

비전-언어-액션 (VLA) 모델: 개념, 진전, 응용및도전과제

Ranjan Sapkota^a, Yang Cao^b, Konstantinos I. Roumeliotis^c, Manoj Karkee^a

^aCornell University, Biological & Environmental Engineering, Ithaca, New York, USA

^bThe Hong Kong University of Science and Technology, Department of Computer Science and Engineering, Hong Kong

^cUniversity of the Peloponnese, Department of Informatics and Telecommunications, Greece

Abstract

비전-언어-액션 (VLA) 모델은 인공지능의 혁신적인 진전을 의미하며, 시각, 자연어 이해, 그리고 구현된 행동을 단일 계산 프레임워크 내에서 통합하는 것을 목표로 합니다. 이기초리뷰는최근비전-언어-액션모델의발전을종합적으로제시하며, 빠르게진화하는이분야의지형을구조화하는다섯가지주제축에체계적으로조직되어있습니다. 우리는 VLA 시스템의개념적기반을확립하는것부터시작하며, 크로스모달학습아키텍처에서비전언어모델 (VLM), 행동계획자, 계층적컨트롤러를긴밀히통합하는범용에이전트로진화한과정을추적합니다. 저희방법론은지난 3 년간발표된 80 개이상의 VLA 모델을포함하는엄격한문헌검토체계를채택합니다. 주요진전분야로는아키텍처혁신, 효율적인학습전략, 실시간추론가속등이있습니다. 우리는자율주행차, 의료및산업로봇공학, 정밀농업, 휴머노이드로봇공학, 증강현실등다양한응용분야를탐구합니다. 이리뷰는실시간제어, 멀티모달액션표현, 시스템확장성, 보이지않는작업에대한일반화, 윤리적배포위험등주요과제들을추가로다룹니다. 최첨단기술을바탕으로, 에이전트 AI 적응, 신체간일반화, 통합신경-기호계획등목표해결책을제안합니다. 우리는 VLA 모델, VLM, 에이전트 AI 가만나사회적으로정렬되고적응적이며범용체화에이전트를강화하는미래지향적로드맵을제시합니다. 따라서이연구는지능형실제로봇공학과인공범용지능을발전시키는데기초참고자료가될것으로기대됩니다. 프로젝트저장소는GitHub(Source Link)

Keywords: 비전-언어-액션, 액션토큰화, 인공지능, 로봇공학, 시각언어모델

1. 소개

비전-언어-액션 (VLA) 모델이개발되기전에는로봇공학과인공지능의발전이주로별도의영역에서이루어졌습니다: 이미지를획득, 해석, 인식할수있는시각시스템입니다 [55, 87, 176], 텍스트를이해하고생성할수있는언어체계들 [213, 177], 그리고이동을제어할수있는액션시스템 [60]. 이고립된시스템들은자체적으로는잘작동했지만, 함께작동하고새로운시나리오에일반화하거나실제문제의복잡성과예측불가능성에적응하는데어려움을겪었습니다 [57, 22].

그림 1 에나타난바와같이, 합성곱신경망 (CNN) 에기반한전통적인컴퓨터비전모델은객체탐지 [25, 157, 26, 24, 174] 나분류 [205, 89, 95, 73] 와같이범위가좁게정의된작업에맞추어설계되었으며, 환경이나목표가조금만달라져도대규모라벨데이터셋과번거로운재학습을필요로했습니다 [200, 77]. 이러한시각모델은그림 1 에나타난바와같이과수원의사과를식별하는등“볼수”는있었지만, 언어를이해하거나시각정보를원하는행동으로변환하는능력은부족했습니다. 대형언어모델 (LLM) 을비롯한언어모델은텍스트기반이해와생성에서큰진전을이루었지만 [28], 물리적세계를지각하거나추론하는능력없이언어처리에만제한되어있었습니다 [98]. 한편, 수작업정책이나강화학습에크게의존하던로봇공학의행동기반시스템은객체조작과같은특정동작을수행할수있었지만 [158], 세심한공학작설계가필요했고사전에설계된시나리오밖으로일반화하기어려웠습니다 [153].

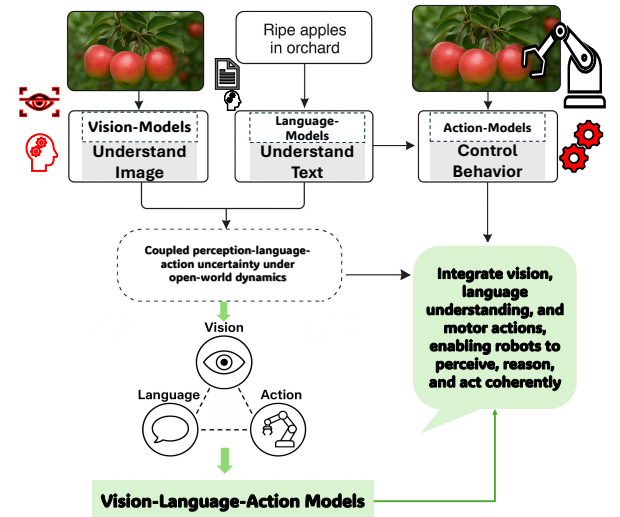


Figure 1: 고립된모달리스에서통합된비전-언어-액션모델로의진화. 통합된지각, 언어, 행동은적응적이고일반화가능한신체적지능을가능하게합니다.

시각과언어를결합하여뛰어난멀티모달이해성능을보인 VLM 의발전에도불구하고 [190, 30, 298, 37, 10, 189], 멀티모달입력에근거하여일반화된행동을생성하거나실행하는통합능력에는여전히공백이남아있었습니다 [156, 137]. 그림 1 에서더자세히시각화된바와같이, 대부분의 AI 시스템은비전-언어, 비전-행동, 언어-행동과같이한두가지양식에특화되어있

*Ranjan Sapkota

Email address: rs2672@cornell.edu (Manoj Karkee)

있고, 세양식을하나의종단간프레임워크로완전히통합하는 데 어려움을 겪었습니다. 따라서로봇은 시각적으로 “사과”를 인식하거나, “사과를 따라”와 같은 텍스트 명령을 이해하거나, 잡기와 같은 사전 정의된 운동 동작을 수행할 수는 있었지만, 이러한 능력을 유연하고 적응적인 행동으로 통합하여 실행하는 데에는 한계가 있었습니다. 그 결과 새로운 작업이나 환경에 유연하게 적응하지 못하는 취약한 파이프라인이 형성되었고, 일반화가 불안정하며 노동 집약적인 엔지니어링이 필요했습니다. 이러한 한계는 체화 AI 의 중요한 병목을 드러냈습니다. 즉, 지각하고 이해하며 행동하는 과정을 공동으로 수행할 수 있는 시스템 없이 는 지능형 자율 행동이 여전히 어려운 목표로 남아 있었습니다.

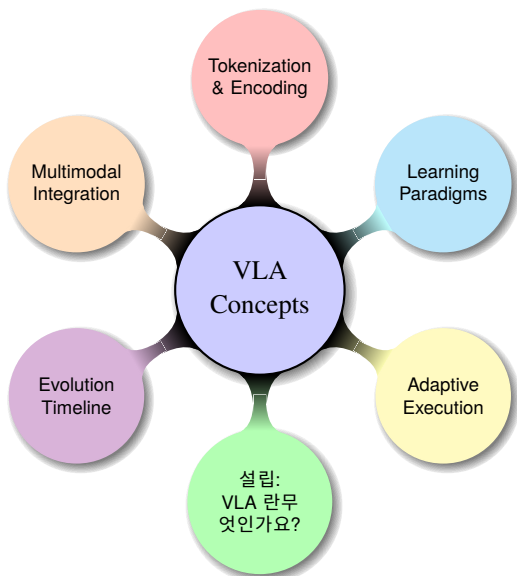


Figure 2: 핵심 VLA 개념의 마인드맵. 각색상구분된 분기는 기본 차원을 강조합니다: 정의 (기초), 역사적 진화, 멀티모달 통합, 토큰화 및 인코딩, 학습 패러다임, 그리고 구현된 환경에서 적응적 실행.

이러한 격차를 메우려는 긴급한 필요성이 VLA 모델의 등장 을 촉진시켰습니다. 2021 년에서 2022 년경 구상된 VLA 모델로, 구글 딥마인드의 Robotics Transformer 2(RT-2) 와 같은 노력으로 선구되었습니다 [299] 는 시각, 추론, 제어를 단일 프레임워크 내에서 통합하는 변혁적 아키텍처를 도입했습니다. 그림에 제시된 한계에 대한 해결책으로서 2, VLA 는 시각 입력, 언어 이해, 운동 제어 기능을 통합하여, 구현된 에이전트가 주변을 인지하고 복잡한 명령을 이해하며 적절한 행동을 동적으로 수행할 수 있게 합니다. 초기 VLA 접근법은 시각 언어 모델을 확장하여 로봇 운동 명령의 수치 또는 기호 표현 동작 토큰을 포함시켜, 쌍으로 연결된 시각, 언어, 궤적 데이터로부터 학습할 수 있게 했습니다 [156]. 이 방법론적 혁신은 로봇이 보이지 않는 물체에 일반화하고, 새로운 언어 명령을 해석하며, 구조화되지 않은 환경에서 단계적 추론을 수행하는 능력을 획기적으로 향상시켰습니다 [107].

VLA 모델은 시각, 언어, 행동을 별도의 영역으로 다루는 오랜 한계를 극복하는 통합 멀티모달 지능 개발의 변혁적 단계를 나타냅니다 [156]. 시각적, 언어적, 행동적 정보를 통합한 인터넷 규모의 데이터셋을 활용함으로써, VLA 는 로봇이 환경을 인식하고 설명할 뿐만 아니라 맥락을 따라 추론하고 복잡하고 역동적인 환경에서 적절한 행동을 수행할 수 있도록 지원합니다 [254]. 그림에 나타난 진화 2 그리고 3 고립된 시각, 언어, 행동 시스템

에서 통합 VLA 패러다임으로, 진정으로 적응적이고 일반화가 가능한 체화된 에이전트 개발로의 근본적인 전환을 포착합니다. 이 패러다임의 변혁적 잠재력을 고려할 때, 현재 문헌에 대한 포괄적이고 비판적으로 정보에 기반한 검토는 시기 적절하고 필수적입니다. 첫째, 이러한 검토는 VLA 를 이전 모델과 구별하는 기본 개념과 아키텍처 원칙을 명확히 하기 위해 필요합니다. 둘째, 이 분야의 급속한 진전과 주요 이정표를 체계적으로 설명하여 연구자와 실무자가 알고리즘 및 기술 발전의 궤적을 이해할 수 있도록 돕습니다. 셋째, 가정용 로봇 공학부터 산업 자동화 및 보조 기술에 이르기까지 VLA 가 이미 변혁적 잠재력을 보여주고 있는 다양한 실제 응용 분야를 파악하기 위해서는 심층 검토가 필수적입니다. 더 나아가, 데이터 효율성, 안전성, 일반화, 윤리적 고려 사항 등 현재의 도전 과제 를 비판적으로 검토함으로써, 광범위한 배포를 위해 해결해야 할 장벽을 식별합니다. 마지막으로, 이러한 통찰을 종합하면 AI 및 로봇 커뮤니티에 새로운 연구 방향과 실질적 고려 사항을 알리고, 협력과 혁신을 촉진합니다.

본 리뷰에서는 VLA 모델의 기본 원리를 체계적으로 분석합니다. 또한 그들의 개발 진행 상황과 기술적 도전 과제에 대해서도 논의합니다. 우리의 목표는 VLA 에 대한 현재의 이해와 응용을 통합하고, 한계를 식별하며 향후 진화 방향을 제안하는 것입니다. 검토는 주요 개념적 기초에 대한 상세한 검토로 시작됩니다 (그림 참조) 2 여기에는 VLA 모델의 정의, 그 역사적 진화, 멀티모달 통합 메커니즘, 그리고 비전, 언어, 행동을 아우르는 통합 토큰화 및 표현 전략이 포함됩니다. 이러한 개념적 설명은 VLA 가 어떻게 구조화되고 다양한 모달리즘에 걸쳐 작동하는지 이해하는 데 기반을 마련합니다.

