

자율주행 자동차를 위한 종단간 학습 기술 중심의 기술 동향

최 정 환*, 박 예 찬*, 전 우 민**, 이 성 진[°]

Research Trends Focused on End-to-End Learning Technologies for Autonomous Vehicles

Jeong-hwan Choi*, Yechan Park*, Woomin Jun **, Sungjin Lee[°]

요 약

자율주행 기술은 미래 자동차 산업의 핵심 기술로 연구되면서, 최근 관련 센서, 부품, 플랫폼 및 인공지능 기술을 통해 관련 기술이 급격한 발전을 이루고 있다. 본 논문은 최근 자율주행 기술의 주요 이슈로 대두되는 종단간 학습 기술에 관한 논문으로 인지, 판단, 제어 각 모듈을 단일 통합 신경망으로 대체하는 방식에 대해 알아본다. 우선 전통적인 모듈형 자율주행 시스템의 인지, 판단, 제어의 모듈 구조, 한계점에 대해 살펴보고 종단간 학습 자율주행 시스템의 구조인 이미테이션 러닝과 강화학습의 주요 기술, 훈련 방법, 입력과 출력 형태, 평가 방법 및 지표에 대해 분석한다. 마지막으로 현재 종단간 학습의 기술 이슈 사항 및 관련 구조에 대해 알아보면서 다양한 연구 사례와 실험 결과에 대해 알아본다.

Key Words : Autonomous Driving, End-to-End Learning, Modular Architecture, Perception, Planning, Control

ABSTRACT

As autonomous driving technology is researched as a core technology in the future automotive industry, recent advancements in related sensors, components, platforms, and artificial intelligence technologies have led to rapid progress in this field. This paper focuses on end-to-end learning technology, which has recently emerged as a major issue in autonomous driving technology. It explores the approach of replacing individual perception, decision-making, and control modules with a single integrated neural network. First, the paper examines the structure and limitations of traditional modular autonomous driving systems, which consist of separate perception, decision-making, and control modules. Then, it analyzes the structure of end-to-end learning autonomous driving systems, highlighting key technologies such as imitation learning and reinforcement learning, training methods, input and output formats, and evaluation methods and metrics. Finally, the paper reviews current technical issues and related structures of end-to-end learning, presenting various research cases and experimental results.

* First Author : Dong-Seoul University, Department of Electric Engineering, wjdghks987@naver.com, 학생회원

[°] Corresponding Author : Dong Seoul University Department of Electronic Engineering, sungjinlee@du.ac.kr, 정회원

* Dong-Seoul University, Department of Electric Engineering, pcy0504@naver.com

** Dong-Seoul University, Department of Electric Engineering, aplus912@naver.com, 학생회원

논문번호 : 202407-125-E-RN, Received June 28, 2024; Revised July 19, 2024; Accepted Aug 7, 2024

I. 서 론

자율주행 기술은 최근 급격한 발전을 이루며 자동차 산업 전반에 걸쳐 혁신을 일으키고 있다¹⁻³⁾. 자율주행차는 도로 안전을 향상시키고 교통 효율성을 극대화하며, 궁극적으로는 인간의 개입 없이 차량이 스스로 운전할 수 있는 미래를 목표로 한다. 이러한 기술의 핵심은 차량이 주변 환경을 인식하고, 이를 바탕으로 최적의 주행 경로를 계획하며, 실제로 차량을 제어하는 일련의 과정을 통합 최적화하는 것이다. 이 모든 과정은 고도로 복잡한 문제로, 다양한 센서 데이터의 실시간 처리, 복잡한 도로 환경에서의 의사결정, 그리고 안전하고 정확한 제어가 필요하다^{2,3)}.

그림 1의 (a) 모듈형 구조(Modular Architecture)에서 보듯이 전통적인 자율주행 시스템은 크게 인지(Perception) 및 위치추정(Localization), 판단(Planning), 제어(Control)의 세 단계로 나뉜다. 각 단계는 고유한 알고리즘과 기술을 사용하여 독립적으로 개발되고 조정된다^{2,3)}. 예를 들어, 인지 단계에서는 카메라, 라이다, 레이더 등의 센서를 통해 주변 환경을 인식하고, 객체를 감지하며, 도로와 차선을 파악한다. 판단 단계에서는 인지된 정보를 바탕으로 차량의 주행 경로를 계획하고, 돌발 상황에 대한 대응 방안을 마련한다. 마지막으로, 제어 단계에서는 결정된 경로와 행동을

실제 차량의 움직임으로 구현하기 위한 제어 명령을 생성한다. 이러한 모듈형 구조 접근법은 각 단계에서의 독립적인 최적화를 가능하게 하지만, 각 모듈 간의 상호작용의 복잡성이 증가하여 데이터 변환, 통합의 어려움이 발생하고 각 모듈의 개별성능으로 전체 성능이 제한되는 단점이 있다.

최근에는 이러한 문제를 해결하고 자율주행 시스템의 성능을 더욱 향상시키기 위해 종단간 학습(End-to-end learning) 방식이 연구, 도입되고 있다⁴⁾. 그림 1의 (b)에서 보듯이 종단간 학습은 입력 데이터에서 최종 출력까지의 모든 과정을 하나의 통합된 신경망 모델로 처리하는 접근법이다. 즉, 카메라 이미지나 라이다 포인트 클라우드와 같은 원시 데이터를 입력받아 차량의 스티어링, 가속, 제동 등의 제어 명령을 직접 출력하는 방식이다. 이러한 접근법은 전체 시스템을 하나의 학습 모델로 통합함으로써 데이터 처리의 일관성을 유지하고, 각 단계 간의 복잡한 상호작용을 간소화할 수 있으며, 인지-판단-제어 각 단계의 복잡한 패턴과 상호작용을 학습하는데 유리한 딥러닝 모델이 사용되어 유기적 연동을 가능케 한다.

종단간 학습 방식은 특히 딥러닝 기술의 발전과 함께 큰 성과를 이루고 있다. 딥러닝 모델은 대규모 데이터셋을 활용하여 다양한 주행 상황에 대한 학습을 통해 높은 수준의 일반화 능력을 갖출 수 있다. 예를 들어,

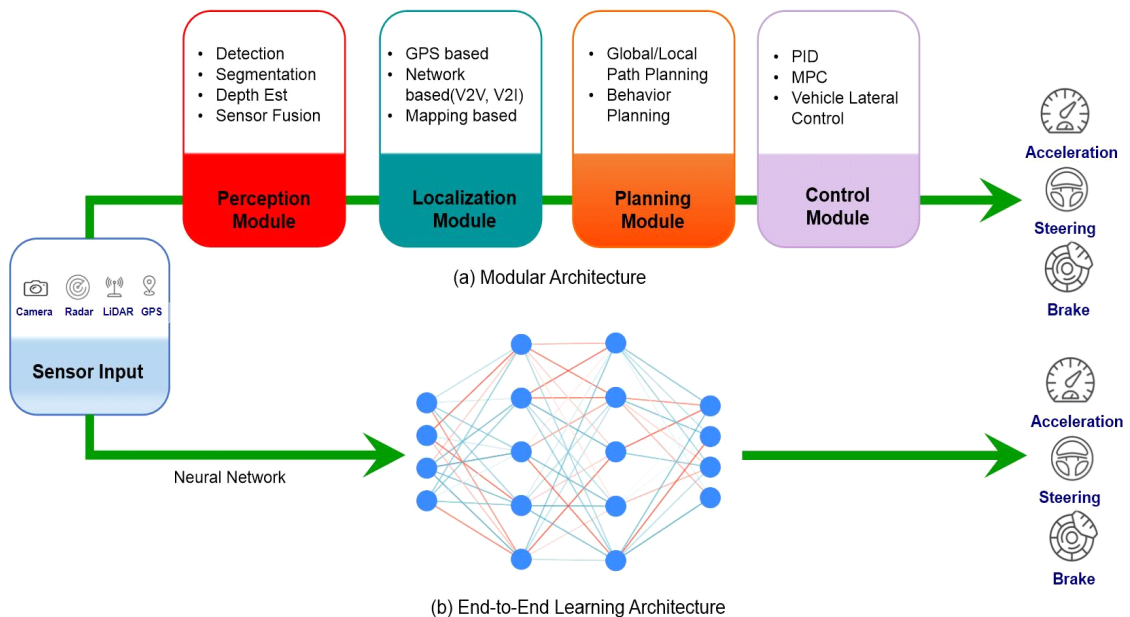


그림 1. 모듈형 구조와 종단간 학습 구조 비교

Fig. 1. Comparison of Modular Architecture and End-to-End Learning Architecture

Convolutional Neural Networks(CNNs)는 이미지 데이터에서 유의미한 특징을 추출하고, Recurrent Neural Networks(RNNs)는 시간에 따른 연속적인 데이터의 패턴을 학습하는 데 유리하다. 또한, Transformer 구조는 Positional Encoding, Self-Attention 기술들을 통해 이미지 데이터 및 연속적 데이터의 특징 모두를 유리하게 학습할 수 있는 구조를 제공한다. 강화학습(Reinforcement Learning) 기술은 자율주행차가 다양한 주행 환경에서 최적의 행동을 선택할 수 있도록 학습시키는 데 사용된다.

그러나 종단간 학습 방식에도 도전 과제가 존재한다^[2,3]. 가장 큰 문제는 학습 모델의 해석 가능성이다. 전통적인 모듈화된 시스템은 각 단계의 결과를 각자의 정확도 메트릭으로 해석할 수 있지만, 종단간 학습 모델은 입력과 출력 사이의 중간 과정의 성능과 판단 이유를 설명하기 어렵다. 이는 특히 안전이 중요한 자율주행차 분야에서 큰 우려를 불러일으킨다. 따라서, 종단간 학습 모델의 해석 가능성을 높이기 위해 자율주행 시스템과 MLLMs을 접목하는 등의 연구가 진행되고 있다.

이 논문에서는 자율주행 기술의 발전을 이끄는 종단간 학습 방식의 기술 동향을 살펴보고, 현재의 도전 과제와 미래의 연구 방향을 제시하고자 한다. 이를 통해 자율주행차의 안전성과 성능을 극대화할 수 있는 새로운 접근법을 모색하고, 궁극적으로는 완전한 자율주행 시대의 도래를 앞당기는데 기여 하고자 한다.

본 논문의 2장에서는 인지 및 위치추정, 판단, 제어로 구성되는 모듈형 구조와 각 모듈에서 수행되는 주요 작업 및 역할에 대해 살펴본다. 3-4장에서는 종단간 학습의 입출력 구조 및 모델 구조, 기타 이슈들에 대해 논의한다. 5장에서는 종단간 학습을 위한 평가 지표, 관련 데이터셋과 시뮬레이터에 대해 소개하며 마지막으로 6장에서는 종단간 학습의 기술들에 대한 고찰 및 향후 발전 방향에 대해 논의한다.

