, C. Li, Y. Li, and Y. J. Lee, "Improved Baselines with Visual Instruction Tuning", Proc. IEEE/CVF Conf. Comput. Vis. Pattern Recognit. (CVPR), Seattle, USA, pp. 26296-26306, Jun. 2024.

저자소개

김 성 현 (Sungheon Kim)



2025년 2월 : GIST
전기전자컴퓨터공학과(학사)
2025년 3월 ~ 현재 :
광주과학기술원 정보컴퓨팅대학
전기전자컴퓨터공학과 석사과정
관심분야 : AI, Deep Learning

박 종 현 (Jonghyun Park)



2021년 2월 : 수원대학교
정보통신공학과(학사)
2023년 8월 : 광주과학기술원
전기전자컴퓨터공학과(공학석사)
2023년 9월 ~ 현재 :
광주과학기술원 정보컴퓨팅대학
전기전자컴퓨터공학과 박사과정

관심분야 : AI, Robotics

김 예 찬 (Yechan Kim)



2021년 8월 : 광주과학기술원
전기전자컴퓨터공학부(공학석사)
2021년 9월 ~ 현재 :
광주과학기술원 정보컴퓨팅대학
전기전자컴퓨터공학과 박사과정
관심분야 : 인공지능, 패턴인식,
원격탐사, 컴퓨터비전

권 철 희 (Cheolhee Kwon)



1998년 2월 : 고려대학교
제어계측공학(공학사)
2000년 2월 : 고려대학교
제어계측공학(공학석사)
2000년 1월 ~ 현재 :
LIG 넥스원 연구위원
관심분야 : AI, Deep Learning

조 규 태 (Kyutae Cho)



2002년 2월 : 숭실대학교
전산학(공학사)
2004년 2월 : 한국과학기술원
전산학(공학석사)
2007년 2월 : 한국과학기술원
전산학(박사수료)
2007년 2월 ~ 현재 :

LIG 넥스원 수석연구원

관심분야 : AI, Deep Learning

진 정 훈 (Junghun Jin)



2002년 2월 : 고려대학교
전자공학(공학사)
2004년 2월 : 고려대학교
영상정보처리(공학석사)
2004년 1월 ~ 현재 :
LIG 넥스원 수석연구원
관심분야 : AI, Deep Learning

김 지 아 (Jia Kim)



2007년 2월 : 이화여자대학교
컴퓨터공학(공학사)
2006년 12월 ~ 현재 :
LIG 넥스원 수석연구원
관심분야 : AI, Deep Learning

전 문 구 (Moongu Jeon)



2001년 6월 : University of
Minnesota, Scientific
Computation(공학박사)
2001년 7월 ~ 2003년 6월 : UCSB,
USA, 박사후연구원
2003년 7월 ~ 2005년 8월 :

Institute for Biodiagnostics, NRC, Canada 연구원

2005년 9월 ~ 현재 : 광주과학기술원 정보컴퓨팅대학

전기전자컴퓨터공학과 교수

관심분야 : 인공지능, 기계학습, 컴퓨터비전