이 설명을 바탕으로 최근 진행 상황과 훈련 효율성 전략에 대한 통합된 관점을 제시합니다 (그림 3). 여기에는 VLA 모델에서 채택되고 확장된 아키텍처 혁신뿐만 아니라, 데이터 효율적인 학습 프레임워크, 매개변수 효율적인 모델링 기법, 그리고 원래 더 넓은 기계 학습 및 로봇 공학 맥락에서 개발된 모델 가속 전략이 포함됩니다. 이러한 발전들은 VLA 시스템을 실제 적용에 맞게 확장하는 데 매우 중요합니다.

이후 VLA 시스템이 현재 직면한 한계에 대한 포괄적인 논의를 제시합니다 (그림 참조) 3 이들 중 다수는 구현된 AI 와 로봇 공학의 광범위한 도전을 반영하지만, 시각, 언어, 행동의 긴밀한 통합으로 인해 뚜렷하고 복잡한 형태 로 나타납니다. 논의할 한계에는 추론 병목 현상, 안전 문제, 높은 계산 요구량, 제한된 일반화, 윤리적 함의가 포함됩니다. 우리는 이러한 시급한 문제들을 강조할 뿐만 아니라, 이를 해결할 수 있는 잠재적 해결책에 대한 분석적 논의도 제공합니다.

이 세 그림은 함께 이 리뷰에서 제시된 텍스트 분석을 뒷받침하는 시각적 묘사를 제공합니다. 개념적 배경, 최근 혁신, 그리고 미해결 과제를 개괄함으로써 이 연구는 향후 연구를 안내하고 보다 견고하고 효율적이며 윤리적으로 기반을 둔 VLA 시스템 개발을 장려하는 것을 목표로 합니다.

그림 4 이 리뷰의 전체 구조와 논리적 흐름을 요약하고, 원고가 VLA 연구를 포괄적이고 체계적으로 분석하도록 어떻게 구성되었는지 보여줍니다. 그림에서 볼 수 있듯이, 논문은 기초 개념에서 최근 발전, 적용, 도전 과제, 미래 연구 방향으로 발전하여 섹션 간 일관된 서술을 보장합니다. 이 아키텍처를 구축하기 위해 주요 키워드 “Vision-Language-Action” 과 “Vision-Language Models”, 그리고 일반적으로 사용되는 약어 “VLA” 를 사용하여 광범위하고 엄격한 문헌 검색이 이루어졌습니다. 이 키워드는 Hugging Face, arXiv, ScienceDirect, Nature, IEEE Xplore, Wiley, Springer Nature 등 주요 학술 및 기술 저장소에서

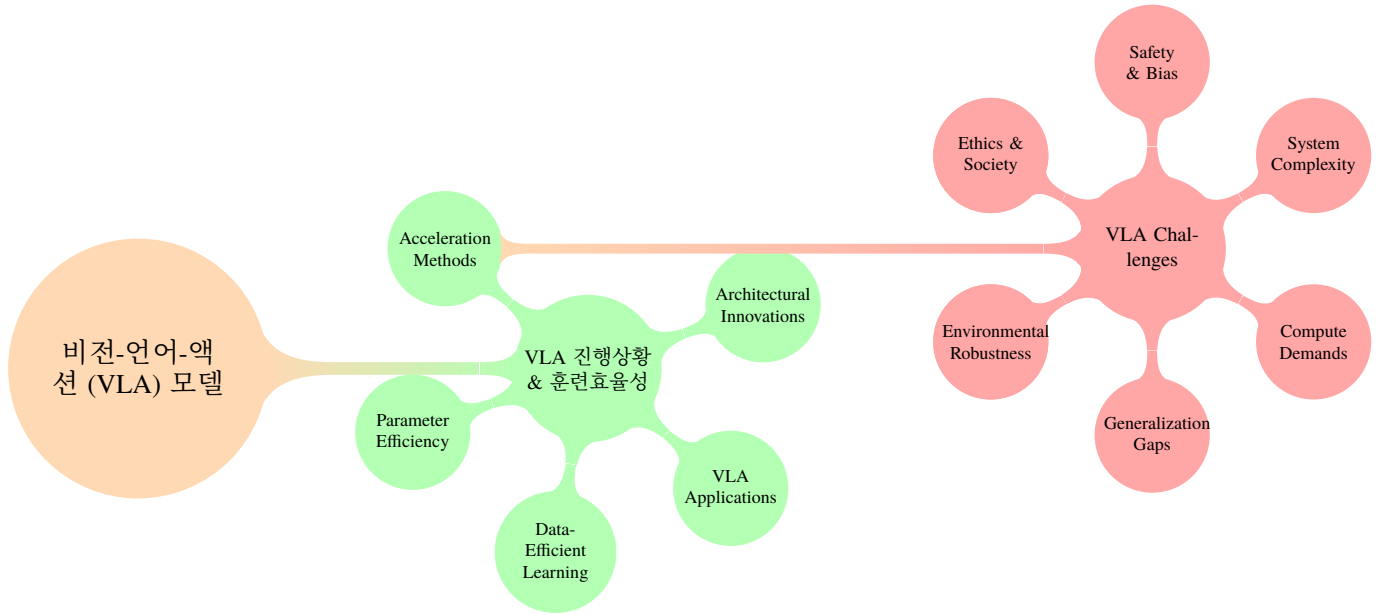


Figure 3: VLA 모델생태계를 보여주는 마인드맵: 학습효율성 (아키텍처혁신, 데이터/매개변수효율성, 가속) 의진전과추론제약, 멀티모달동작, 안전성, 일반화, 윤리와함께, 확장가능한실제배포를위해극복해야할주요과제들.

후보연구를검색하는데사용되었습니다. 결과말뭉치는수작업심사를통해관련성, 기술적깊이, 그리고이번검토범위와의일치를더욱정제했습니다. VLA 개념, 방법론적진전, 응용영역, 미해결과제이해에직접적으로기여한논문만이보존되었습니다. 이다단계필터링과정은기초연구와최첨단연구를균형있게포괄하면서도주변적이거나느슨하게관련된연구를피할수있게했습니다. 그결과, Figure 4 논문의주제구성을반영할뿐만아니라기저리뷰방법론을구현합니다.

본논문은구조적이고계층적인조직을따르며, VLA 모델의기초, 진화및합의를체계적으로발전시키며, 이는그림에요약되어있습니다 4. 발표는체화된지능과 VLA 의등장을동기부여하는소개로시작되며, 이어서멀티모달통합, 토큰화, 학습패러다임, 실시간제어등핵심원칙을확립하는개념섹션이어어집니다. 이러한토대를바탕으로진행섹션은아키텍처혁신, 교육및효율성향상, 그리고매개변수효율적인가속전략을살펴봅니다. 응용섹션에서는이러한발전을휴머노이드로봇, 자율주행차, 의료, 농업, 산업시스템, 인터랙티브 AR 내비게이션등실제영역에기반을두고있습니다. 이어서실시간추론, 안전성, 일반화, 시스템통합, 윤리적고려사항등주요과제에대한집중분석을진행합니다. 논문은지속적인학습, 확장성, 해석가능성, 내체적지능분야에서교차연구방향을정제한미래로드맵으로마무리됩니다.

2. 비전-언어-액션모델의개념

VLA 모델은시각적입력을공동으로처리하고, 자연어명령을해석하며, 동적환경에서작동하는물리적으로봇하드웨어에구현할수있는실행가능한동작표현을생성하는지능형시스템의한군을나타냅니다. 기술적으로 VLA 는비전인코더 (예: CNN, ViT), 언어모델 (예: LLM, 트랜스포머), 정책모델또는플래너를결합하여작업조건제어를달성합니다. 이모델은일반적으로교차주의, 연결임베딩, 토큰통합과같은시각언어모델에서확립된멀티모달융합기법을기반으로하며, 감각관찰,

언어지시, 액션표현을정렬하도록확장합니다.

전통적인가시모터파이프라인과달리, VLA 는의미론적그라운드 (semantic grounding) 를지원합니다 [154], 문맥인식추론을가능하게합니다 [228], affordance detection [79], 그리고시간계획 [141]. 일반적인 VLA 모델은카메라또는센서데이터를통해환경을관찰하고, 언어로표현된목표 (예: ”빨간사과를집어라”) 를해석한뒤 (그림 5 참조), 자동화된시스템이실행할수있는저수준또는고수준동작시퀀스를출력합니다. 최근발전은모방학습, 강화학습, 또는검색증강모듈을통합하여표본효율성과일반화를향상시키고있습니다. 본리뷰는 VLA 모델이기초융합아키텍처에서로봇공학, 내비게이션, 인간-AI 협업전반에걸쳐실제세계에배포할수있는범용에이전트로어떻게진화했는지살펴봅니다.

VLA 모델은시각적지각, 언어이해, 신체액션생성을하나의프레임워크로통합하는멀티모달인공지능시스템입니다. 이모델은로봇이나 AI 에이전트가감각입력 (예: 이미지, 텍스트) 을해석하고, 맥락적의미를이해하며, 실제환경에서자율적으로작업을수행할수있게합니다. 이는고립된하위시스템이아닌종단간학습과행동을통해이루어집니다. 개념적으로그림 5에나타난바와같이, VLA 모델은이전로봇및 AI 시스템의능력을제한했던시각인식, 언어이해, 운동실행간의역사적간극을연결합니다.

2.1. 발전과연대표

2022 년부터 2025 년까지 VLA 모델의급속한발전은세가지뚜렷한진화단계কে보여줍니다:

1. **기초통합 (2022–2023).** 초기 VLA 는멀티모달융합구조를통해기본적인시각-운동조정을확립했습니다. CLI-Port [202] 는 CLIP 임베딩과모션프리미티브를결합했고, Gato [181] 는 604 개과제에서범용역량을보였으며, RT-1 [19] 은대규모모방학습을통해높은조작성공률을달성했습니다. VIMA [112] 는트랜스포머기반계획기를통해시간적추론을도입했으며, 2023 년에는 RT-2 [299]

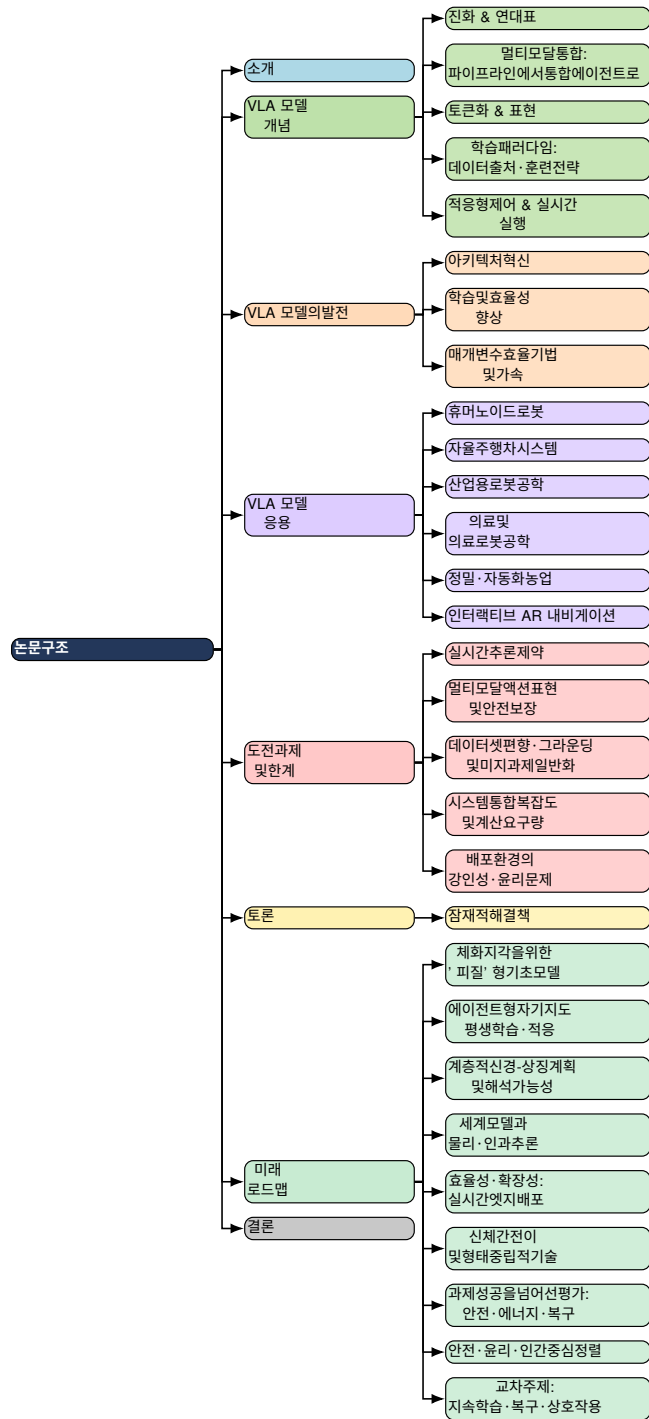


Figure 4: 이 논문의 전체 구성과 논리적 흐름. 서론과 VLA 개념에서 출발해 최근 진전, 응용, 도전과제, 논의 및 미래 로드맵으로 이어지는 구조를 요약한다.