II. 모듈형 구조

자율주행 기술은 주변의 다양한 센서 정보를 입력으로 받아 사람의 개입 없이 시스템 스스로 하나의 완결하고 통합된 판단을 내리고 차량을 운전하는 기능을 제공한다. 그림 2는 차량과 환경, 그리고 운전자가 어떻게 상호작용하는지를 보여준다. 자율주행 시스템은 운전자를 대체하여 동적인 환경 변화를 센서 통해 인지하고, 차량의 경로계획, 판단, 차량 제어 역할을 수행해야 한다^[4]. 본 섹션에서는 모듈형 구조를 통해 이런 자율주행 시스템의 세부 기술들에 대해 알아본다.

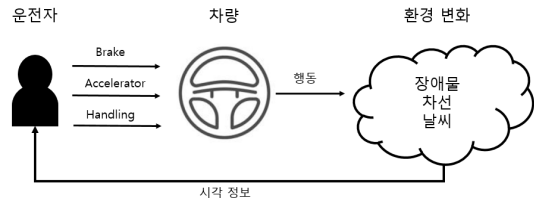


그림 2. 운전자-차량-환경 폐쇄 루프 시스템 모델
Fig. 2. Driver-Vehicle-Environment Closed-Loop System Model

2.1 모듈형 파이프라인 구성 요소

모듈형 구조는 시스템을 여러 개의 독립적이고 자율적인 기능 단위의 모듈로 나누어서 기능들을 각 모듈이 맡아 수행하는 접근방식이다. 예측할 수 없는 교통 상황에서도 모듈들이 상호작용하며 완전한 자율주행 기능을 해내는 것이 모듈형 구조의 목표이다^[5].

다음은 모듈형 구조의 인지 모듈, 위치추정 모듈, 경로계획 모듈, 제어 모듈을 각각 섹션 2.1.1), 섹션 2.1.2), 섹션 2.1.3), 섹션 2.1.4)에 소개한다.

2.1.1 인지 모듈

인지 모듈은 카메라, 라이다 등의 센서로 차량과 차선, 보행자 같은 주변 환경을 인식하고 객체를 감지하는 모듈이다. 인지 기술은 객체탐지(Object Detection), 영상분할(Image Segmentation) 등을 중심으로 2D 및 3D 영상 데이터에 따라 다양한 활용 기술들, 즉, 차선 인식^[6](Lane Detection), 깊이 인식(Depth Estimation), 3차원 복원(3D Reconstruction), 조감도 영상분할(Bird's Eye View Segmentation) 등이 연구되었다. 또한, 사용되는 센서에 따라 카메라(Camera) 기반, 라이다(Lidar) 기반, 센서 융합(Sensor Fusion) 기반 등의 기술이 연구되었다.

2D 객체탐지 기술은 RPN(Region Proposal Network) 유무에 따라 크게 2단계 탐지(Two-Stage Detection)와 1단계 탐지(One-Stage Detection) 기술로 분류되며, 일반적으로 1단계 탐지가 자율주행의 실시간 처리에 더 유리하다.

마찬가지로, 영상분할, 차선인식, 깊이인식, 3차원 복원, 조감도 영상분할 연구에서도 정확도 지표와 더불어 실시간 처리는 중요한 요소로 작용하여 이런 딥러닝 연산 경량화는 중요 이슈로 연구되었다^[7].

단안 카메라(Monocular) 3D인지 연구에서의 Pseudo-LIDAR^[8]는 2D 이미지에서 깊이를 추정하고 3D 좌표로 변환하여 “가상 라이다 포인트 클라우드” 맵을 구성한다. 이밖에 인코더-디코더 구조에 기반한 다양한 기법들이 정확도 및 지연시간 감소를 위해 연구

되었다.

반면, LiDAR는 무질서한 점들로 이루어진 클라우드 형태의 데이터로, 유용한 특징을 추출하기 위한 방법들이 연구되었다. 복셀 기반 VoxelNet^[9]과 PointPillars^[10]는 포인트 클라우드를 일정 간격의 복셀로 나눈 후, CNN을 적용하는 방법이다. 포인트 기반 PointRCNN^[11]은 백본에서 3D 바운딩 박스를 예측하고, 다음 단계에서 관심 영역 풀링(ROI Pooling)을 수행한다. 포인트-복셀 기반 방법인 BADet^[12]는 포인트 클라우드의 복셀 특징, 픽셀 특징을 집계하여 더 풍부한 지역적 특징을 추출할 수 있다.

최근에는 카메라-라이다 간 센서 융합을 통해 카메라의 풍부한 텍스처 정보와 라이다의 정확한 거리 정보를 결합하여 영상인지를 수행한다^[13]. 연구 [Simple-BEV]에서는 센서퓨전 기반 조감도 영상분할 연구 SimpleBEV를 수행하여 nuScenes 데이터셋 상에서 60 IoU (vehicle) 성능을 얻었다.

2.1.2 위치 추정 모듈

위치추정 모듈은 자율주행차의 정확한 위치를 추정하는 모듈이다. 첫 번째로 GPS 센서를 통한 위치추정 방식, 두 번째로 V2V(Vehicle to Vehicle) 혹은 V2I(Vehicle to Infrastructure)를 통한 네트워크 기반 위치 추정방식, 세 번째로 맵핑 기반 위치추정 방식이 있다. GPS 센서를 통한 방식은 차량의 주변환경에 따라 그 정확도에 안정성이 확보되지 않아, V2V와 V2I를 결합한 V2X(Vehicle to Everything)를 통해 그 값을 보정하거나^[14] 세 번째 방식인 맵핑 기반 위치추정 방식을 통해 보정하기도 한다. 특히 맵핑 기반 위치추정을 위해 고정밀 LiDAR를 통해 3D 포인트 클라우드 맵을 만들고 차량의 위치를 추정하는 방법인 NDT(Normal Distribution Transform) 스캔 매칭과 차량의 현재 속도, 방향, 시간 등의 Dead Reckoning results를 결합하여 차량의 위치를 정밀 추정하는 시도가 있었다^[15].

2.1.3 경로계획 모듈

경로계획 모듈은 경로를 계획하고 최적의 주행 경로를 결정하는 모듈이다. 보통 위치추정 모듈에서 나온 위치 정보를 기반으로 경로를 계획한다. 해당 경로계획 모듈은 전역 경로계획(Global Path Planning), 지역 경로계획(Local Path Planning), 고수준 작업인 행동계획(Behavior Planning)의 기술로 분류된다.

전역 경로계획은 목적지까지의 전체 경로를 계획한다. 도로 지도와 같은 대규모 정보에서 해당 전역 경로 계획 알고리즘을 활용한다. 지역 경로계획은 차량의 현

재 위치에서 단기적인 경로계획을 의미한다. 센서 데이터로 인지되는 주변 환경 정보에 따라 경로를 실시간으로 수정한다. 보통 전역 경로계획과 지역 경로계획을 사용한다^[16].

행동계획^[17]은 차량이 특정 상황에서 어떻게 행동할지를 결정하는 과정을 의미한다. 이는 주로 Mission Planning, Motion Planning을 포함하며 여기에 최적의 차선 변경과 교차로에서의 우선순위 판단 등의 동작이 보조적으로 수행된다. 그림 3 Mission Planning은 차량의 네트워크 정의 파일(Route Network Definition File, RNDF), 실시간 교통 정보 데이터를 사용하여 출발지에서 목적지까지의 계획을 수행한다. 이는 전역 경로계획과 비슷한 역할을 하지만, 정적 데이터만 사용하는 것이 아닌, 동적 데이터를 사용하여 예기치 않은 상황에 대한 대응과 변화하는 상황에 따라 경로를 재계획하는 유연성이 포함된다.

Motion Planning은 행동계획에서 선택한 행동을 구체적으로 실행하기 위해 생성되는 세부 경로이다. 이 결정을 따르기 위해 어떤 경로와 판단을 따라야 하는지, 어떤 속도와 조향 각도로 이동해야 하는지를 계산한다.

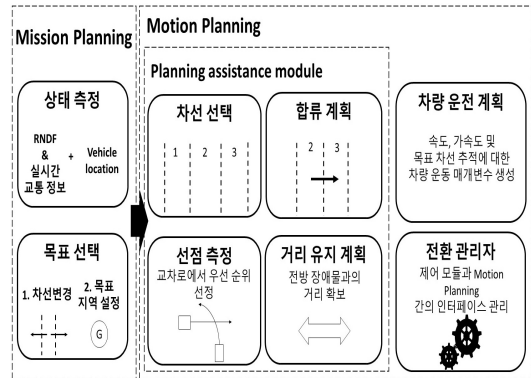


그림 3. 행동 계획의 예시
Fig. 3. Example of Behavior Planning

2.1.4 제어 모듈

제어 모듈은 차량을 제어하여 주어진 경로를 따라가는 모듈이다. 제어 모듈은 브레이크와 엑셀, 조향을 조작하여 차량이 안정적으로 주행 경로를 유지하게 해준다. 제어 모듈에는 저수준 제어와 고수준 제어로 나뉘어진다.

저수준 제어란 스티어링, 가속기 등의 차량의 직접적인 물리적 제어를 수행하는 것을 의미한다. 저수준 제어 중 하나인 비례-미분-적분(Proportional-Integral-Differentia, PID) 알고리즘을 사용한 사례^[4]를 보면 PID 제어가 자

동차 스티어링과 속도 등의 하드웨어에 직접 연결되어 동작하는 것을 확인할 수 있다. 저수준 제어의 또 다른 제어 방식으로 Vehicle Lateral Control이 있다. Pure pursuit^[18]은 자동차 뒷바퀴와 추종 목표점 사이의 전방 주시 거리로 조향각을 제어한다. Stanley^[19]는 자동차 앞바퀴를 기준으로 목표점까지의 경로와 얼마나 벗어나 있는지의 경로 교차 오차와 목표점과 방향이 얼마나 틀어져 있는지의 조향 방향 오차를 고려하여 조향 각을 제어한다.

고수준 제어란 자율주행 시스템이 상황에 맞춰 직접 의사결정과 전략적 제어를 수행하는 것을 의미한다. 모델예측제어(MPC)^[15]는 유한한 미래 구간 내에서 미래의 상태를 예측하고 다음 행동에 대한 최적의 제어 입력을 계산한다. 상대적으로 저수준 제어보다 유동적이고 변화에 강하다.