가시각적 사고 연쇄 추론을 가능하게 하고 Diffusion Policy [40] 가 확산 과정을 통한 고급 확률적 행동 예측을 발전시켰습니다. 이러한 기초 모델은 저수준 제어를 다루었지만, 복잡한 작업을 재사용 가능한 미기반 하위 행동으로 분해하고 새로운 맥락에서 재결합하는 구성적 추론은 여전히 제한적이었으며, 이는 이후의 유연성 중심 연구를 촉진했습니다 [287, 100].

2. 전문화와 체화된 추론 (2024). 2 세대 VLA 는 도메인 특화 유도 편향을 도입했습니다. Deer-VLA [269] 는 검색 보조

학습을 통해 few-shot 적응을 개선했고, Uni-NaVid [280] 는 3D 장면 그래프 통합을 통해 내비게이션을 최적화했으며, ReVLA [48] 는 메모리 효율적인 가역 아키텍처를 도입했습니다. Occllama [239] 는 물리 기반 주의 메커니즘으로 부분 관측 가능성을 다루었습니다. 동시에 Arai 등 [5] 은 객체 중심 위험 해소를 통해 구성적 이해를 향상시켰고, Zhou 등 [293] 은 멀티모달 센서 융합을 통해 VLA 의 적용 범위를 자율주행으로 확장했습니다. 이러한 발전은 새로운 벤치마킹 방법론의 필요성을 부각했습니다 [254].

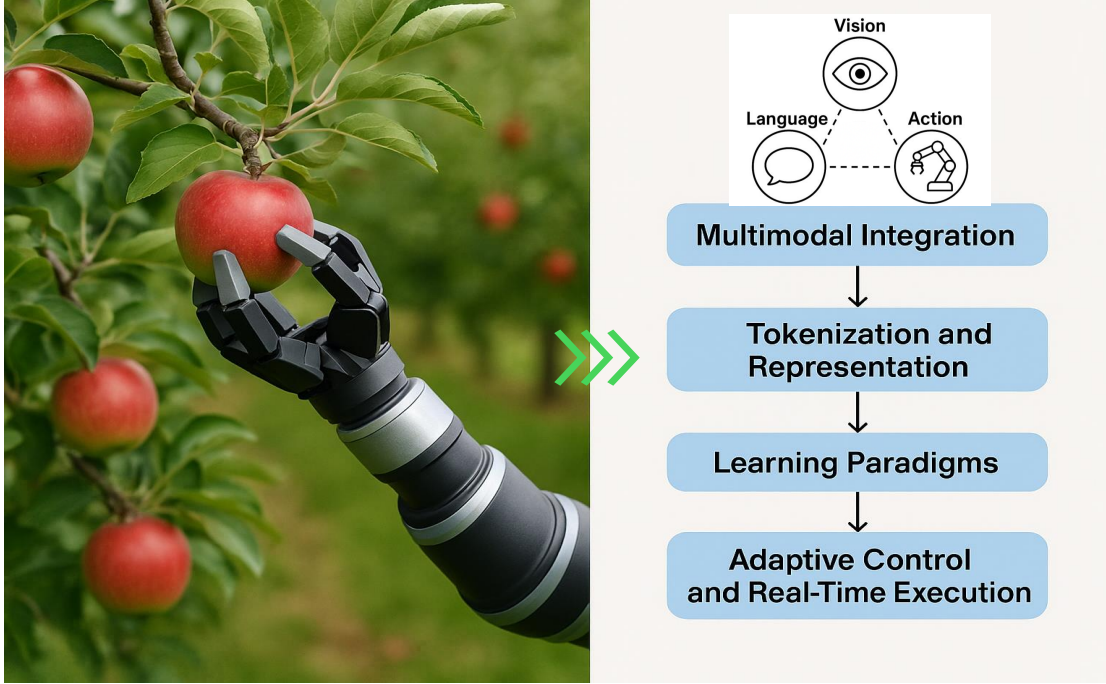


Figure 5: VLA 모델의 기초 개념 (사과따기 시나리오에서) 이 그림은 로봇팔이 과수원에서 자율적으로 익은 사과를 따는 모습을 보여주며, VLA 모델의 안내를 받는다. 오른쪽에는 VLA 모델의 네 가지 주요 단계인 멀티모달 통합, 토큰화 및 표현, 학습 패러다임, 적응 제어 및 실시간 실행을 설명한 플로우차트가 있습니다.

3. **일반화 및 안전 중요 배포 (2025).** 최신 시스템은 강인성과 인간 간 정렬을 우선시합니다. SafeVLA [274] 는 위험 인식의 사결정을 위해 형식적 검증을 통합했고, Humanoid-VLA [53] 는 계층적 VLA 를 통한 전신 제어를 입증했으며, Edge VLA [20] 는 임베디드 배포를 위한 계산 효율성을 최적화했습니다. CogACT [131] 는 인과 추론을 위해 신경-기호 (neuro-symbolic) 논리를 결합했습니다. 또한 Chain-of-Affordance [129] 의 어포던스 (행동 유도성) 체이닝과 GR00T N1 [14] 의 시뮬레이션-실제 (Sim-to-Real) 전이는 신체 간 일반화 문제를 다루며, ShowUI [139] 는 자연어 그라운드팅을 통해 VLA 와 인간 개입 (Human-in-the-loop) 인터페이스를 연결합니다.

그림 6 은 2022 년부터 2025 년까지 개발된 45 개의 VLA 모델의 진화 과정을 종합적으로 보여줍니다. CLIPort [202], Gato [181], RT-1 [19], VIMA [112] 와 같은 초기 VLA 시스템은 사전 학습된 비전-언어 표현을 작업 조건부 조작 및 제어 정책과 결합하여 기초를 마련했습니다. 이후 ACT [287], RT-2 [299], VoxPoser [100] 는 시각적 사고 연쇄 추론과 어포던스 그라운드팅을 통합했습니다. Diffusion Policy [40] 와 Octo [218] 는 확률 모델링과 확장 가능한 데이터 파이프라인을 도입했습니다. 2024 년에는 Deer-VLA [269], ReVLA [48], UniNaVid [280] 가 도메인 특화와 메모리 효율적 설계를 추가했고, Occllama [239] 와 ShowUI [139] 는 부분 관측성과 사용자 상호 작용을 다루었습니다. 이러한 흐름은 QUAR-VLA [54] 와 RoboMamba [143] 와 같은 로봇 공학 중심 VLA 로 이어졌습니다. 최근에는 SafeVLA [274], Humanoid-VLA [53], Mo-ManipVLA [246] 가 검증, 전신 제어, 메모리 시스템을 통합하면서 일반화와 배포 가능성을 강조합니다. GR00T N1 [14] 과 SpatialVLA [175] 는 시뮬레이션-실제 (Sim-to-Real) 전이와 공간 그라운드팅을 한층 강화했습니다. 이 연대표는 VLA 가 모둠 학습에서 범용적이고 안전한 체화된 지능으로 발전해온 과정

정을 보여줍니다.

2.2. 멀티모달 통합: 고립된 파이프라인에서 통합 에이전트로

VLA 모델의 등장에서 중심적인 진전은 멀티모달 통합, 즉 시각, 언어, 행동을 통합된 아키텍처 내에서 공동 처리할 수 있다는 점에 있습니다. 전통적인 로봇 시스템은 시각, 자연어 이해, 제어를 개별 모듈로 취급하며, 종종 수동으로 정의된 인터페이스 나 데이터 변환을 통해 연결되었습니다 [140, 21, 219]. 예를 들어, 고전적인 파이프라인 기반 프레임워크는 상징적 라벨을 출력하는 시각 모델이 필요했고, 이를 플래너가 도메인 별 수작업 엔지니어링을 통해 특정 행동에 매핑했습니다 [178, 117]. 이러한 접근법은 적응력이 부족했고, 모호하거나 보이지 않는 환경에서는 실패했으며, 미리 정의된 템플릿 이상의 지침을 일반화할 수 없었습니다.

반면, 현대 VLA 는 대규모 사전 학습 인코더와 트랜스포머 기반 아키텍처를 사용하여 다양한 양식을 종단간 융합합니다 [244]. 이 변화는 모델이 동일한 계산 공간 내에서 시각적 관찰과 언어적 지시를 해석할 수 있게 하여 유연하고 맥락 인식 있는 추론을 가능하게 합니다 [128]. 예를 들어, 그림 5에 나타난 '익은 사과 따기' 과제에서 비전 인코더는 일반적으로 비전 트랜스포머 (ViT) 또는 ConvNeXt 로, 장면을 파싱하여 관련 객체 (예: 과일, 잎사귀, 배경) 를 위치화하고 분류하며, 고정된 색상 가정이 아닌 학습된 질감, 형태, 맥락적 특징에 기반해 속성도와 관련된 시각적 단서를 추론합니다 [243]. 한편, 언어 모델은 종종 T5, GPT, 또는 BERT 의 변형으로, 명령어를 고차원 임베딩으로 인코딩합니다. 이 표현들은 교차주의 또는 공동 토큰화 방식을 통해 융합되어, 액션 정책을 이해하는 통합된 잠재 공간을 만들어냅니다 [86].

이 멀티모달 시너지는 CLIPort 에서 처음으로 효과적으로 입증되었습니다 [202] 이 명령어는 테이블탑 장면의 RGB 이미지와 자연어 명령어 (예: '빨간 칸 위에 파란 블록을 놓다') 를 입

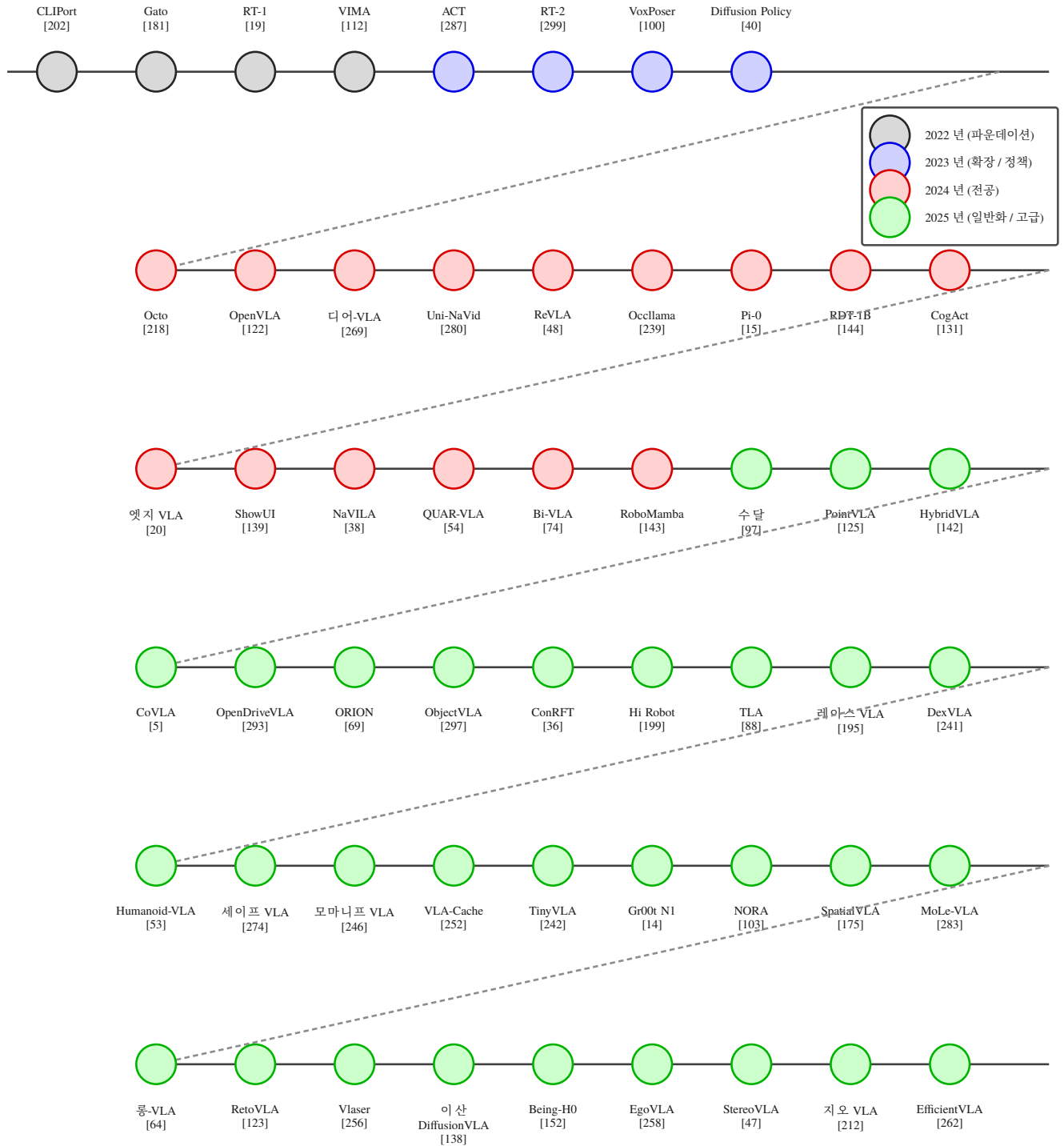


Figure 6: 2022 년부터 2025 년까지연도별로구성된 Vision-Language-Action(VLA) 모델의종합연대표로, 고대비색상코딩과주제별그룹화가적용되었습니다.

력으로받아의미론적그라운딩을위해 CLIP 으로인코딩하고, 합성곱전송디코더를통해픽셀단위의픽애플레이스액션분포를출력합니다. 언어임베딩에서각-운동정책을직접조건화함으로써 CLIPort 는명시적인어구문분석을제거하고중단간언어조건조작을가능하게합니다. 마찬가지로 VIMA [112] 이 접근법은트랜스포머인코더를사용하여객체중심시각토큰과 명령어토큰을공동처리함으로써공간추론과제전반에걸쳐 few-shot 일반화를가능하게했습니다.

최근의발전은시간적·공간적기반을포함시켜이용합을더욱촉진합니다. VoxPoser [100] 사전학습된시각언어모델과 고전모션플래너를구성하여 3D 객체선택의모호함을해결하기위해복셀수준의추론을사용하며, 특히과제별훈련데이터없이제로샷조작을달성합니다. 반면, RT-2 는 [299] 시각적언어토큰과액션표현을통합된트랜스포머내에서융합하며, 대규모인터넷시각언어말뭉치와 RT-1 데이터셋의 10 만개 이상의실제로봇시연을공동학습하여보이지않는명령에대한

제로샷일반화를가능하게합니다. 또다른주목할만한기여는 Octo [218] 입니다. 이모델은 Open X-Embodiment 데이터셋을통해다양한로봇과환경에서수집된 400 만개이상의로봇계획을학습한메모리증강트랜스포머를소개하며, 장기적의사결정을지원하고공동지각-언어-행동학습의확장성을입증합니다.

무엇보다도, VLA 는실제그라운드에서직면한문제에대한강력한해결책을제공합니다. 예를들어, Occllama [239] 는주의기반메커니즘을통해가림객체참조를처리하며, ShowUI [139] 는비전문가사용자가음성이나타이핑된입력으로에이전트를지휘할수있는자연어인터페이스를시연합니다. 이러한능력들은통합이표면수준의융합에만국한되지않기때문에가능하며; 오히려양상전반에걸친의미적, 공간적, 시간적정렬을포착합니다.

2.3. 토큰화와표현: VLA 가세상을인코딩하는방법

VLA 모델을기존시각언어아키텍처와차별화하는핵심혁신은토큰기반표현프레임워크로, 지각보다전체적인추론을가능하게합니다 [161, 286], 언어적및물리적행동공간 [136]. 트랜스포머와같은자기회귀생성모델에서영감을받은현대 VLA 는모든양식 (시각, 언어, 상태, 행동) 를통합하는이산토큰을사용해세계를인코딩하여공유임베딩공간에통합합니다 [142]. 이를통해모델은”무엇을해야하는가”(의미적추론) 뿐만아니라”어떻게해야하는가”(제어정책실행) 를완전히학습가능하고조합적으로이해할수있게합니다 [248, 150, 221].

- **접두사토큰: 인코딩문맥및명령어:** 접두사토큰은 VLA 모델의맥락적중추역할을합니다 [252, 107]. 이토큰들은환경장면 (이미지또는비디오를통해) 과함께제공되는자연어지시를압축된임베딩으로인코딩하여모델내부표현을프라이밍합니다 [17].

예를들어, 그림 7에나타난바와같이, ”빨간트레이위에초록색블록을쌓기” 와같은작업에서는복잡하게잡힌탁상이미지를 ViT 나 ConvNeXt 같은비전인코더를통해처리하고, 명령어는 LLM(예: T5 또는 LLaMA) 에의해임베드됩니다. 이토큰들은이후모델이목표와환경배치에대한초기이해를확립하는접두사토큰의연속으로변환됩니다. 이공유표현은교차모달그라운드 (cross-modal grounding) 를가능하게하여, 시스템이공간참조 (예: ”왼쪽에”, ”파란색옆”) 와객체의미론 (”녹색블록”) 을두양식모두에서해결할수있게합니다.

- **상태토큰: 로봇구성삽입:** 외부자극을인지하는것외에도, VLA 는자신의내부물리적상태를인지해야합니다 [242, 143]. 이는상태토큰을사용하여에이전트의구성관절위치, 힘-토크측정값, 그리퍼상태, 엔드이펙터자세, 심지어인근물체의위치에대한실시간정보를인코딩함으로써달성됩니다 [126]. 이토큰들은특히조작이나이동시상황인식과안전을보장하는데매우중요합니다 [211, 105].

그림 8는 VLA 모델이상태토큰을활용하여조작과내비게이션모두에서동적이고맥락인식적인의사결정을가능하게하는방식을보여줍니다. 그림 8a 에서는로봇팔이부서지기쉬운물체근처에부분적으로뻗은모습이보여집니다. 이러한상황에서상태토큰은관절각도, 그리퍼자세, 종단효과기근접성과같은실시간고유수용성정보를인코

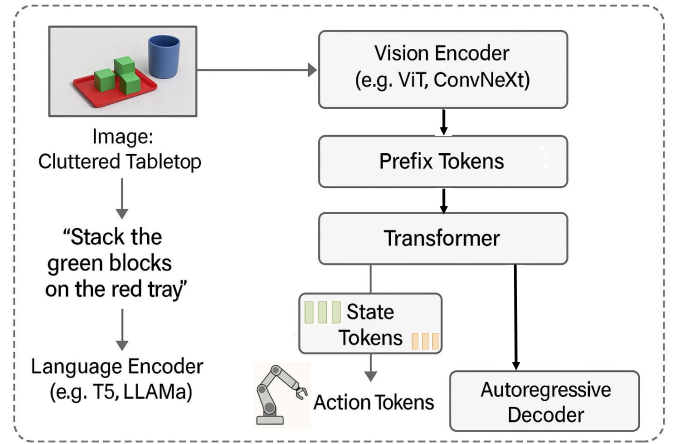


Figure 7: VLA 모델에서종단간토큰화및표현과정정보를보여주는다이아그램입니다. 시각적입력 (예: 복잡하게보이는탁상) 은비전인코더 (예: ViT) 에의해인코딩되며, 자연어명령어 (예: ” 초록블록을쌓기”) 는언어인코더 (예: T5) 로처리됩니다. 이시스템은접두사, 상태, 액션토큰을트랜스포머를통해융합하고, 모터동작을자기회귀적으로예측합니다.

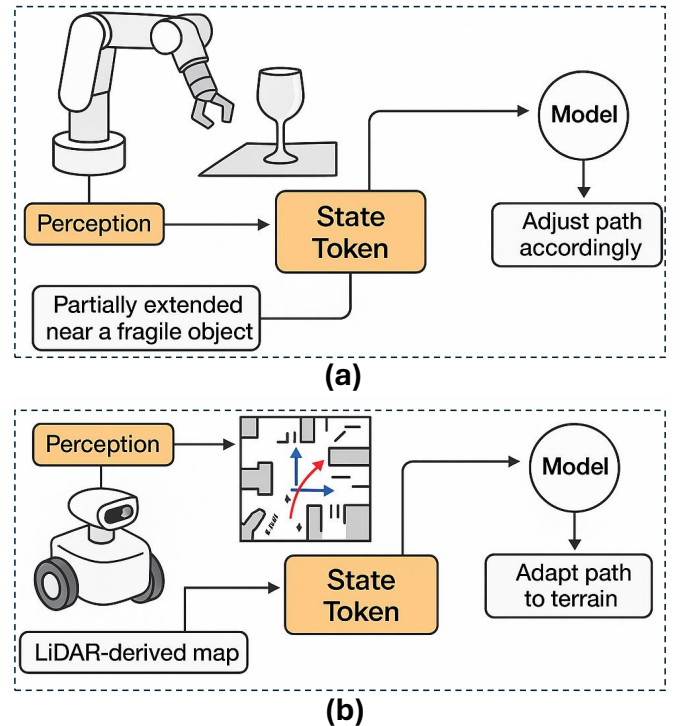


Figure 8: VLA 모델이실제상황에서접두사, 상태, 액션토큰을어떻게활용하는지보여줍니다. 로봇조작에서상태토큰은취약한물체근처에서팔확장을감지하여경로조정을가능하게합니다. 내비게이션에서는 LiDAR 와주행거리측정데이터를나타냅니다. 사과따기과제는접두사토큰이목표이해를안내하는반면, 액션토큰은목표잡기와실행을위한동작시퀀스를생성하는방식을보여줍니다.

딩함으로써중요한역할을합니다. 이토큰들은시각및언어기반접두사토큰과지속적으로융합되어트랜스포머가물리적제약을추론할수있게합니다. 따라서모델은충돌이임박했다고추론하고팔궤적을재조정하거나출력힘조절등모터명령을조정할수있습니다. 그림 8b 에묘사된이동식로봇플랫폼에서는상태토큰은주행거리, LiDAR

스캔, 관성센서데이터와같은공간적특징을캡슐화합니다. 이장비들은지형에의식하는이동과장애물회피에필수적입니다. 트랜스포머모델은이상태표현을환경및교육맥락과통합하여변화하는환경에동적으로적응하는내비게이션동작을생성합니다. 복잡한환경에서물체를잡거나올통불통한지형을자율적으로이동하는경우, 상태토큰은상황인식을위한구조화된메커니즘을제공하여자동회귀디코더가내부로봇구성과외부감각데이터를모두반영하는정밀하고맥락기반액션시퀀스를생성할수있게합니다.