2.1.5 모듈형 구조 한계와 종단간 학습 대두

앞서 언급되었던 것처럼, 모듈형 구조는 인지, 위치 추정, 경로계획, 제어와 같은 여러 단계가 별도로 학습되고 조정된다. 하지만 모듈형 구조는 한 모듈의 오류가 전체 시스템에 전파하는 오류 전파 문제^[20]와 각 모듈 간 상호작용의 통합이 어렵다는 단점^[2,3]이 있어 그 성능적 제한이 발생한다. 반면, 종단간 학습 기반 구조는 모듈형 구조의 각 기능들을 하나의 신경망 모델로 통합하여 학습하는 방식으로 이를 통해 데이터 처리 파이프라인을 단순화하고, 학습 과정에서 모든 단계를 동시에 최적화할 수 있다^[21]. 또한, 각 기능별로 개별적인 모델을 사용하는 대신 하나의 모델을 공유함으로써 계산 자원을 절약할 수 있으며, 데이터를 기반으로 직접 학습하기 때문에 더 많은 데이터를 추가하면 모델의 성능이 향상될 수 있다. 결과적으로, 모델은 데이터로부터 주행 방식을 스스로 학습하게 된다^[22].

3, 4장에서는 이런 종단간 학습의 입출력 구조, 모델 구조에 대해 서술한다.

Ⅲ. 종단간 학습 기반 입출력 형식

종단간 학습 기반 자율주행은 Camera, LiDAR, Navigations와 같은 정보가 제어 신호를 생성하는 백본 모델의 입력으로 사용되며, 이는 가속, 조향, 브레이크와 같은 동작을 출력으로 내보낸다.

3.1 INPUT Modal

이번 섹션에서는 종단간 학습 기반 자율주행에 필수적인 입출력 Modal에 대해 탐구한다.

입력 Modal에는 Camera, LiDAR, Multi-Modal, 및 네비게이션에 대해 설명하고, 출력 Modal은 WayPoint, 조향 및 속도, 비용 함수에 대해 설명한다.

3.1.1 Camera

카메라 기반^[23]종단간 학습 주행은 카메라 영상을 이용해 주행 명령을 직접 생성한다. 이는 복잡한 주행 시나리오 처리와 실시간 주행에서 높은 성능을 보인다.

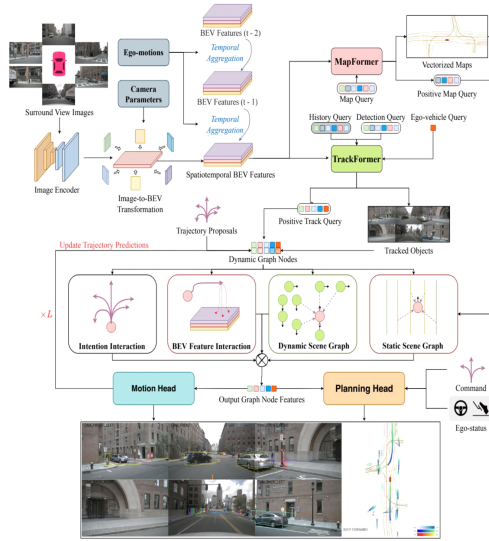
단안 카메라는 입력 영상에서 직접 제어 신호를 생성하는 방식으로 자율주행을 구현한다^[24]. 초기 연구^[25]에서는 단안 카메라를 이용해 출력 조향과 인간 운전자의 명령 사이의 평균 제곱 오차(MSE)를 최소화하여 차선을 따라 주행하는 데 성공했다. 또한 YouTube 비디오를 통해 시각적 특징과 주행 책을 학습하는 방식도 연구되었다^[26]. 최근에는 이미지 시퀀스를 고수준의 토큰으로 변환하고, 트랜스포머 네트워크에 입력하여 주행 명령을 예측하는 모델도 개발되었다^[27]. 단안 카메라는 저렴하고 여러 상황에서 사용하지만, 깊이 감지 기능에 한계가 있다.

스테레오 카메라는 다양한 시각에서 정보를 융합해 주행 정책 효율성을 높인다^[28]. 이 연구^[29]는 RGB 이미지와 깊이(Depth) 데이터를 입력으로 사용하였고 조정부 모방 학습(CIL)^[30]을 CNN 구조로 사용하였다. 스테레오 카메라는 물체를 식별하거나 분할과 같은 유용한 색상 및 텍스처 속성을 가지나 날씨 및 조명 조건에 취약하다.

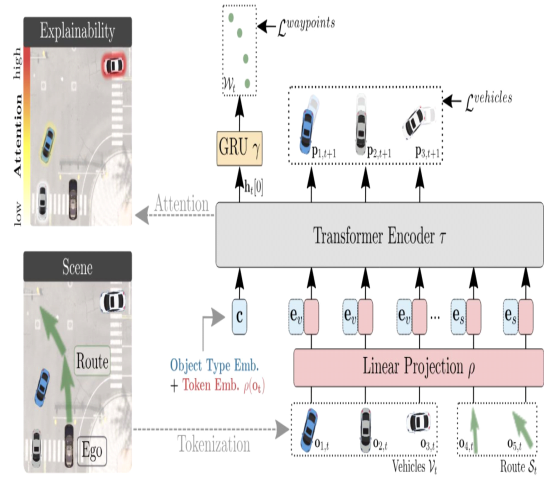
서라운드 뷰 카메라는 도시 주행 및 교차로 시나리오에서 단안 카메라로 발생하는 시야각 감소 및 주변부 왜곡 단점을 개선한다^[31,32]. SparseDrive^[33]는 주행 환경에서 필요한 정보만을 포함하여 데이터의 간결성과 처리성을 높이고, “승자가 모든 것을 가져가는” 손실함수를 사용하여 주행의 안정성을 높였다.

3.1.2 LiDAR

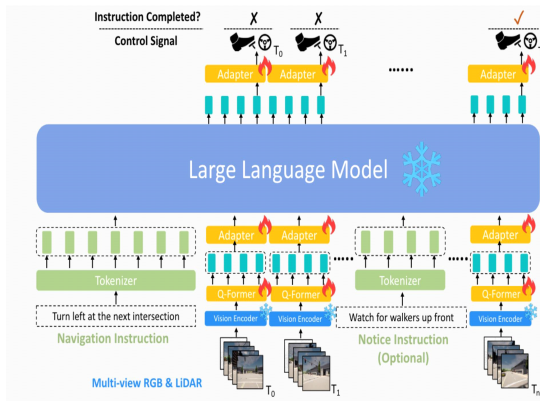
LiDAR 데이터를 사용하는 이유는 빛이 부족한 환경에서도 정확한 거리 측정을 제공하며, 가장 포괄적인 공간 정보를 제공하기 때문이다^[34]. 하지만 LiDAR만 사용하였을 때 데이터의 적은 해상도로 인해 물체가 작은 클래스에 대해 검출하지 못하는 한계가 있다^[35]. 또한 HDMapNet^[36] 연구에서는 LiDAR만 사용한 경우에 전체적인 클래스에서 낮은 성능을 보여주었다. 그러므로 해당 연구에서는 LiDAR의 정확한 위치 표현력과 카메라의 풍부한 텍스처 표현력을 결합한 카메라-LiDAR 융합 모델을 통해 인식 성능을 개선하는 연구가 수행되었다.



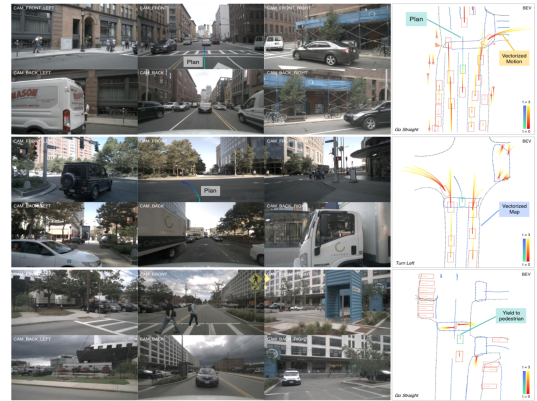
(a) GraphAD^[41]



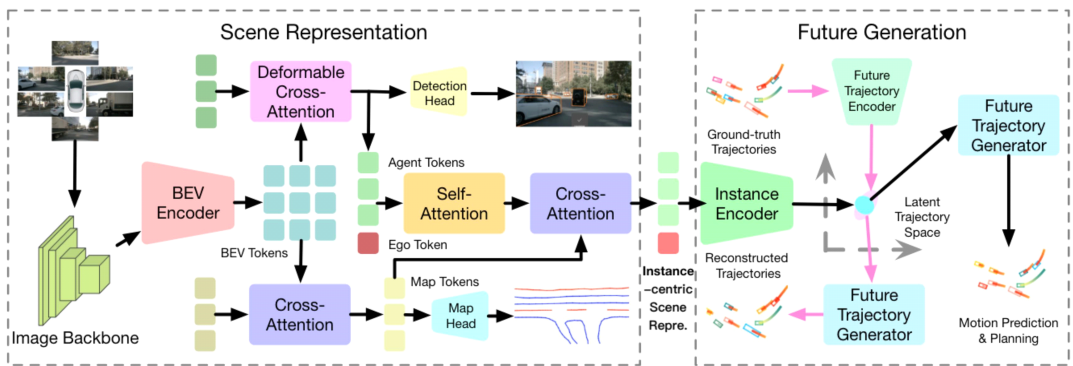
(b) PlanT^[37]



(c) LMDrive^[48]



(d) VAD^[42]



(e) GenAD^[52]

그림 4. 다양한 입력종단간 학습 구조

Fig. 4. Diverse Input-Output End-to-End Learning Architecture

3.1.3 Multi-Modal

Multi-Modal 시스템은 다양한 센서 데이터를 결합하여 보다 정확하고 신뢰성 있는 정보를 추출하는 방법이다^[29]. 자율주행차에서는 카메라, LiDAR, GPS 등 다양한 센서를 사용하여 주행 환경을 인식하고 융합하여 종합적인 주행 정보를 제공한다. Camera, LiDAR Fusion을 사용한 모델^[37]은 트랜스포머 모델을 사용하여 성능 비교 결과 TransFuser보다 10.36점, 2.7배 더 빠른 결과가 나왔고 LAV보다 24.92점 더 높은 성능을 보였다. 여기에서 점수는 CARLA^[2] 자율주행 시뮬레이터에서 평가된 주행 성능 점수를 의미한다.

Multi-Modal의 범주에는 그림 5. 처럼 초기 융합(Early fusion), 중간 융합(Mid Fusion), 후기 융합(Late Fusion)이 있다.

초기 융합은 학습이 가능한 종단간 학습 시스템에 입력하기 전에 센서 데이터를 결합하는 방식이다. 융합 전에 전처리가 필요하고 계산 효율성이 높으며, RGB 색상과 Depth 채널을 연결하는데 성능이 우수하다. EarlyBird 모델^[38]은 카메라에서 추출한 이미지 특징을 BEV로 투영하여 하나의 고차원 BEV 특징 맵으로 결

합하는 초기 융합 방식을 사용하여 객체탐지 및 탐지 추적의 정확성을 크게 향상시키고 가림 현상과 탐지 누락 문제를 효과적으로 해결한다.