- **액션토큰: 자기회귀적제어생성:** VLA 토큰파이프라인의마지막계층은액션토큰을포함합니다 [121, 122], 이모델이자기회귀적으로생성하여운동제어의다음단계를나타냅니다 [242]. 각토큰은관절각도업데이트, 토크값, 바퀴속도, 고수준이동원시등저수준제어신호에대응합니다 [81]. 추론과정에서모델은접두사와상태토큰을조건으로한단계씩토큰을해독하여 VLA 모델을언어기반정책생성기로효과적으로전환합니다 [67, 209]. 이공식은실제작동시스템과의원활한통합을가능하게하며, 가변길이의동작시퀀스를지원합니다 [11, 99], 그리고강화또는모방학습프레임워크를통한모델미세조정을가능하게합니다 [285]. 특히 RT-2 같은모델이있습니다 [299] 그리고 PaLM-E [58] 이설계를예로들수있는데, 지각, 지시, 구현이통합된토큰스트림으로통합된것입니다.

예를들어, 그림 9에묘사된사과따기과제에서모델은과수원이이미지와텍스트명령어를포함하는접두사토큰을받을수있습니다. 상태토큰은로봇의현재팔자세와그리퍼가열려있는지닫혀있는지설명합니다. 그후액션토큰이단계별로예측되어로봇팔을사과쪽으로유도하고, 그리퍼방향을조정하며, 적절한힘으로잡기를실행합니다. 이접근법의장점은전통적으로텍스트생성에사용되던트랜스포머가문장을생성하는것과유사한방식으로물리적동작의연속을생성할수있게해준다는점입니다. 다만여기서문장은움직임입니다.

로봇공학에서 VLA 패러다임을구체화하기위해, 그림 9는시각, 언어, 고유수용감각상태에특화된멀티모달정보가어떻게인코딩되고융합되어실행가능한액션시퀀스로변환되는지를보여주는구조화된파이프라인을제시합니다. 이중단간루프는로봇이“푸른잎근처의익은사과를따라”와같은복잡한작업을해석하고정밀하며상황민감적인조작을수행할수있게합니다. 시스템은먼저 **멀티모달입력획득** 단계에서시작하며, 여기서시각적관찰(예: RGB-D 프레임), 자연어명령, 실시간로봇상태정보(예: 관절각도또는속도) 라는세가지데이터스트림이수집됩니다. 이들은사전학습된모델을사용하여각각이산임베딩으로토큰화됩니다 [51, 282]. 그림 9에나타난바와같이, 이미지는 Vision Transformer(ViT) 백본을거쳐 비전토큰을생성하고, 명령어는 BERT 나 T5 와같은언어모델로파싱되어 언어토큰을생성하며, 상태입력은경량 MLP 인코더를통해압축된 상태토큰으로변환됩니다.

이토큰들은교차모달주의메커니즘을통해융합되며, 모델은객체의미, 공간배치, 물리적제약을공동으로추론합니다 [76]. 이렇게융합된표현은의사결정을위한맥락적기반을형성합니다 [96, 149]. 그림 9에서이는 **멀티모달퓨전 (multimodal fusion)** 단계로표시됩니다. 융합임베딩은일반적으로일련의액션토큰을생성하는자기회귀디코더로전달됩니다.

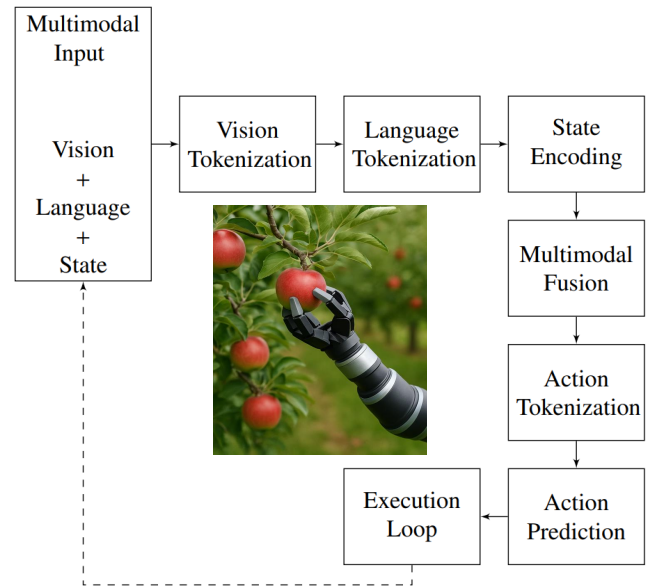


Figure 9: VLA 가세상을어떻게인코딩하는과정을보여주는과정입니다. VLA 는시각, 언어, 센서입력을토큰으로변환하고, 교차주의를통해융합하며, 트랜스포머를통한액션시퀀스를예측하고, 실시간피드백으로작업을실행하여로봇이장면을해석하고, 지시를따르며, 동적으로행동을적응할수있게합니다.

이토큰은관절변위, 그리퍼힘조절, 또는고수준모션프리미티브(예: “파지자세로이동”, “손목회전”) 를나타낼수있습니다. 예측된액션토큰은이후저수준제어명령어로변환되고, 외부하드웨어의존실행루프를통해실행됩니다. 이루프는로봇컨트롤러와상호작용하여다음 VLA 추론단계에필요한갱신된상태관찰을되돌려줌으로써지각-행동사이클을닫습니다. 이러한폐루프메커니즘덕분에모델은교란, 물체이동, 가림에실시간으로동적으로적응할수있습니다 [276, 155, 251].

명확한구현세부사항을제공하기위해, 알고리즘 1은 VLA 토큰화과정을형식화합니다. RGB-D 프레임 I , 자연어명령 T , 관절각벡터 θ 가주어지면, 알고리즘은순차적으로실행가능한액션토큰집합을생성합니다. 이미지 I 는 ViT 를통해처리되어 400 개의시각토큰집합 V 로변환됩니다. 병렬적으로명령 T 는 BERT 모델에의해 12 개의의미언어토큰시퀀스 L 로인코딩됩니다. 동시에로봇상태 θ 는관절각, 엔드이펙터자세, 고유수용신호를포함하며, 다층퍼셉트론 (MLP) 에의해 64 차원상태임베딩 S 로인코딩됩니다. 이토큰들은교차주의모듈을통해융합되어공유 512 차원표현 F 를생성하며, 이표현은실제행동에필요한의미, 의도, 상황인식을포착합니다. 마지막으로 FAST 와같은정책디코더 [171] 는융합특징을 50 개의개별액션토큰으로매핑하며, 이토큰들은모터명령 $\tau_{1:N}$ 으로디코딩됩니다.

디코딩과정은트랜스포머기반아키텍처를사용하여구현되며, 이는코드스니펫에서확인할수있습니다. 액션예측코드. A 트랜스포머디코더는 12 개의레이어, 모델차원 512, 그리고 8 개의어텐션헤드로구현되어있습니다. 융합된멀티모달토큰은컨텍스트로제공되며, 디코더는액션토큰을한단계씩자동회귀하여예측합니다. 각예측토큰은전체멀티모달컨텍스트와이전에생성된모든행동을조건으로한다음제어결정을나타냅니다. 결과적으로생성된액션토큰시퀀스는실행을위한

연속모터명령계적으로 디토큰화됩니다. 이구현은 LLM 에서 텍스트 생성이 작동하는 방식을 반영하지만, 여기서 ”문장” 은 자연어 생성 기법을 물리적 행동 합성을 위해 새롭게 해석한 움직임계적입니다.

함께, 그림 9, 알고리즘 1, 그리고 의사코드는 VLA 가 일관되고 해석 가능한 토큰 공간 내에서 시각, 지시, 구현을 통합하는 방식을 보여줍니다. 이러한 모듈화 덕분에 프레임워크는 작업과 로봇 형태에 걸쳐 일반화할 수 있어, 사과 따기, 가정 작업, 모바일 내비게이션과 같은 실제 응용 분야에 신속히 적용할 수 있습니다. 특히, 토큰화 단계의 명확성과 분리가 성능 덕분에 아키텍처가 확장 가능해져 VLA 시스템에서 토큰 학습, 계층적 계획, 기호적 기반에 대한 추가 연구가 가능해집니다.

Algorithm 1 VLA 토큰화 파이프라인

```

1: Input: RGB-D frame  $I$ , text command  $T$ , joint angles  $\theta$ 
2:  $V \leftarrow \text{ViT}(I)$  ▷ 400 vision tokens
3:  $L \leftarrow \text{BERT}(T)$  ▷ 12 language tokens
4:  $S \leftarrow \text{MLP}(\theta)$  ▷ 64-dim state encoding
5:  $F \leftarrow \text{CrossAttention}(V, L, S)$  ▷ 512-dim fused token
6:  $A \leftarrow \text{FAST}(F)$  ▷ 50 action tokens
7: Output: Motor commands  $\tau_{1:N}$ 

```

액션 예측 코드

```

# Python-like pseudocode
def predict_actions(fused_tokens):
    transformer = Transformer(num_layers=12,
                              d_model=512,
                              nhead=8)
    action_tokens = transformer.decode(fused_tokens,
                                      memory=fused_tokens)
    return detokenize(action_tokens)

```

2.4. 학습 패러다임: 데이터 소스와 교육 전략

VLA 모델 훈련은 웹의 이미지와 로봇 데이터셋의 작업 기반 정보를 통합하는 하이브리드 학습 패러다임을 필요로 합니다 [35]. 앞서 보았듯이, VLA 의 멀티모달 아키텍처는 언어 이해, 시각 인식, 운동 제어를 지원하는 다양한 형태의 데이터에 노출되어야 합니다. 이는 일반적으로 두 가지 주요 데이터 소스를 통해 이루어집니다.

첫째, 그림 10에 묘사된 바와 같이, 대규모 인터넷 기반 말뭉치가 모델의 의미론적 사전 지식 (semantic prior) 의 중추를 형성합니다. 이 데이터셋에는 이미지-캡션 쌍 (예: COCO, LAION-400M), 지시-따르기 데이터셋 (예: HowTo100M, WebVid), 시각적 질문 답변 말뭉치 (예: VQA, GQA) 가 포함됩니다. 이러한 데이터셋은 시각 및 언어 인코더의 사전 학습을 가능하게 하여 모델이 객체, 행동, 개념의 일반적인 표현을 획득하는데 도움을 줍니다 [2]. 이 단계는 CLIP 스타일의 대조 학습이나 언어 모델링 손실과 같은 대조적이거나 가스막힌 모델링 목표를 사용하여 공유 임베딩 공간 내에서 시각과 언어 양식을 정렬하는데 사용됩니다 [186, 260]. 중요한 점은, 이 단계가 VLA 예기본’ 세계에 대한 이해’ 를 제공하여 구성 일반화, 물체 그라운드, 제로샷 전송을 용이하게 한다는 것입니다 [33, 16].

하지만 의미론적 이해만으로는 물리적 작업 수행에 충분하지 않습니다 [44, 231, 137]. 따라서 두 번째 단계는 모델을 체화

된 경험에 기반시키는 데 초점을 맞춥니다 [231]. 실제로 로봇이 나고 충실도 시뮬레이터에서 수집된 로봇 궤적 데이터셋을 사용하여 언어와 시각이 어떻게 행동으로 연결되는지 모델을 가르칩니다 [67]. 여기에는 RoboNet 과 같은 데이터셋이 포함됩니다 [45], Bridge 데이터 [61], 그리고 RT-X [227] 이 기능은 자연어 명령에 따라 비디오-액션 쌍, 공동 궤적, 환경 상호 작용을 제공합니다 [159]. 시범 데이터는 운동 감각 교육, 원격 조작, 스크립트 정책 등에서 나올 수 있습니다 [115, 13]. 이 단계는 일반적으로 감독 학습 (예: 행동 복제) 을 사용합니다 [68], 강화 학습 (RL), 또는 모방 학습을 통해 자기회귀 정책 디코더가 융합된 시각 언어 상태 임베딩을 기반으로 액션 토큰을 예측하도록 훈련시키는 데 사용됩니다 [83].