중간 융합은 네트워크 내에서 일부 전처리 단계 또는 Modal에 대한 특징 추출이 완료된 후 Modal을 결합하는 방식이다. 특징 추출 자체는 종단간 학습 모델의 일부이며, 학습이 가능하다. Multi-View Fusion^[39]은 LiDAR 포인트 클라우드의 BEV와 원근 뷰가 제공하는 정보를 활용하여 중간 융합 한다.

TransFuser^[40]은 Self-Attention 기법을 사용해 이미지와 LiDAR를 여러 해상도에서 다중 트랜스포머 모듈을 사용하여 중간 융합 한다. 또한 속도와 가속도와 같은 차량의 상태 측정치는 시각적인 입력과 주로 중간 융합된다^[46]. 그림 4. (a)의 GrapAD^[41]는 멀티 뷰 이미지 데이터를 BEV로 변환 후 이미지, 카메라 매개변수, 차차 위치 정보 등을 중간 융합으로 통합한다. 최근에는 대형 언어 모델(LLM)과 통합하여, 사용하는 모델^[42]도 개발되었다.

후기 융합은 여러 데이터 분기의 결과를 병합하는 방식이다. 즉, 각 입력 Modal에 대한 출력을 별도로 계

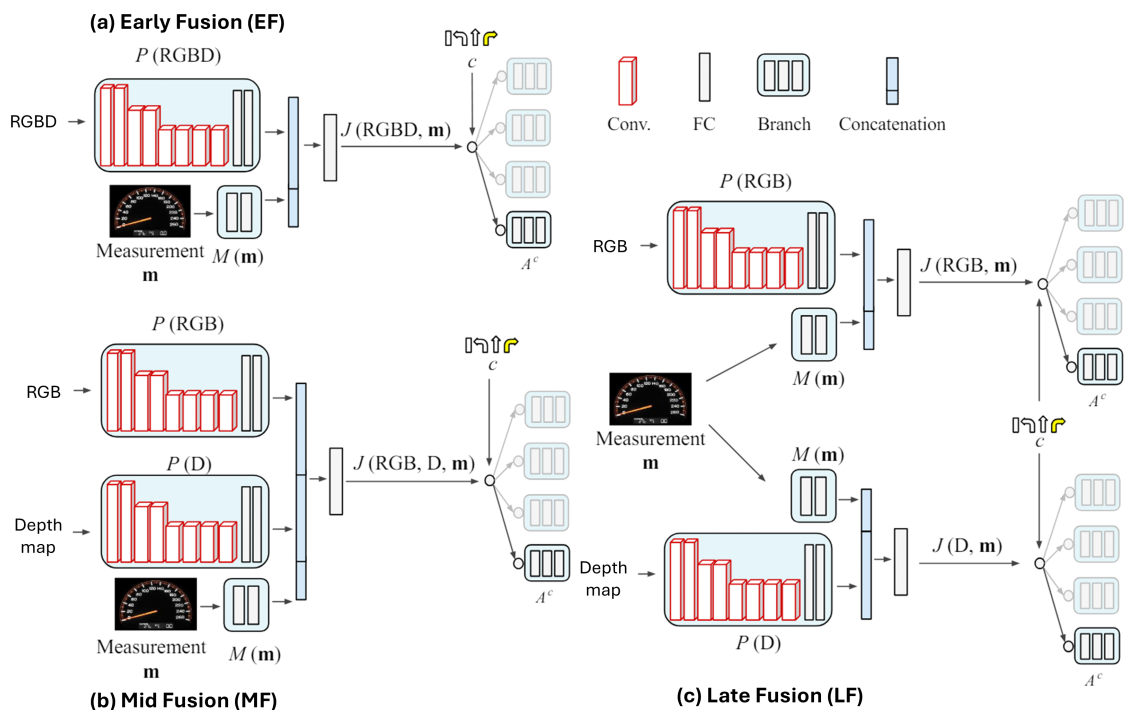


그림 5. RGB 이미지와 Depth의 다양한 Multi Modal^[29], (a). Early Fusion: RGB와 Depth 채널이 백본의 입력으로 사용, (b). middle Fusion: RGB와 Depth채널의 특징 Layer가 융합, (c). Late Fusion: RGB와 Depth 채널의 각각의 출력이 융합
Fig. 5. RGB Image and Depth in Various Multi-Modal^[29], (a) Early Fusion: RGB and Depth channels are used as inputs to the backbone, (b) Middle Fusion: Features from RGB and Depth channels are fused at an intermediate layer, (c) Late Fusion: The outputs of the RGB and Depth channels are fuse

산한 후 별도의 출력을 특정 방식으로 결합한다. 후기 융합은 칼만 필터^[43] 또는 Export 혼합^[44]을 사용하는 양상불이다. 후기 융합은 좌회전, 우회전, 직진과 같은 네비게이션 명령을 입력하는데 초기 결합보다 더 효율적인 것으로 나타났다^[30].

3.1.4 네비게이션

종단간 학습 주행 시스템에서 네비게이션 Modal은 자율주행차가 목적지에 도달하기 위해 필요한 경로와 방향을 안내하는 역할을 한다. 이 입력 Modal은 차량이 도로 네트워크 내에서 효과적으로 움직이도록 지원하며, 네비게이션 입력은 경로계획(Path Planning), 네비게이션 명령(Navigation Commands)과 텍스트 명령(Textual Commands)으로 비롯될 수 있다.

경로계획^[45]은 차량이 현재 위치에서 목적지까지 최적의 경로를 계산하는 시스템으로, 네비게이션 화면과 같은 시각적인 형태로 모델에 삽입할 수 있다. ChauffeurNet^[46]은 HD 맵을 기반으로 작동하며 원하는 경로를 이미지로 제공한다. 또한, 그림 4. (b)의 PlanT^[37]는 목표 위치의 입력을 기반으로 지점 간 네비게이션을 활용한다.

네비게이션 명령은 차량이 주행 중 특정 명령을 따를 수 있도록 제공되는 지시사항이다. 이는 운전 모델을 특정 방식으로 운전하도록 요청할 수 있으며, 좌회전 또는 우회전 명령을 삽입하는 거나 차선의 오른쪽 또는 다른 차량을 따라가도록 지시할 수 있다^[46]. Huang et al^[47]은 카메라와 라이다의 Multi-modal 센서 데이터와 네비게이션 명령이 네트워크의 입력으로 들어가 제어 신호를 출력한다.

텍스트 명령은 사용자나 외부 시스템이 네비게이션 시스템에 제공하는 텍스트 형태의 지시사항이다. 그림 4. (c)는 자연어 처리 기술을 사용하여 명령을 이해하고 실행하는 시스템을 포함한다^[48].

3.2 OUTPUT Modal

3.2.1 WayPoint

미래 경로 지점과 원하는 경로를 예측하는 것은 높은 수준의 출력 Modal이다. WayPoint는 차량이 따라야 할 구체적인 목표 지점을 제공하며, 보다 정밀하게 경로 계획을 가능하게 한다. WayPoint는 원시 이미지 데이터에서 인식 네트워크를 훈련시키고, 관찰적 모방 학습을 사용하여 나온 WayPoint에서 제어를 예측하는 또 다른 네트워크를 학습하는 모델^[49]을 사용하거나, Motion Planner^[50]를 사용해 미래 경로를 설명하는 일

련의 경로 지점을 생성한다.

TransFuser^[40]은 자동 회귀 WayPoint 네트워크를 만들어 현재 위치와 목표 위치를 입력으로 받아, 다음 WayPoint를 예측하고, WayPoint Transfer Module^[41]은 Lane 및 Landmark 감지 기술을 사용해 WayPoint를 자동으로 추출한다. 그림 4. (d)의 VAD는 벡터화된 지도를 통해 차량이 이동할 주행 경로와 주행 방향을 보여준다^[42]. WayPoint 정보를 기반으로 차량의 제어 시스템은 필요한 조향, 가속, 감속 등의 제어 신호를 생성한다. 이때 사용하는 대표적인 알고리즘이 PID 제어로 이를 사용하여 차량이 다음 WayPoint로 이동하기 위해 필요한 방향과 속도를 계산한다^[43].

3.2.2 조향 및 속도

대부분의 종단간 학습 모델은 특정 시점에서 조향 각도와 속도를 출력으로 내보낸다^[30].

조향 제어는 주어진 경로를 정확히 따를 수 있도록 방향을 조절하는 기능이다. 조향 각도의 회귀 문제를 분류 문제로 변환하고 시간적인 종속성을 학습할 수 있도록 하는 C-LSTM 모델을 이용해 조향 각도를 예측하거나^[44], 카메라의 이미지를 입력으로 받아 조향 각도를 예측하는 PilotNet^[45] 모델이 개발되었다. 속도 제어는 차량이 안전하고 효율적으로 주행할 수 있도록 속도를 조절하는 기능이다. 이는 다양한 주행 시나리오에서 최적의 속도를 유지하는 데 필수적이다. 일부 연구에서는 속도와 조향을 동시에 예측하는 종단간 학습모델을 개발하여, 차량의 전반적인 주행 성능을 개선 하였다^[46]. DNN 구조로부터 얻은 정보를 네트워크 유형을 3개로 나눠 직접적인 제어 동작을 생성할 수 있고^[47], 주행 경로 명령을 One-Hot Encoding시켜 벡터로 넣어 이를 직진, 좌회전, 우회전에 맵핑시키는 방식도 있다^[48].

3.2.3 비용 함수

차량의 안전한 조작을 위해 많은 경로와 경로 지점이 존재하며, 비용 함수는 다양한 주행 시나리오에서 특정 주행 경로 또는 행동의 바람직함을 평가하여 주행 경로의 비용을 계산하고 가장 낮은 비용의 경로를 선택하도록 돕는다^[21].

충돌이 없는 경로를 선택하기 위해 도로 위의 객체들을 개별로 인식할 수 있다. Neural Motion Planner^[49]은 자율주행차가 각 위치의 적합성을 정의하는 Cost Volume을 사용하였다. ST-P3^[50]는 분할 맵으로 표현된 확률 필드와 교통 규칙 등의 사전 지식을 최대한 활용하여 최소 비용 경로를 선택하는 비용 함수를 사용한다.

안전성 비용 함수와 관련하여 안전 맵을 사용하여 안전 집합 내에서의 행동을 분석하고 위험한 운전 상황을 예측하고 피하는 방법을 찾는다^[51]. 예를 들어, 안전 맵은 도로의 특정 구역이 얼마나 안전한지 점수로 표시할 수 있으며, 이 점수를 바탕으로 가장 안전한 경로를 선택하고 장애물을 피할 수 있다.

최근에는 그림 4. (e)의 GenAD처럼, 주변 환경을 BEV 인코더를 통해 토큰으로 변환한 후, Future Generation에서 주변 객체의 미래 궤적 분포를 학습하고 잠재적인 미래 상황을 예측하여 경로에 대한 안전성을 평가하는 연구가 있다^[52].

IV. 종단간 학습의 모델 구조

종단간 학습 방법에는 그림 6. 에 보듯이 모방 학습(Imitation Learning, IL)과 강화 학습(Reinforcement Learning, RL)이 있으며 이 두 접근 방식은 자율주행 시스템이 환경을 이해하고 적절한 행동을 학습하는 데 중요한 역할을 한다. 모방 학습은 인간 운전자의 행동을 모방하는데 중점을 두며, 강화 학습은 주행 과정에서 보상을 기반으로 최적의 행동을 학습한다. 본 섹션에서는 이러한 학습 방법의 세부적인 내용을 다룬다.

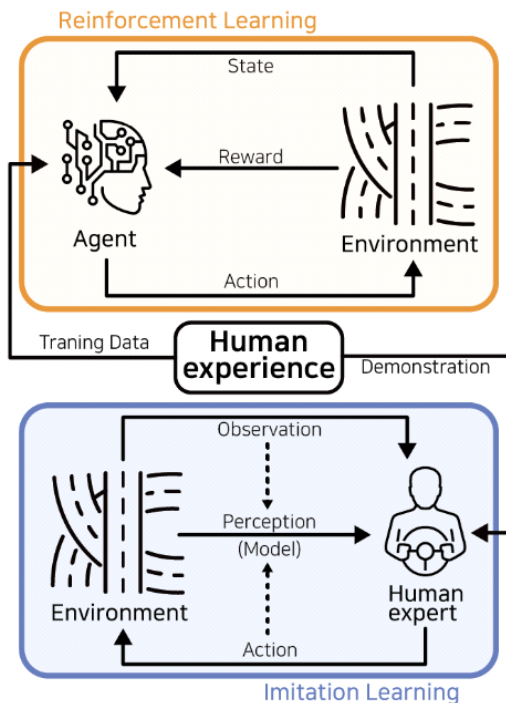


그림 6. 모방학습과 강화학습의 비교
Fig. 6. Comparison of Imitation Learning and Reinforcement Learning

4.1 모방 학습

모방 학습^[20,53]은 에이전트가 전문가의 행동을 관찰하고 이를 바탕으로 자율주행 시스템이 유사한 상황에서 동일한 행동을 하도록 훈련하는 방법이다. 이 접근 방식은 자율주행 시스템이 복잡한 환경에서도 인간과 유사한 주행 능력을 갖추도록 하는 데 중요한 역할을 한다. 모방 학습은 행동 복제(Behavior Cloning, BC)와 역강화 학습(Inverse Reinforcement Learning, IRL)으로 나뉜다.

4.1.1 BC

BC^[54]는 지도 학습의 한 형태로, 주어진 입력에 대해 전문가가 취한 행동을 학습하여 에이전트의 정책 π_θ 를 전문가의 정책 π_E 와의 손실을 일치시키는 것을 목표로 한다. 이를 통해 수집된 데이터셋의 지도 학습을 통해 계획 손실을 최소화 할 수 있어, 다양한 분야의 제어 시스템을 구축하는 데 성공적으로 사용되었다^[55].

$$\min_{\theta} E\|\pi_{\theta} - \pi_E\| \quad (1)$$

BC는 전문가의 행동이 관측에 의해 완전히 설명될 수 있다고 가정하며, 훈련 데이터 셋을 기반으로 입력에서 출력으로 맵핑하는 모델을 학습한다. 초기의 자율주행을 위한 행동 복제는 카메라 입력으로부터 제어 신호를 생성하기 위해 종단간 학습 신경망을 사용했다^[56]. 이후 다중 센서 입력^[57], Waypoint를 기반으로 제어 신호로 맵핑하는 정책^[40]과 같은 개선모델들이 개발되어 왔다. 하지만 BC는 훈련 받은 데이터 분포와 실제 환경에서 접하게 되는 데이터의 분포가 달라 모델의 성능이 저하되는 공 변량 변화(Convariate Shift)문제가 있다^[55]. 이 문제를 해결하기 위해 전문가의 정책을 반복적으로 에이전트의 정책과 결합하여 새로운 데이터셋을 생성하는 DAgger^[58]를 도입하였다. 또한 모델이 특정 행동의 원인과 결과를 혼동하는 인과 혼동(Causal Confusion)문제도 있다. 이 문제는 데이터 증강(Data Augmentation), 앙상블 학습(Ensemble Learning)과 같은 방법^[24]과 모방 정책을 객체 인식 방식으로 정규화하는 방식^[59]으로 개선되었다.

4.1.2 IRL

IRL^[60]는 에이전트의 행동 데이터를 기반으로 그 행동을 유도하는 보상 함수 r_θ 를 추정하는 것을 목표로 한다.

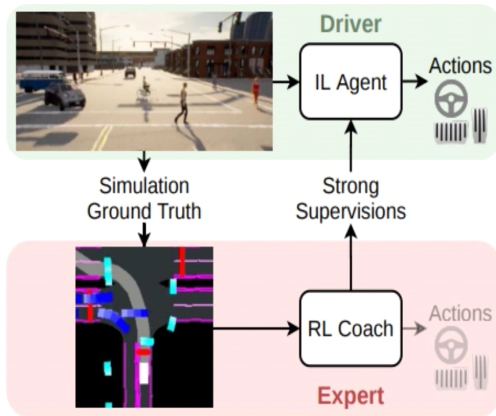
$$\max_{\theta} (E_{\pi_E}[G_t|r_{\theta}] - E_{\pi}[G_t|r_{\theta}]) \quad (2)$$

초기 연구에서는 최대 마진 접근(maximum margin approaches)방식을 통해 최적의 보상을 식별했다^[61]. 이후 전문가의 행동을 설명할 수 있는 보상 모호성을 제거할 수 있는 최대 엔트로피를 도입하였다^[62]. 이 방법은 불확실성을 최대화하면서 관찰된 행동이 가장 높은 확률을 가지도록 하는 방식이다. 자율주행 자동차의 계획에 IRL을 활용하면 인간 운전자의 반응을 예측하고 최적의 행동을 결정할 수 있다^[63]. 또한, 비용함수를 자동으로 학습하는 IRL 모듈을 만들어 후보 경로 (Trajectory)를 평가하여 최종 경로를 결정할 수 있다^[64]. 하지만, 단점으로는 IRL 알고리즘을 실행하는 데 비용이 많이 들고, 계산 요구량이 많으며, 훈련 중에 불안정하고 작은 데이터셋에서 수렴하는 데 시간이 더 오래 걸릴 수 있다는 단점이 있다.

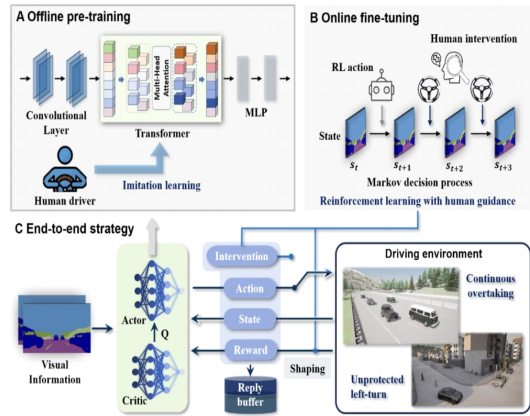
4.2 강화 학습

강화 학습은 에이전트가 환경과 상호작용하면서 최적의 행동을 학습하는 기계 학습 방법이다. 강화학습은 선택 가능한 선택들 중 보상을 최대화하는 것을 목표로 하며^[65], 네트워크는 행동에 따라 보상이나 패널티를 받아 운전 결정을 내린다. 강화 학습은 모방 학습보다 데이터 효율성이 낮지만, CARLA와 같은 일부 시뮬레이터에서는 쉽게 구현할 수 있으며, 에이전트는 더 다양한 시나리오를 탐색할 수 있다^[66]. 심층 신경망 훈련의 발전으로 고차원 감각입력에서 직접 학습하는 Deep Q Network(DQN)^[67]을 중단간 학습 모델에 적용하여 자

율주행에서 큰 진전을 이루었다^[68]. 또한, Actor-Critic 모델 기반인 Deep Deterministic Policy Gradients(DDPG)^[69], Soft Actor Critic(SAC)^[69]를 사용해 연속적인 행동 공간에서도 작동할 수 있고 가능한 랜덤하게 행동하면서 작업을 성공적으로 수행할 수 있으며, 정책을 효율적으로 수행하기 위한 Proximal Policy Optimization(PPO)^[70]도 자율주행에 적용되었다. 강화 학습은 네트워크가 시뮬레이터 정보를 사용할 수 있는 경우에도 효과적으로 적용할 수 있는데^[71], 그림 7. (a)에서 Roach^[72](Reinforcement Learning Coach)는 IL에이전트에게 감독을 제공하며 IL에이전트는 운전 환경을 관찰, 행동을 수행하여 규칙 기반인 CALRA Autopilot과 비교하여 운전 점수 지표가 23% 향상되었고, 그림 7. (b)의 PTA-RLHG^[73]는 사전 훈련된 트랜스포머 모듈과 인간 가이드 강화 학습을 통해 기존의 Hug-RL^[74]과 비교하여 성공률이 14.7% 향상되었고 충돌률이 24.7% 감소하였으며 이동 거리 지표가 약 3.7미터 증가하였다. 강화 학습을 실제 세계와 결합하는 것은 자율주행 기술의 상용화에 필수적이지만, 상당한 어려움을 수반한다^[75]. 시뮬레이션 환경에서 높은 성능을 발휘한 모델이라도, 실제 도로 환경에서는 예기치 못한 변수와 복잡한 상호작용이 발생할 수 있다. 안전성과 신뢰성을 확보하기 위해 실제 주행 테스트와 시뮬레이션을 반복적으로 결합하여 모델을 개선해야 한다. 이 과정은 많은 시간비용이 소모되지만, 자율주행을 현실 세계에 구현하기 위해서는 필수적이다.



(a) Roach^[72]



(b) Pre-trained Transformer-Enabled^[73]

그림 7. 에이전트를 최적으로 훈련시키기 위한 RL기반 학습 방법
Fig. 7. RL-based Learning Method for Optimally Training an Agent

4.3 종단간 학습의 기타 이슈 및 구조

본 섹션에서는 종단간 학습의 기타 이슈로, 종단간 학습을 통해 유발되는 작업 간 충돌 문제 및 해석 가능성 향상을 위한 연구를 설명한다.

우선, 종단간 학습 모델의 다중 작업 간의 충돌이슈를 위해, 다중 작업 3D 인지 모델 HENet^[76]이 연구되었다. 이는 단기적인 프레임에는 큰 이미지의 인코더를 사용하고 장기적인 프레임에는 작은 이미지의 인코더를 사용하여, 각 프레임의 데이터를 독립적인 인코더와 디코더로 처리하여 각 작업간의 충돌을 최소화하고 분업화 함으로서 정보를 효율적으로 결합하게 하여 해결하였다.