최근 작품들은 점점 더 다양한 내용을 채택하고 있습니다. 다단계 또는 다중 작업 훈련 전략. 예를 들어, 모델은 종종 마스크 언어 모델링을 이용해 시각 언어 데이터셋에서 사전 학습된 후, 토큰 수준의 자기회귀 손실을 이용해 로봇 시연 데이터를 미세 조정하는 경우가 많습니다 [122, 295, 252]. 다른 이들은 커리큘럼 학습을 사용하는데, 간단한 작업 (예: 객체 밀기) 이 더 복잡한 작업 (예: 다단계 조작) 보다 먼저 진행됩니다 [288]. OpenVLA 와 같이 도메인 적응을 더욱 활용하는 방법도 있습니다 [122] 또는 합성 및 실제 배포 간의 간극을 매우기 위한 시뮬레이션-투리얼 전송 [125]. 의미적 사전과 작업 실행 데이터를 통합함으로써, 이러한 학습 패러다임은 VLA 모델이 작업, 도메인, 구현체에 걸쳐 일반화할 수 있게 하며, 이는 견고한 실제 세계에서 의 동작이 가능한 확장 가능한 명령어 추적에 이 전의 중추를 형성합니다.

공동 파인 튜닝을 통해 이 데이터셋들은 정렬됩니다 [232, 63]. 모델은 시각 및 언어적 입력을 적절한 액션 시퀀스로 매핑하는 것을 학습합니다 [175]. 이후 훈련 패러다임은 모델이 객체 활용 가능성 (예: 사과를 잡을 수 있음) 과 행동 결과 (예: 들어 올리는 데 힘과 궤적이 필요함) 을 이해하는데 도움을 줄 뿐만 아니라, 새로운 시나리오로의 일반화를 촉진합니다 [129]. 주방 조작과제로 훈련된 모델은 객체 위치 파악, 잡기, 언어 지시 사항 준수 등의 일반 원리를 배웠다면 야외 과수원에서 사과를 따는 방법을 추론할 수 있을 것입니다.

최근 아키텍처, 예를 들어 구글 딥마인드의 RT-2 (Robotics Transformer 2) [299] 이 원리를 실제로 입증했습니다. RT-2 는 액션 생성을 텍스트 생성의 한 형태로 다루며, 각 액션 토큰이 로봇의 제어 공간 내 개별 명령에 대응합니다. 이 모델은 웹 규모의 멀티모달 데이터와 수천 개의 로봇 시연을 모두 기반으로 학습되었기 때문에, 새로운 명령어를 유연하게 해석하고 새로운 객체와 작업에 대해 제로샷 일반화를 수행할 수 있는데, 이는 전통적인 제어 시스템이나 초기 멀티모달 모델에서는 거의 불가능했던 일입니다.

2.5. 적응형 제어 및 실시간 실행

VLA 의 또 다른 강점은 센서의 실시간 피드백을 활용해 실시간으로 동작을 조정하는 적응형 제어에 있습니다 [195]. 이는 과수원, 가정, 병원과 같은 역동적이고 구조화되지 않은 환경에서 특히 중요하며, 예: 바람이 사과를 움직이거나, 조명 변화, 인간의 존재 등으로 인해 과제 매개변수가 바뀔 수 있습니다. 실행 중에 상태 토큰은 실시간으로 업데이트되어 센서 입력과 공동 피드백을 반영합니다 [252]. 그 후 모델은 계획된 행동을 그에 맞게 수정할 수 있습니다. 예를 들어, 사과 따기 시나리오에서 목표 사과가 약간 이동하거나 다른 사과가 시야에 들어오면, 모델은 장면을 동적으로 재해석하고 잡는 궤적을 조정합니다. 이 능력은

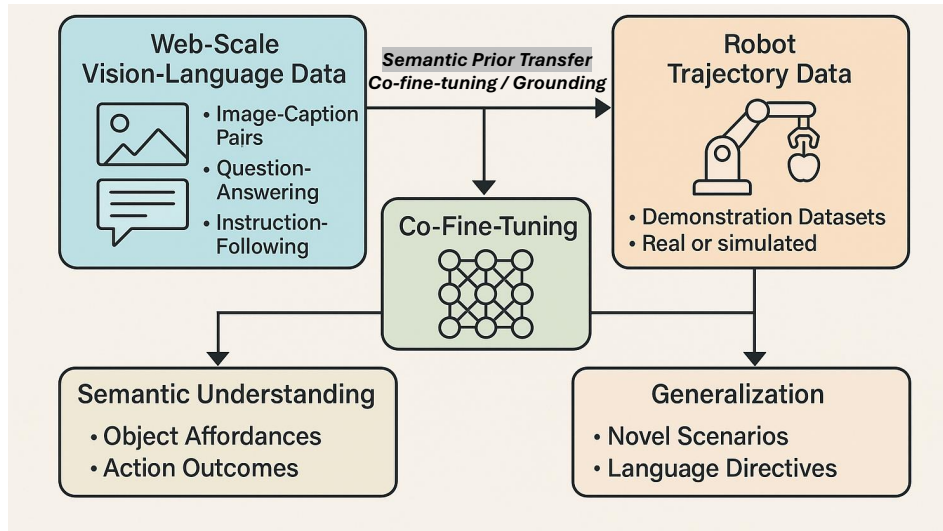


Figure 10: 학습패러다임: VLA 를위한데이터소스와훈련전략.

인간과 유사한 적응력을 모방하며, 파이프라인 기반 로봇 대비 VLA 시스템의 핵심 장점입니다.

3. 비전-언어-액션 모델의 진전

VLA 모델의 시작은 2022년 11월에 출시된 ChatGPT를 비롯한 트랜스포머 기반 LLM의 놀라운 성공에 의해 촉발되었습니다. ChatGPT는 전례 없는 의미 추론 능력을 입증했습니다 [179]. 이 돌파구는 연구자들이 언어 모델을 멀티모달 영역으로 확장하도록 영감을 주었으며, 로봇 공학을 위한 지각과 행동을 통합했습니다. 2023년까지 GPT-4는 텍스트와 이미지를 모두 처리하는 멀티모달 기능을 도입했고, 이는 이후 언어 중심의 멀티모달 기초 모델을 물리적 액션 표현과 제어 인터페이스를 통합하는 노력을 촉진했습니다 [1]. 동시에 CLIP(2022) 같은 VLM도 있습니다. [202] 그리고 Flamingo (2022) [3] 대조 학습을 통해 견고한 시각-텍스트 정렬을 확립하여 제로샷 객체 인식에 가능하게 하고 CLIP과 같은 VLM 모델의 기초를 마련했습니다. 이 모델은 대규모 라벨 데이터셋을 활용하여 이미지와 텍스트 설명을 정렬했으며, 이는 액션을 통합하는 데 중요한 전조였습니다.

중대한 발전 중 하나는 RT-1의 13만 개 시연과 같은 대규모 로봇 데이터셋의 생성으로, 시각, 언어, 행동 구성 요소의 공동 훈련에 필수적인 행동 기반 데이터를 제공했습니다 [19]. 이 데이터셋은 다양한 작업과 환경을 포착하여 모델이 일반화 가능한 행동을 학습할 수 있게 했습니다. 2023년 구글의 RT-2 로아키텍처적 돌파구가 이어졌습니다 [18] 이 모델은 시각, 언어, 액션 토큰을 통합하여 로봇 제어를 자기회귀적 수열 예측 작업으로 다루는 획기적인 VLA 모델입니다 (RT-2 블로그 (Source Link)). RT-2는 이산 코사인 변환 (DCT) 압축과 바이트-쌍인 코딩 (BPE)을 사용하여 이산화 동작을 수행하며 63% 새로운 물체에서의 성능 향상. 교차주의 트랜스포머와 같은 멀티모달 융합 기술은 Vision Transformer(ViT) 처리된 이미지 (예: 400 패치 토큰)와 언어 임베딩을 통합하여, “로봇이” 그릇 왼쪽에 있는 빨간 컵을 선택하세요”와 같은 복잡한 명령을 실행할 수 있게 합니다. 또한 UC 버클리의 Octo 모델 (2023)은 OpenX-Embodiment 데이터셋에서 80만 건의 로봇 시연을 기반으로

93M 매개변수와 확산 디코더를 활용한 오픈소스 접근법도 도입하여 연구 분야를 더욱 확장시켰습니다 [218].

3.1. VLA 모델의 아키텍처 혁신

2023년부터 2024년까지 VLA 모델은 큰 아키텍처 발전과 훈련 방법론의 정교화를 거쳤습니다. 듀얼 시스템 아키텍처는 핵심 혁신으로 부상했으며, 이는 NVIDIA의 GR00T N1(2025)에서 대표적입니다 [14] 이 프로그램은 시스템 1(저수준 제어를 위한 10ms 지연 시 빠른 Diffusion Policy)과 시스템 2(고수준 작업 분해를 위한 LLM 기반 플래너)를 결합했습니다. 이러한 분리는 전략적 계획과 실시간 실행 간의 효율적인 조정을 가능하게 하여 역동적인 환경에서 적응력을 높였습니다. 스탠포드의 OpenVLA(2024) 같은 다른 모델도 있습니다 [122], 97만 개의 실제 로봇 시연을 기반으로 학습된 7B 매개변수 오픈소스 VLA를 도입했으며, 이중비전 인코더 (DINOv2 [164] 그리고 SigLIP [271]) 그리고 Llama 2 언어 모델 [223], RT-2-X(55B)와 같은 대형 모델보다 더 좋은 성능을 보입니다 [122]. 교육 패러다임은 웹 규모의 시각 언어 데이터 (예: LAION-5B)에 대한 공동 파인 튜닝을 활용하도록 진화했습니다 [194] 로봇 궤적 데이터 (예: RT-X) [227], 의미 지식과 물리적 제약 조건을 일치시키는데 [194]. UniSim과 같은 합성 데이터 생성 도구는 견고한 학습에 필수적인 가린 객체와 같은 사진처럼 사실적인 시나리오를 만들어 데이터 부족 문제를 해결했습니다 (UniSim [264]). 파라미터 효율성은 미세 조정을 위한 저랭크 적응 (LoRA) 어댑터를 통해 향상되었습니다 [93] 이 기술은 완전한 재학습 없이도 메인 적응을 가능하게 하여 GPU 시간을 70시간 줄였습니다%. Physical Intelligence의 pi 0 모델 (2024)에서 볼 수 있는 확산 기반 정책의 도입 [15] 행동 다양성을 개선했지만 상당한 계산 자원이 필요했습니다. 이러한 발전은 VLA 기술을 민주화하여 협업을 촉진하고 혁신을 가속화했습니다.

최근 VLA 모델은 효율성, 모듈성, 견고성을 균형 있게 갖춘 세 가지 주요 아키텍처 패러다임으로 수렴했습니다: 초기 융합 모델, 이중 시스템 아키텍처, 그리고 자기 보정 프레임워크입니다. 이 각각의 혁신은 실제로 로봇 시스템에서 그라운드링, 일반화, 동작 신뢰성이라는 구체적인 도전 과제를 해결합니다.

1. 초기 융합 모델: VLA 접근법의 한 종류는 입력 단계에서 비전과 언어 표현을 융합한 후 정책 모듈로 전달하는 데 중점을

됩니다. Huang 등의 EF-VLA 모델 [96] 국제학습표현학회 (ICLR 2025) 에서 발표된 이 문서는 CLIP 이 확립한 표현적 정렬을 유지함으로써 이러한 경향을 잘 보여줍니다 [202]. EF-VLA 는 이미지-텍스트 쌍을 받아 들여 CLIP 의 동결 인코더로 인코딩하고, 동작 예측전에 트랜스포머 백본 초기에 생성된 임베딩을 융합합니다. 이 설계는 CLIP 사전 학습 과정에서 학습된 의미적 일관성을 유지하여 과적합을 줄이고 일반화를 향상시킵니다. 특히 EF-VLA 는 구성 조작 과제에서 20% 의 성능 향상을 기록했으며, 이전에 본적 없는 목표 설명에 대해서도 85% 의 성공률을 달성했습니다. 시각 언어 중추를 고정함으로써 계산 효율성을 유지하고 치명적인 망각을 방지하며, 도메인 별 훈련은 경량 정책 또는 액션 모듈에 한정되어 모델의 범용 시각의 미 표현을 희생하지 않고 작업 적응을 가능하게 합니다.