반면, 해석 가능성 향상 이슈를 위해 LLM(Large Language Model) 및 VAE(Variational Auto-Encoder) 기반으로 연구되었다. 특히, LLM 기반 연구에서는 다중 모달 대형 언어 모델(MLLMs)^[77]을 통해 영상 데이터 뿐 아니라 텍스트 데이터를 입력으로 받아, 결합하여 해석하는 텍스트 입력을 통해 결과 해석력을 추가하였다. DriveGPT4^[78]는 MLLMs를 자율주행에 응용한 종단간 학습 모델로서 다중 프레임 비디오를 입력으로 받아, 차량의 행동을 분석하고 텍스트 질의를 처리한다. 이에 사용자가 제기하는 다양한 질문에 효과적으로 답변하여 해석력을 추가하였다. 반면, VAE^[79]에 기반하여 자율주행의 해석 가능성을 높이기 위한 연구^[80]로는 영상입력에 대한 VAE 추출 Context Vector로 부터 인식 결과에 대한 해석 내용을 NCP(Neural Circuit Policy)로 도출하여 특정 변수가 전체 시스템의 동작에 어떻게 영향을 미치는지에 대한 통찰을 제공하였다.

V. 종단간 학습의 성능 평가

자율주행의 최종 성능을 평가하기 위한 방법으로 오프라인(개방형) 루프 평가와 온라인(폐쇄형) 루프 평가로 구분할 수 있으며, 두 평가 지표 중 하나만을 사용하거나 모두를 사용하는 경우도 있다^[81]. 본 장에서는 종단간 학습을 위한 성능평가 지표들을 설명하고 이에 기반한 종단간 학습의 성능을 제시하여 그 가능성 및 한계에 대해 서술한다.

5.1 오프라인(개방형) 루프 평가

오프라인(개방형) 루프 평가는 실제 환경에서 실험한다는 의미로서, 현실세계의 주행 시나리오 및 환경 데이터를 사용하여 평가한다. 평가 방법은 시스템의 예측된 출력과 실제 출력 간의 오차율을 두고 평가한다. 주요 평가 지표로는 예측 궤적과 실제 궤적 간의 유클리

드 거리를 구한 후 평균을 취하고 가장 작은 값을 결정하는 MinADE(Average Displacement Error)와 예측 궤적의 마지막 점과 실제 궤적의 마지막 점간의 최소 변위 오차 MinFDE(Final Displacement Error)가 있다^[23]. 실제 궤적과 예측 궤적 간의 유클리드 거리를 구한 후, 제곱 합을 취하는 L2 error와 자율주행차가 주행 중 얼마나 충돌하였는지의 Collision rate 지표 방법이 주로 사용된다.

5.2 온라인(폐쇄형) 루프 평가

온라인(폐쇄형) 루프 평가는 가상환경 시뮬레이터 상에서 평가한다는 의미로서, 온라인 평가로 지칭하기도 한다. 다양한 가상 주행 시나리오를 포함하고 있기에 오프라인(개방형) 루프 평가의 단점인 제한된 현실적 환경 때문에 일반화가 힘들다는 점을 극복한다. 주요 평가 지표로는 경로 완료율, 위반 점수, 운전 점수 등이 있다^[20].

주요 시뮬레이터로 TORCS (The Open Racing Car Simulator)와 CARLA가 있다. TORCS(The Open Racing Car Simulator)^[82]는 멀티 에이전트 기반 자동차 시뮬레이터로서 차량 역학의 모든 기본 요소를 반영하였다. 강화 학습 에이전트, 유전 알고리즘 같은 기계 학습 AI 알고리즘을 테스트할 수 있다. 그러나 TORCS는 트랙 레이싱 시뮬레이션으로 설계되었기 때문에 보행자나 교통 규칙 등이 있는 복잡한 도시 환경 시뮬레이션에는 적합하지 않다. 반면, CARLA는 도시 환경 자율주행에 중점을 두고, 상세한 환경 시뮬레이션 및 센서 통합을 제공한다. 기상 조건 및 교통 밀도 제어와 여러 환경 인지 모드로 다양한 데이터를 생성한다. 자율주행 시뮬레이션을 평가하기 위해 CARLA 환경에서 모듈형 파이프라인, 모방 학습 기반과 강화 학습 기반 종단간 자율주행을 접목하였고, 경로 완료율을 통해 성능을 검증하였다. CARLA 시뮬레이션으로 데이터셋을 만들 수도 있다. V2X 기반 자율주행 데이터셋 DeepAccident^[83]는 CARLA를 사용하여 실제 운전사고 시나리오 데이터를 생성하였다. DeepAccident는 다양한 차량 안전사고를 포함한 V2X 기반 데이터셋으로, 현실 세계에서 수집하기 어려운 차량 충돌 사고 시나리오를 가상 세계에서 만들 수 있었다.

5.3 종단간 학습의 성능

초기 종단간 학습 연구^[56]에서는 3개의 카메라 영상을 입력으로 받아 9계층 CNN기반의 백본을 사용하여 핸들 제어만을 차량 제어 파라미터로 사용하여 실험하였다. 그 결과 600초의 주행시간동안 총 10번, 각 횟수

당 6초의 사람개입이 있었으며, 이는 90%의 자동화를 달성할 수 있었다.

GraphAD^[41]는 nuScenes 데이터셋에서 평균 충돌률 0.12%을 기록하였다.

PlanT^[37]는 CARLA 시뮬레이터의 Longest6 벤치마크에서 주행 점수(경로 완료율, 위반 점수) 81.36을 달성하며, 추론 시간 10.79ms의 성능을 보여주었다.

연구 LMDrive^[48]는 Vicuna-v1.5의 LLM모델을 통해 1Km당 약 0.15의 Vehicle Collision, 0.01의 Pedestrian Collision, 0.28의 Layout Collision 성능을 도출하였다.

VAD^[42]는 nuScenes 데이터셋에서 0.07 Collision rate, 4.5 FPS 성능을 도출하였다.

GenAD^[52]는 nuScenes 데이터셋에서 Collision rate는 0.43%의 성능을 보이며 3D 객체탐지 성능으로 0.19 mAP의 성능을 도출하였다.

위 성능 수치들의 분석을 통해 알 수 있듯이, 종단간 학습을 적용하더라도 완전 자율주행을 위한 안전성(자동화율, 충돌 수, 경로 완료율) 및 실시간성(추론시간, FPS) 확보에는 아직 더 많은 연구를 필요로 한다는 것을 알 수 있다.

또한, 해당 종단간 학습 연구들은 모듈형 기반 인지, 판단, 제어 구조들의 Header를 제거하고 Feature 출력을 Cascade 형태로 연결한 Hybrid 구조로 연구되는 것을 알 수 있다. 이런 Hybrid 구조를 통해 모듈형 구조의 장점인 중간 모듈 성능 확인 및 모듈별 사전 학습 기법을 적용하여 전체 성능의 신뢰성을 확보하는 것을 알 수 있다.

VI. 결 론

본 논문에서는 자율주행 시스템 구현을 위한 두 가지 접근법, 즉, 모듈형 구조 방식과 종단간 학습 방식의 주요 기술들을 여러 논문들을 통해 알아보고 시스템 성능 측면에서 분석 하였다.

실시간으로 주변 환경을 인지하고 판단하며 최종 목적지까지 안전하게 주행해야 하는 자율주행 시스템은 시스템 최적화 관점에서 경량화 및 여러 태스크 통합이 진행되어야 하며, 이런 관점에서 전체 혹은 부분 종단간 학습 방법 적용이 필요하겠다.

반면, 각 모듈을 통합하여 종단간 학습을 구성하기에 앞서 각 모듈의 정확도 및 실시간 성능 확보가 선행하여 통합하는 것이 종단간 학습을 안정적으로 구성하는 방법을 알 수 있었다.

추후, 이런 종단간 학습 방식과 모듈형 접근방식을

Hybrid구조로 구성하여 시스템을 구현하는 것이 성능적 안정성을 확보할 수 있는 방안이며 향후 자율주행 시스템의 고도화를 위해서는 관련 기술 최적화에 대한 지속적 노력이 필요하겠다.

References

- [1] G. Kaljavesi, et al., "CARLA-autoware-bridge: Facilitating autonomous driving research with a unified framework for simulation and module development," *arXiv preprint arXiv:2402.11239*, 2024.
(<https://doi.org/10.48550/arXiv.2402.11239>)
- [2] A. Dosovitskiy, et al., "CARLA: An open urban driving simulator," *arXiv preprint arXiv:1711.03938*, 2017.
(<https://doi.org/10.48550/arXiv.1711.03938>)
- [3] A. E. Sallab, et al., "Deep reinforcement learning framework for autonomous driving," *arXiv preprint arXiv:1704.02532*, 2017.
(<https://doi.org/10.48550/arXiv.1704.02532>)
- [4] M. Bansal, et al., "ChauffeurNet: Learning to drive by imitating the best and synthesizing the worst," *arXiv preprint arXiv:1812.03079*, 2018.
(<https://doi.org/10.48550/arXiv.1812.03079>)
- [5] J. Levinson, et al., "Towards fully autonomous driving: Systems and algorithms," in *IEEE Intell. Veh. Symp. (IV)*, pp. 163-168, 2011.
(<https://doi.org/10.1109/IVS.2011.5940562>)
- [6] D. Kwak, et al., "Rethinking real-time lane detection technology for autonomous driving," *J. KICS*, vol. 48, no. 5, pp. 589-599, 2023.
(<https://doi.org/10.7840/kics.2023.48.5.589>)
- [7] Y. Wang, et al., "A comprehensive review of modern object segmentation approaches," *Foundations and Trends in Comput. Graphics and Vision*, vol. 13, no. 2-3, pp. 111-283, Oct. 2022.
(<http://dx.doi.org/10.1561/06000000097>)
- [8] X. Weng, et al., "Monocular 3d object detection with pseudo-lidar point cloud," in *Proc. IEEE/CVF Int. Conf. Comput. Vision Wkshps.*, 2019.
(<https://doi.org/10.1109/ICCVW.2019.00114>)