2. 이중 시스템 아키텍처: 인간 인지의 이중 과정 이론에서 영감을 받은 NVIDIA 의 GROOT N1(2025) 모델이 있습니다 [14] 두 가지 보완적인 하위 시스템을 구현합니다: 빠른 반응 모듈 (시스템 1) 과 느린 추론 계획기 (시스템 2). 시스템 1 은 10ms 지연 시간으로 동작하는 확산 기반 제어 정책을 포함하며, 중단 효과기 안정화나 적응형 그래스핑과 같은 미세하고 저수준 제어에 이상적입니다. 반면 System 2 는 작업 계획, 기술 구성, 고수준 순서 작성에 LLM 을 사용합니다. 플래너는 장기 목표 (예: "테이블 정리") 를 원자 하위 작업으로 분석하며, 저수준 컨트롤러는 실시간 실행을 보장합니다. 이 분해는 다시간 척도의 추론과 향상된 안전성을 가능하게 하며, 특히 신속한 반응과 속도가 공존해야 하는 환경에서 더욱 그렇습니다. 다단계가정 조작에 대한 벤치마크 테스트에서 GROOT N1 은 RT-1, RT-2, OpenVLA 와 같은 모놀리식 모델보다 성공률에서 17% 더 높은 성능을 보였으며, 충돌 실패율을 28% 감소시켰습니다.

3. 자기 교정 프레임워크: 세 번째 아키텍처 진화는 기존의 추론 파이프라인에 명시적인 장애 감지 및 복구 메커니즘을 추가하는 자가 수정 VLA 모델의 등장입니다. SC-VLA(2024) 는 이전의 중단 간도는 계층적 VLA 설계와 유사한 표준 빠른 추론 경로를 유지하지만, 실행 실패나 불일치가 감지될 때 선택적으로 재평가하고 복구 조치를 생성하기 위해 선택적으로 활성화되는 느린 수정 경로를 추가하기도 합니다. 이 프레임워크에서는 경량 트랜스포머를 사용해 퓨즈드 임베딩에서 직접 포즈나 동작을 예측하는 것이 기본 동작입니다. 실패가 감지되면, 예를 들어 실패한 잡기나 장애물 충돌이 발생하면, 모델은 사고 연쇄 추론을 수행하는 2 차 과정을 호출합니다 [281, 270]. 이 경로는 내부 LLM (또는 외부 전문가 시스템) 에 쿼리하여 실패 모드를 진단하고 고수준 전략을 생성합니다 [59]. 예를 들어, 로봇이 가려진 물체를 반복적으로 잘못 식별하면, LLM 은 능동적 시점 변경이나 그리퍼 재배치를 제안할 수 있습니다. 페루프 실험에서 SC-VLA 는 작업 실패율을 35

4. VLA 모델의 아키텍처 디자인 공간 VLA 모델은 풍부한 아키텍처 설계와 기능 강조점을 보여주며, 이는 중단 간 파이프라인과 모듈식 파이프라인, 계층적 또는 평면 정책 구조, 저수준 통제와 고수준 계획 간의 균형 차원을 따라 체계적으로 조직할 수 있습니다 (표 참조) 1). CLIPort 와 같은 중단 간 VLA [202], RT-1 [19], 그리고 OpenVLA [122] 원시 감각 입력을 단일 통합 네트워크를 통해 직접 모터 명령으로 처리합니다. 반면, VLATest 와 같은 컴포넌트 중심 모델은 [237] 그리고 Chain-of-Affordance [129] 시각, 언어 기반, 행동 모듈을 분리하여 개별 하위 모듈에서 목표 개선을 가능하게 합니다.

계층적 아키텍처는 전략적 사결정과 반응적 통제를 분리하여 복잡하고 장기적인 과제를 해결하기 위해 등장했습니다. 예를 들어, CogACT [131] 그리고 NaVILA [38] LLM 기반 계획

자가 하위 목표를 저수준 컨트롤러에게 발행하는 2 단계 계층 구조를 사용하여 시스템 2 의 추론과 시스템 1 실행의 강점을 결합합니다. 마찬가지로, ORION 도 [69] 장기 맥락 집계를 위한 QT-포머와 생성 계획 계획기를 통합하여 응집된 프레임워크를 만듭니다.

저수준 정책 강조는 확산 기반 컨트롤러 (예: Pi-0 [15], Dex-GraspVLA [291]) 는 부드럽고 다양한 모션 분포를 생성하는데 뛰어나지만, 종종 더 많은 계산 비용이 발생합니다. 반면, 고수준 플래너 (예: FAST Pi-0 Fast [171], CoVLA [5]) 빠른 하위 목표 생성 또는 거친 계획 예측에 집중하며, 세밀한 제어는 전문 모듈이나 전통적인 모션 플래너에 위임합니다. HybridVLA 와 같은 중단 간 듀얼 시스템 모델 [142] 그리고 Helix [217] 모듈 해석 가능성을 유지하면서 두 구성 요소를 공동 학습시켜 이러한 구분 을 흐리게 합니다.

표 1 최근 VLA 이 이러한 균형을 어떻게 균형 있게 관리하는 지도 강조합니다. OpenDriveVLA 와 같은 모델 [293] 그리고 CombatVLA [34] 동적이고 안전에 중요한 영역에서는 계층적 계획을 우선시하는 반면, Edge VLA 와 같은 경량의 엣지 타겟 시스템은 선호합니다 [20] 그리고 TinyVLA [242] 실시간 저수준 정책을 강조하고 고차원 추론을 희생하세요. 이 분류 프레임워크는 VLA 의 설계 공간을 명확히 할 뿐만 아니라, 임베디드 배포에 최적화된 완전한 중단 간 투엔드 계층 모델과 같은 충분히 탐구되지 않은 조합을 정확히 찾아내어 로봇 공학, 자율 주행 등 다양한 분야에서 VLA 시스템의 기능과 적용 가능성을 발전시킬 것을 약속함으로써 향후 개발을 안내합니다.

표의 분류 1 이 는 다양한 VLA 아키텍처를 비교할 수 있는 명확한 틀을 제공하고, 중단 간 통합과 계층적 분해 같은 설계 선택 이 작업 수행, 확장성, 적응성에 어떤 영향을 미치는 지 강조하기 때문에 중요합니다. 모델을 저수준 정책 실행과 고수준 계획 등 차원별로 분류함으로써, 연구자들은 기존 접근법의 강점과 한계를 보다 명확히 파악하고 아키텍처 혁신의 기회를 발견할 수 있습니다. 예를 들어, 고속과 일수확이나 정밀 실패와 같은 농업로봇 작업은 빠르고 반응적인 저수준 제어를 강조하는 아키텍처의 이점을 누리는 반면, 과수원 내 비게이션, 다중 행커 버리지 계획, 장기 수작물 모니터링과 같은 응용 분야는 더 강력하고 고수준 계획 및 추론 능력이 필요합니다. 따라서 이 분류법은 특정 사용 사례에 적합한 VLA 아키텍처를 선택하는데 도움을 주며, 응답성과 인지 계획의 균형을 맞추는 하이브리드 시스템 개발을 이끌어 궁극적으로 구현 AI 의 발전을 가속화합니다.

또한, 최근 VLA 모델의 발전을 종합하기 위해 표 2 2022 년부터 2025 년까지 개발된 주목할 만한 시스템들을 요약합니다. 초기 융합, 이중 시스템 처리, 자기 보정 피드백 루프와 같은 아키텍처 혁신을 바탕으로 다양한 설계 철학과 훈련 전략을 통합합니다. 각 테이블 항목은 모델의 아키텍처 구성 요소—비전 인코더, 언어 인코더, 액션 디코더와 함께 모델 기능을 기반으로 평가하는데 사용되는 주요 학습 데이터셋—을 명시적으로 열거합니다. 예를 들어 CLIPort [202] 그리고 RT-2 [299] 의 미 임베딩을 액션 정책과 일치시켜 초기 기초를 마련했으며, 최근의 프레임워크는 Pi-Zero, CogACT [131], 그리고 GROOT N1 [14] 확산 기반 또는 고주파 컨트롤러를 통한 확장 가능한 아키텍처를 도입합니다. 여러 모델은 인터넷 규모의 시각 언어 말뭉치와 로봇 계획 데이터셋을 활용한 멀티모달 사전 학습을 활용하여 일반화와 제로 샷 능력을 향상시킵니다 [297, 291, 289, 257]. 이 표 비교는 실제 및 시뮬레이션 환경에서 VLA 설계의 기능 다양성, 도메인 적용 가능성, 그리고 새로운 트렌드를 이해하려는 연구자들의 기준점이 됩니다.

Table 1: 아키텍처패러다임과과학적우선순위에기반한구조적분류를보여주는 VLA 모델의분류학. 우리는모델을중단간실행, 계층적계획-통제분해, 구성요소중심모듈화를지원하는점과, 저수준운동정책과고수준작업계획자중하나로차별화합니다.

모델명	연도	중단간	계층형	구성요소 중심	저수준 정책	고수준 플래너
CLIPort [202]	2022	✓	✗	✗	✓	✗
RT-1 [19]	2022	✓	✗	✗	✓	✗
Gato [181]	2022	✓	✗	✗	✓	✗
VIMA [112]	2022	✓	✗	✗	✓	✗
Diffusion Policy [40]	2023	✓	✗	✗	✓	✗
ACT [287]	2023	✓	✗	✗	✓	✗
VoxPoser [100]	2023	✓	✗	✗	✓	✗
Seer [80]	2023	✓	✗	✗	✓	✗
Octo [218]	2024	✓	✗	✗	✓	✗
OpenVLA [122]	2024	✓	✗	✗	✓	✗
CogACT [131]	2024	✗	✓	✗	✓	✓
VLA Test [237]	2024	✗	✗	✓	✗	✗
NaVILA [38]	2024	✗	✓	✗	✓	✓
RoboNurse-VLA [132]	2024	✓	✗	✗	✓	✗
Mobility VLA [42]	2024	✗	✓	✗	✓	✓
RevLA [48]	2024	✗	✗	✓	✗	✗
Uni-NaVid [280]	2024	✗	✓	✗	✓	✓
RDT-1B [144]	2024	✓	✗	✗	✓	✗
RoboMamba [143]	2024	✓	✗	✗	✓	✗
Chain-of-Affordance [129]	2024	✗	✗	✓	✗	✗
Edge VLA [20]	2024	✗	✗	✓	✗	✗
ShowUI-2B [139]	2024	✓	✗	✗	✓	✗
Pi-0 [15]	2024	✓	✗	✗	✓	✗
FAST (Pi-0 Fast) [171]	2025	✗	✗	✓	✓	✗
OpenVLA-OFT [121]	2025	✓	✗	✗	✓	✗
CoVLA [5]	2025	✗	✓	✗	✓	✓
OpenDriveVLA [293]	2025	✗	✓	✗	✓	✓
ORION [69]	2025	✗	✓	✗	✓	✓
UAV-VLA [191]	2025	✗	✓	✗	✓	✓
CombatVLA [34]	2025	✓	✗	✗	✓	✗
HybridVLA [142]	2025	✗	✓	✗	✓	✓
NORA [103]	2025	✓	✗	✗	✓	✗
SpatialVLA [175]	2025	✗	✗	✓	✓	✗
MoLe-VLA [283]	2025	✗	✗	✓	✓	✗
JARVIS-VLA [130]	2025	✓	✗	✗	✓	✗
UP-VLA [279]	2025	✓	✗	✗	✓	✗
Shake-VLA [120]	2025	✗	✗	✓	✓	✗
DexGraspVLA [291]	2025	✗	✓	✗	✓	✓
DexVLA [241]	2025	✗	✓	✗	✓	✓
Humanoid-VLA [53]	2025	✓	✗	✗	✓	✗
ObjectVLA [297]	2025	✓	✗	✗	✓	✗
Long-VLA [64]	2025	✓	✗	✗	✓	✗
RetoVLA [123]	2025	✗	✗	✓	✓	✗
Vlaser [256]	2025	✗	✓	✗	✓	✓
Discrete Diffusion VLA [138]	2025	✓	✗	✗	✓	✗
Being-H0 [152]	2025	✓	✗	✗	✓	✗
EgoVLA [258]	2025	✓	✗	✗	✓	✗
StereoVLA [47]	2025	✓	✗	✗	✓	✗
GeoVLA [212]	2025	✗	✓	✗	✓	✓
EfficientVLA [262]	2025	✗	✗	✓	✓	✗

Table 2: 대표적인비전-언어-액션 (VLA) 모델의간결한요약. 각행은주요인코더/디코더, 훈련데이터, 그리고주요요유기능을보고합니다.