- [9] Y. Zhou and T. Oncel, "Voxelnet: End-to-end learning for point cloud based 3d object detection," *IEEE/CVF Conf. CVPR*, pp. 4490-4499, 2018.
(<https://doi.org/10.1109/CVPR.2018.00472>)
- [10] A. H. Lang, S. Vora, H. Caesar, L. Zhou, J. Yang, and O. Beijbom, "Pointpillars: Fast encoders for object detection from point clouds," *IEEE/CVF Conf. CVPR*, pp. 12697-12705, 2019.
(<https://doi.org/10.48550/arXiv.1812.05784>)
- [11] S. Shi, X. Wang, and H. Li, "Pointcnn: 3d object proposal generation and detection from point cloud," *IEEE/CVF Conf. CVPR*, pp. 770-779, 2019.
(<https://doi.org/10.48550/arXiv.1812.04244>)
- [12] R. Qian, X. Lai, and X. Li, "BADet: Boundary-aware 3D object detection from point clouds," *Pattern Recognition*, vol. 125, 108524, 2022.
(<https://doi.org/10.1016/j.patcog.2022.108524>)
- [13] R. Qian, et al., "3D object detection for autonomous driving: A survey," *Pattern Recognition*, vol. 130, 108796, 2022.
(<https://doi.org/10.1016/j.patcog.2022.108796>)
- [14] S. Kuutti, et al., "A survey of the state of the art localization techniques and their potentials for autonomous vehicle applications," *IEEE Internet of Things J.*, vol. 5, no. 2, pp. 829-846, 2018.
(<https://doi.org/10.1109/JIOT.2018.2812300>)
- [15] A. Naoki, et al., "Autonomous driving based on accurate localization using multilayer LiDAR and dead reckoning," *IEEE 20th Int. Conf. Intell. Trans. Syst. (ITSC)*, 2017.
(<https://doi.org/10.1109/ITSC.2017.8317797>)
- [16] X. Li, et al., "Real-time trajectory planning for autonomous urban driving: Framework, algorithms, and verifications," *IEEE/ASME Trans. Mechatronics*, vol. 21, no. 2, pp. 740-753, 2015.
(<https://doi.org/10.1109/TMECH.2015.2493980>)
- [17] C. Urmson, et al., "Autonomous driving in urban environments: Boss and the urban challenge," *J. Field Robotics*, vol. 25, no. 8, pp. 425-466, 2008.
(<https://doi.org/10.1002/rob.20255>)
- [18] R. C. Coulter, "Implementation of the pure pursuit path tracking algorithm," *Carnegie Mellon University, The Robotics Institute*, p. 92-01, 1992.
(<https://doi.org/10.48550/arXiv.2305.20026>)
- [19] T. Sebastian, et al., "Stanley: The robot that won the DARPA Grand Challenge," *J. Field Robotics*, vol. 23, no. 9, pp. 661-692, 2006.
(<https://doi.org/10.1002/rob.20147>)
- [20] P. Wu, et al., "Trajectory-guided control prediction for end-to-end autonomous driving: A simple yet strong baseline," *arXiv preprint arXiv:2206.08129v2*, 2022.
(<https://doi.org/10.48550/arXiv.2206.08129>)
- [21] J. Hawke, et al., "Urban driving with conditional imitation learning," *arXiv preprint arXiv:1912.00177*, 2019.
(<https://doi.org/10.48550/arXiv.1912.00177>)
- [22] C. Li, et al., "End-to-end autonomous driving: Challenges and frontiers," *arXiv preprint arXiv:2306.16927*, 2023.
(<https://doi.org/10.48550/arXiv.2306.16927>)
- [23] Y. Hu, J. Yang, L. Chen, et al., "Planning-oriented autonomous driving," *IEEE/CVF Conf. CVPR*, pp. 17853-17862, 2023.
(<https://doi.org/10.48550/arXiv.2212.10156>)
- [24] K. Ishihara and V. Hautamaki, "Multi-task learning with attention for end-to-end autonomous driving," *IEEE/CVF Conf. CVPR*, pp. 2902-2911, 2021.
(<https://doi.org/10.1109/CVPRW53098.2021.00325>)
- [25] C. Wen, J. Lin, T. Darrell, D. Jayaraman, and Y. Gao, "Fighting copycat agents in behavioral cloning from observation histories," *Published at NeurIPS* p. 9, 2020.
(<https://doi.org/10.48550/arXiv.2010.14876>)
- [26] Q. Zhang, Z. Peng, and B. Zhou, "Learning to drive by watching youtube videos: Action-conditioned contrastive policy pretraining," in *ECCV 17th Eur. Conf.*, pp. 111-128, Tel Aviv, Israel, Springer, Oct. 2022.

- (https://doi.org/10.1007/978-3-031-19809-0_7)
- [27] S. Chen, et al., “VADv2: End-to-end vectorized autonomous driving via probabilistic planning,” *arXiv preprint arXiv:2402.13243*, 2024.
(<https://doi.org/10.48550/arXiv.2402.13243>)
- [28] J. Zhang, Z. Huang, and E. Ohn-Bar, “Coaching a teachable student,” *IEEE/CVF Conf. CVPR*, pp. 7805-7815, 2023.
(<https://doi.org/10.1109/CVPR52729.2023.00754>)
- [29] Y. Xiao, F. Codevilla, A. Gurram, O. Urfalioglu, and A. M. López. “Multimodal end-to-end autonomous driving,” *arXiv preprint arXiv:1906.03199*, 2019.
(<https://doi.org/10.48550/arXiv.1906.03199>)
- [30] F. Codevilla, et al., “End-to-end driving via conditional imitation learning,” in *IEEE ICRA*, pp. 1-9, 2018.
(<https://doi.org/10.1109/ICRA.2018.8460487>)
- [31] S. Hecker, D. Dai, and L. Van Gool, “End-to-end learning of driving models with surround-view cameras and route planners,” *ECCV*, pp. 449-468, 2018.
(https://doi.org/10.1007/978-3-030-01234-2_27)
- [32] V. R. Kumar, et al., “OmniDet: Surround view cameras based multi-task visual perception network for autonomous driving,” *IEEE Robotics and Automat. Lett.*, vol. 6, no. 2, pp. 2830-2837, Apr. 2021.
(<https://doi.org/10.1109/LRA.2021.3062324>)
- [33] W. Sun, et al., “SparseDrive: End-to-end autonomous driving via sparse scene representation,” *arXiv preprint arXiv:2402.19620*, 2024.
(<https://doi.org/10.48550/arXiv.2405.19620>)
- [34] C. R. Qi, H. Su, K. Mo, and L. J. Guibas. “Pointnet: Deep learning on point sets for 3d classification and segmentation,” *IEEE/CVF Conf. CVPR*, pp. 652-660, 2017.
(<https://doi.org/10.48550/arXiv.1612.00593>)
- [35] Y. Cheong, et al., “Study on point cloud based 3D object detection for autonomous driving,” *J. KICS*, vol. 49, no. 1, pp. 31-40, 2024.
(<https://doi.org/10.7840/kics.2024.49.1.31>)
- [36] Q. Li, et al., “Hdmapnet: An online hd map construction and evaluation framework,” *2022 ICRA IEEE*, 2022.
(<https://doi.org/10.1109/ICRA46639.2022.9812383>)
- [37] K. Renz, et al., “PlanT: Explainable planning transformers via object-level representations,” *arXiv preprint arXiv:2210.14222*, 2022.
(<https://doi.org/10.48550/arXiv.2210.14222>)
- [38] T. Teepe, et al., “EarlyBird: Early-fusion for multi-view tracking in the bird’s eye view,” *arXiv preprint arXiv:2310.13350*, 2023.
(<https://doi.org/10.48550/arXiv.2310.13350>)
- [39] Y. Zhou, et al., “End-to-end multi-view fusion for 3d object detection in lidar point clouds,” *arXiv preprint arXiv:1910.06528*, 2019.
(<https://doi.org/10.48550/arXiv.1910.06528>)
- [40] K. Chitta, et al., “Transfuser: Imitation with transformer-based sensor fusion for autonomous driving,” *arXiv preprint arXiv:2205.15997*, 2022.
(<https://doi.org/10.48550/arXiv.2205.15997>)
- [41] Y. Zhang, et al., “GraphAD: Interaction scene graph for end-to-end autonomous driving,” *arXiv preprint arXiv:2403.19098*, 2024.
(<https://doi.org/10.48550/arXiv.2403.19098>)
- [42] Y. Duan, et al., “Prompting multi-modal tokens to enhance end-to-end autonomous driving imitation learning with LLMs,” *arXiv preprint arXiv:2404.04869*, 2024.
(<https://doi.org/10.48550/arXiv.2404.04869>)
- [43] P. L. Houtekamer and H. L. Mitchell, “Data assimilation using an ensemble kalman filter technique,” *Monthly Weather Rev.*, vol. 126, no. 3, pp. 796-811, 1998.
([https://doi.org/10.1175/1520-0493\(1998\)126<0796:DAUAEK>2.0.CO;2](https://doi.org/10.1175/1520-0493(1998)126<0796:DAUAEK>2.0.CO;2))
- [44] R. A. Jacobs, M. I. Jordan, S. J. Nowlan, and G. E. Hinton, et al., “Adaptive mixtures of local experts,” *Neural Computation*, vol. 3, no. 1, pp. 79-87, 1991.
(<https://doi.org/10.1162/neco.1991.3.1.79>)
- [45] B. Zhou, P. Krähenbühl, and V. Koltun. “Does computer vision matter for action?,” *arXiv preprint arXiv:1905.12887*, 2019.
(<https://doi.org/10.48550/arXiv.1905.12887>)

- [46] S. Chowdhuri, T. Pankaj, and K. Zipser, "MultiNet: Multi-modal multi-task learning for autonomous driving," in *2019 IEEE Winter Conf. Appl. Comput. Vision (WACV)*, pp. 1496-1504, 2019.
(<https://doi.org/10.1109/WACV.2019.00164>)
- [47] Z. Huang, et al., "Multi-modal sensor fusion-based deep neural network for end-to-end autonomous driving with scene understanding," *arXiv preprint arXiv:2005.09202v3*, 2020.
(<https://doi.org/10.48550/arXiv.2005.09202>)
- [48] H. Shao, et al., "LMDrive: Closed-loop end-to-end driving with large language models," *IEEE/CVF Conf. CVPR*, pp. 15120-15130, 2024.
(<https://doi.org/10.48550/arXiv.2312.07488>)
- [49] G. Li, et al., "Oil: Observational imitation learning," *arXiv preprint arXiv:1803.01129*, 2018.
(<https://doi.org/10.48550/arXiv.1803.01129>)
- [50] P. d. Haan, et al., "Causal confusion in imitation learning," in *Advances in NIPS*, pp. 11693-11704, 2019.
(<https://doi.org/10.48550/arXiv.1905.11979>)
- [51] M. Aldibaja, et al., "Waypoint transfer module between autonomous driving maps based on LiDAR directional sub-images," *Sensors*, 2024.
(<https://doi.org/10.3390/s24030875>)
- [52] B. Jiang, et al., "VAD: Vectorized scene representation for efficient autonomous driving," *arXiv preprint arXiv:2303.12077v3*, 2023.
(<https://doi.org/10.48550/arXiv.2303.12077>)
- [53] P. Zhao, et al., "Design of a control system for an autonomous vehicle based on adaptive-pid," *Int. J. Advanced Robotic Syst.* vol. 9, no. 2, p. 44, 2012.
(<https://doi.org/10.5772/51314>)
- [54] H. E. Eraqi, et al., "End-to-end deep learning for steering autonomous vehicles considering temporal dependencies," *arXiv preprint arXiv:1710.03804*, 2017.
(<https://doi.org/10.48550/arXiv.1710.03804>)
- [55] M. Bojarski, C. Chen, J. Daw, and A. Degirmenci, et al., "The NVIDIA PilotNet experiments," *arXiv preprint arXiv:2010.08776*, 2020.
(<https://doi.org/10.48550/arXiv.2010.08776>)
- [56] R. Michelmore, M. Kwiatkowska, and Y. Gal, "Evaluating uncertainty quantification in end-to-end autonomous driving control," *arXiv preprint arXiv:1811.06817*, 2018.
(<https://doi.org/10.48550/arXiv.1811.06817>)
- [57] J. Pedro, et al., "End-to-end deep neural network architectures for speed and steering wheel angle prediction in autonomous driving," *Electr.*, vol. 10, no. 11, p. 1266, 2021.
(<https://doi.org/10.3390/electronics10111266>)
- [58] J. Hawke, et al., "Urban driving with conditional imitation learning," *arXiv preprint arXiv:1912.00177*, 2019.
(<https://doi.org/10.1109/ICRA40945.2020.9197408>)
- [59] W. Zeng, et al., "End-to-end interpretable neural motion planner," *IEEE/CVF Conf. CVPR*, pp. 8652-8661, 2019.
(<https://doi.org/10.1109/CVPR.2019.00886>)
- [60] S. Hu, et al., "ST-P3: End-to-end vision-based autonomous driving via spatial-temporal feature learning," in *Computer Vision-ECCV 2022: 17th Eur. Conf.*, pp. 533-549, Tel Aviv, Israel, Springer, Oct. 2022.
(https://doi.org/10.1007/978-3-031-19839-7_31)
- [61] H. Shao, et al., "Safety-enhanced autonomous driving using interpretable sensor fusion transformer," in *PMLR*, pp. 726-737, 2023.
(<https://doi.org/10.48550/arXiv.2207.14024>)
- [62] W. Zheng, et al., "GenAD: Generative end-to-end autonomous driving," *arXiv preprint arXiv:2402.11502*, 2024.
(<https://doi.org/10.48550/arXiv.2402.11502>)
- [63] Y. Duan, Q. Zhang, and R. Xu, "Prompting multi-modal tokens to enhance end-to-end autonomous driving imitation learning with LLMs," *arXiv preprint arXiv:2404.04869*, 2024.
(<https://doi.org/10.48550/arXiv.2404.04869>)

- [64] S. Levine, et al., “Learning hand-eye coordination for robotic grasping with large-scale data collection,” in *Int. Symp. Experimental Robotics (ISER)*, 2016. (<https://doi.org/10.48550/arXiv.1603.02199>)
- [65] S. Ross and D. Bagnell, “Efficient reductions for imitation learning,” in *Int. Conf. Artificial Intell. and Statistics, ser. JMLR Proc.*, vol. 9, pp. 661-668, 2010.
- [66] M. Bojarski, et al., “End to end learning for self-driving cars,” *arXiv preprint arXiv:1604.07316*, 2016. (<https://doi.org/10.48550/arXiv.1604.07316>)
- [67] D. Chen and P. Krähenbühl, “Learning from all vehicles,” *arXiv preprint arXiv:2203.11934*, 2022. (<https://doi.org/10.48550/arXiv.2203.11934>)
- [68] J. Zhang and K. Cho, “Query-efficient imitation learning for end-to-end autonomous driving,” *arXiv preprint arXiv:1605.06450*, 2016. (<https://doi.org/10.48550/arXiv.1605.06450>)
- [69] J. Park, et al., “Object-aware regularization for addressing causal confusion in imitation learning,” *35th Conf. Neural Inf. Processing Syst., NeurIPS* pp. 3029-3042, 2021. (<https://doi.org/10.48550/arXiv.2110.14118>)
- [70] A. Y. Ng and S. J. Russell, “Algorithms for inverse reinforcement learning,” in *Int. Conf. Machine Learn.*, pp. 663-670, 2000.
- [71] T. P. Lillicrap, et al., “Continuous control with deep reinforcement learning,” *arXiv preprint arXiv:1509.02971*, 2015. (<https://doi.org/10.48550/arXiv.1509.02971>)
- [72] B. D. Ziebart, et al., “Maximum entropy inverse reinforcement learning,” in *Aaai*, vol. 8, pp. 1433-1438, 2008.
- [73] D. Sadigh, et al., “Planning for autonomous cars that leverage effects on human actions,” in *Robotics: Sci. and Syst.*, vol. 2, pp. 1-9, Ann Arbor, MI, USA, 2016. (<https://doi.org/10.15607/RSS.2016.XII.029>)
- [74] Z. Huang, H. Liu, J. Wu, and C. Lv, “Conditional predictive behavior planning with inverse reinforcement learning for human-like autonomous driving,” *arXiv preprint arXiv:2212.08787v3*, 2023. (<https://doi.org/10.48550/arXiv.2212.08787>)
- [75] R. S. Sutton and A. G. Barto, “*Reinforcement learning: An introduction*,” MIT Press, 2018.
- [76] B. Jaeger and A. Geiger, “An invitation to deep reinforcement learning,” *arXiv preprint arXiv:2312.08365*, 2023. (<https://doi.org/10.48550/arXiv.2312.08365>)
- [77] P. Wolf, et al., “Learning how to drive in a real world simulation with deep q-networks,” in *2017 IEEE Intell. Veh. Symp. (IV)*, pp. 244-250, 2017. (<https://doi.org/10.1109/IVS.2017.7995727>)
- [78] L. Chen, et al., “Conditional dq-n-based motion planning with fuzzy logic for autonomous driving,” *IEEE Trans. Intell. Transport. Syst.* vol. 23, no. 4, pp. 2966-2977, 2020. (<https://doi.org/10.1109/tits.2020.3025671>)
- [79] J. Chen, S. E. Li, and M. Tomizuka, “Interpretable end-to-end urban autonomous driving with latent deep reinforcement learning,” *IEEE Trans. Intell. Transport. Syst.*, vol. 23, no. 6, pp. 5068-5078, 2021. (<https://doi.org/10.1109/TITS.2020.3046646>)
- [80] X. Zhang, Y. Jiang, Y. Lu, and X. Xu, “Receding-horizon reinforcement learning approach for kinodynamic motion planning of autonomous vehicles,” *IEEE Trans. Intell. Veh.*, vol. 7, no. 3, pp. 556-568, 2022. (<https://doi.org/10.1109/TIV.2022.3167271>)
- [81] C. Zhang, R. Guo, W. Zeng, Y. Xiong, B. Dai, R. Hu, M. Ren, and R. Urtasun, “Rethinking closed-loop training for autonomous driving,” in *ECCV*, pp 264-282, 2022. (https://doi.org/10.1007/978-3-031-19842-7_16)
- [82] Z. Zhang, et al., “End-to-end urban driving by imitating a reinforcement learning coach,” in *Proc. IEEE/CVF Int. Conf. Computer Vision*, pp. 15222-15232, 2021. (<https://doi.org/10.1109/ICCV48922.2021.01494>)
- [83] D. Hu, et al., “Pre-trained transformer-enabled Strategies with human-guided

- fine-tuning for end-to-end navigation of autonomous vehicles,” *arXiv preprint arXiv:2402.12666v1*, 2024.
(<https://doi.org/10.48550/arXiv.2402.12666>)
- [84] J. Wu, Z. Huang, et al., “Prioritized experience-based reinforcement learning with human guidance for autonomous driving,” *arXiv preprint arXiv:2109.12516*, 2021.
(<https://doi.org/10.48550/arXiv.2109.12516>)
- [85] D. Chen, et al., “Learning to drive from a world on rails,” *arXiv preprint arXiv:2105.00636*, 2021.
(<https://doi.org/10.48550/arXiv.2105.00636>)
- [86] Z. Xia, et al., “HENet: Hybrid encoding for end-to-end multi-task 3d perception from multi-view cameras,” *arXiv preprint arXiv:2404.02517*, 2024.
(<https://doi.org/10.48550/arXiv.2404.02517>)
- [87] Y. Wang, et al., “Exploring the reasoning abilities of multimodal large language models (mllms): A comprehensive survey on emerging trends in multimodal reasoning,” *arXiv preprint arXiv:2401.06805*, 2024.
(<https://doi.org/10.48550/arXiv.2401.06805>)
- [88] X. Zhenhua, et al., “Drivept4: Interpretable end-to-end autonomous driving via large language model,” *arXiv preprint arXiv:2310.01412*, 2023.
(<https://doi.org/10.48550/arXiv.2310.01412>)
- [89] D. P. Kingma, et al., “Auto-encoding variational bayes,” *arXiv preprint arXiv:1312.6114*, 2013.
(<https://doi.org/10.48550/arXiv.1312.6114>)
- [90] B. Anass, et al., “Exploring latent pathways: Enhancing the interpretability of autonomous driving with a variational autoencoder,” *arXiv preprint arXiv:2404.01750*, 2024.
(<https://doi.org/10.48550/arXiv.2404.01750>)
- [91] C. Chen, et al., “Deepdriving: Learning affordance for direct perception in autonomous driving,” *IEEE/CVF ICCV*, pp. 2722-2730, 2015.
(<https://doi.org/10.1109/ICCV.2015.312>)
- [92] B. Wymann, et al., “TORCS, The Open Racing Car Simulator,” <http://www.torcs.org>, 2014.

- [93] T. Wang, et al., “Deepaccident: A motion and accident prediction benchmark for v2x autonomous driving,” in *Proc AAAI Conf: Artificial Intell.*, pp. 5599-5606, 2024.
(<https://doi.org/10.1609/aaai.v38i6.28370>)

최 정 환 (Jeong-hwan Choi)



2023년 2월: 동서울대학교 전자공학과 졸업
2024년 3월~현재: 동서울대학교 전자공학과 학사과정
<관심분야> 로봇공학, 강화학습, 자율주행

박 예 찬 (Yechan Park)



2023년 2월: 동서울대학교 전자공학과 졸업
2024년 3월~현재: 동서울대학교 전자공학과 학사과정
<관심분야> 영상처리, 자율주행, 시뮬레이션, 인공지능

전 우 민 (Woomin Jun)



2023년 2월: 동서울대학교 전자공학과 졸업
2024년 3월~현재: 동서울대학교 전자공학과 학사과정
<관심분야> 딥러닝, 영상인식, 자율주행, 조감도 인식

이 성 진 (Sungjin Lee)



2011년 8월: 연세대학교 전기전자공학과 박사 졸업
2012년 9월~2016년 7월: 삼성전자 DMC연구소 책임연구원
2016년 7월~현재: 동서울대학교 전자공학과 부교수
<관심분야> 딥러닝, 영상인식, 자율주행

[ORCID:0000-0003-3159-8394]