모델 (참고문헌)	아키텍처 (비전 / 언어 / 액션)	학습데이터	핵심강점 / 차별성
CLIPort [202]	CLIP-ResNet50 + Transporter-ResNet / CLIP-GPT / LingUNet	Self-collected [SC]	의미적 CLIP 특징을트랜스포터공간추론과정정렬하여정밀한 SE(2) 조작을지원합니다.
RT-1 [19]	EfficientNet / Universal Sentence Encoder / Transformer (이산화된동작)	RT-1-Kitchen [SC]	토큰화된행동을이용한다중작업주방조작을위한초기대규모트랜스포머정체.
RT-2 [299]	ViT-22B 또는 ViT-4B / PaLI-X 또는 PaLM-E / 심볼튜닝 (액션토큰)	VQA + RT-1-Kitchen	인터넷규모의 VQA 클로봇데이터와함께공동조정하여구현된작업에대한창발적일반화를가능하게합니다.
Gato [181]	ViT / SentencePiece / Transformer (통합토큰 스트림)	Self-collected [SC]	공유토큰화와단일트랜스포머를통해로봇공학, 언어, 아타리를통합하는범용에이전트.
VIMA [112]	ViT + Mask R-CNN / T5 / 트랜스포머	VIMA-Data [SC]	여러구성작업유형 (6 가지프롬프트양식) 에서의프롬프트기반 VL 그라운딩.
ACT [287]	ResNet-18 / — / CVAE-트랜스포머	ALOHA [SC]	시간양상블은세밀한제어로부터드러운양손모방을가능하게합니다.
Octo [218]	CNN / T5-베이스 / 확산트랜스포머	Open X-Embodiment (OXE)	대규모다중로봇정체는 4M+ 궤적을따라여러로봇구현체에걸쳐훈련되었습니다.
VoxPoser [100]	ViLD + MDETR / GPT-4 / MPC (LLM 가이드 계획)	Zero-shot	작업특화교육없이제약인지모션계획을위한 LLM+VLM을구성합니다.
Diffusion Policy [40]	ResNet-18 / — / U-Net 또는확산트랜스포머	Self-collected [SC]	확산모델링은견고한시간-운동조절을위한멀티모달액션분포를포착합니다.
OpenVLA [122]	DINOv2 + SigLIP / Prismatic-7B / 심볼튜닝	OXE + DROID	오픈소스 RT-2 유사 VLA; 효율적인 LoRA 적응과광범위한일반화를지원합니다.
π^0 (Pi-Zero) [15]	PaliGemma VLM / PaliGemma (멀티모달) / 300M 확산액션모델	Pi-Cross-Embodiment	경량일반로봇컨트롤러 (보고됨) ~ 총 3B 등급) 은강력한로봇간, 오픈월드일반화, 그리고양손능력을갖추고있습니다.
π^0 -Fast [171]	PaliGemma VLM / PaliGemma / FAST 토큰화가 적용된자기회귀트랜스포머	Pi-Cross-Embodiment	압축된주파수-공간동작토큰을통한고주파실시간제어 (최대 15x 더빠른추론).
OpenVLA-OFT [121]	SigLIP + DINOv2 (멀티뷰) / Llama-2 7B / 병렬디코딩 + 액션정제 (L1 회귀)	LIBERO; 양손 ALOHA	병렬디코딩과청크작업을통한레시피미세조정; 보고된수치는 97.1% LIBERO 성공과 26x 고주파양손제어를위한더빠른추론.
RDT-1B [144]	멀티뷰 RGB 인코더 / 트랜스포머언어모듈 / 확산트랜스포머 (통합액션공간)	46 개의데이터셋 (>100 만에피소드) + ALOHA 파인튜닝	1.2B 확산기초모델은강한언어조건화와제로샷전송을통한민첩한양손조작을위한모델입니다.
Helix ¹	시스템 2: 멀티모달추론을위한오픈소스 VLM (7-9 Hz) / 통합의미론 / 시스템 1: 트랜스포머시간-운동정책 (200 Hz, 삼체전체)	Figure robot E2E (pixels+language→actions)	휴머노이드중심이중속도 VLA 는실시간고자유도제어, 민첩성, 그리고제로샷일반화를통한다중로봇협력조작을가능하게합니다.
CogACT [131]	DINOv2 ViT-L/14 + SigLIP ViT-So400M/14 / Llama-2 Prismatic-7B / DiT-베이스 (300M 확산)	OXE 부분 집합; Realman & Franka 과제	확산액션트랜스포머를가진구성요소기반 VLA; 보고된 +59.1% OpenVLA 에비해실제성공과보이지않는로봇/물체에대한강한적응.
Chain-of-Affordance (CoA) [129]	어포던스인식시각인코더 / 트랜스포머추론프롬프트 / 자기회귀 + 확산정책 (어포던스조건부)	LIBERO; Real+Sim 조작	순차적어포던스 (행동유도성) 추론 (객체 → 파지 → 공간 → 동작은공간계획과장애물회피를향상시킵니다; OpenVLA 보다 LIBERO 성과가더강하다고보고했습니다.
Edge VLA (EVLA) [20]	SigLIP + DINOv2 / Qwen2 (0.5B) / 비자기회귀적공동제어예측	Bridge; OXE; 120 만개의텍스트-이미지쌍	엣지 최적화 VLA(예: Jetson 급) 는 30-50 Hz 추론과 OpenVLA 와비교할만한저전력성능.
ShowUI-2B [139]	UI 가이드시각적토큰선택 / 인터리브 V-L-A 스트리밍 / 트랜스포머 GUI 동작예측기	256K GUI 명령어팔로잉	디지털자동화를위한컴팩트 2B VLA; 강력한스크린샷그라운드와 GUI/웹내비게이션, 효율적인토큰선택.
GR00T N1 [14]	NVIDIA Eagle-2 VLM / 통합고수준계획 / 확산트랜스포머 (DiT)	인간시연 + 로봇궤적 + 시물레이션 + 인터넷영상	계획과확산실행을결합한다단계제어와광범위한신체간일반화를위한범용휴머노이드이중시스템설계.
Seer [80]	그라운딩최적화시각백본 / 트랜스포머언어 / 자기회귀동작헤드	LIBERO	조작을위한강한시각적기반; LIBERO 에서는경쟁력이 있으나, 일반적으로 OpenVLA-OFT 와같은최신파인튜닝변형모델보다는낮은성능을보입니다.
DiffusionVLA [240]	Transformer visual encoder / autoregressive reasoning / diffusion action head	LIBERO; 공장분류; 제로샷빈피킹	확산제어는견고성과해석가능성을향상시킵니다; 일부 구성에서 CoA 보다공간일반화가약하다고보고되었습니다.
NaVILA [38]	CLIP + CNN / LLaMA-2 / 계층적제어: 위상 계획 + RL 이동	실제다리로봇내비게이션데모	지형일반화를위한모듈식계층구조; 보고된 88% 자연어를통한실제내비게이션성공.
RoboNurse-VLA [132]	SAM2 + RGB-D / LLaMA-2 + 음성-텍스트변환/포즈회귀 + 그리퍼분류기	수술인수인계동영상 + 음성안내	실시간수술도구인수인계는새로운도구와역동적인수술실장면에강력하게대응합니다.

다음페이지에서계속

¹<https://www.figure.ai/news/helix>

이전페이지에서계속

모델 (참고문헌)	아키텍처 (비전 / 언어 / 액션)	학습데이터	핵심강점 / 차별성
Mobility VLA [42]	롱컨텍스트 ViT + 목표이미지인코더 / T5 기반 명령어인코더 / 그래프플래너 + 시각적목표현지화	MINT 데이터셋 (VL 명령어투어)	멀티모달투어에서언어위상지도는넓고보이지않는공간에대한내비게이션일반화를가능하게합니다.
TinyVLA [242]	FastViT / 컴팩트 128D 언어 / 확산디코더 (50M)	Mini-ALOHA + SC tasks	대규모사전훈련이필요하지않습니다; 더빠른추론보고 (5x) 이합수는제한된컴퓨터예산내에서강력한정밀도로이루어집니다.
QUAR-VLA [54]	CLIP + 고유수용성임베딩 / BERT + 그라운딩어댑터 / 트랜스포머전신디코더	QUART (이동 + 조작)	강력한시물레이션-실제 (Sim-to-Real) 전이과세밀한명령어정렬을가진쿼드루프중심 VLA 입니다.
ChatVLA [295]	위상정렬비전인코더 / Prismatic MoE LLM / 통합 V-L-A 플래너	통합채팅액션 (웹 + 로봇)	공동 VQA+ 기획; 잊어버리는것을줄이고조작을통해대화형작업실행을지원합니다.
PointVLA [125]	CLIP + 3D 포인트클라우드융합 / LLaMA-2 / 공간토큰융합을이용한트랜스포머	few-shot 공간작업 (실제 + 시물레이션)	사전학습된 2D 지식을유지하면서 3D 구조를주입하여장기과제공간추론을향상시킵니다.
VLA-Cache [252]	SigLIP + 토큰메모리버퍼 / Prismatic-7B / 동적토큰재사용트랜스포머	ALOHA + sim/real 퓨전	효율성을위해정적시간적토큰을캐시합니다; 보고된 40-50% 성능저하가거의없는빠른추론.
HybridVLA [142]	CLIP + DINOv2 / LLaMA-2 / 하이브리드확산 + 자기회귀양상블	RT-X + synthetic fusion	동적양상블은다양한환경에서견고성을높이고더강력한sim-to-real 일반화를지원합니다.
MoLe-VLA [283]	다단계 ViT + STAR 라우터 / CogKD 강화트랜스포머 / 희소트랜스포머 (동적라우팅)	RLBench + 실제조작	선택적충활성화는효율을제공합니다 (보고된 5.6x 속도 증가) 및더높은성공률 (보고된 +8%).
UAV-VLA [191]	항공이미지 / GPT 명령어파싱 / 트랜스포머 경로계획용 ViT	위성 + UAV 영상지침	대규모미지도환경을위한확장가능한언어기반을갖춘제로샷항공임무계획.
DexGraspVLA [291]	객체중심공간 ViT / 트랜스포머그래프추론 / 확산그래프컨트롤러	민첩잡기벤치마크 (시물 + 실제)	보도 >90% 다양한물체에서제로샷성공; 조명/배경변화와보이지않는조건에견고합니다.
GraspVLA [46]	멀티뷰 DINOv2 + SigLIP / VLM 이박스 + 그라스프 / 플로우매칭액션전문가 (PAG) 를예측합니다	SynGrasp-1B; GRIT	합성사전학습잡기 VLA; 시물레이션-실제 (Sim-to-Real) 전이를개선하고통태일객체클래스와선호도에대한제로/few-shot 일반화를지원합니다.
Interleave-VLA [62]	InternVL2.5 + OWLv2 / Qwen2.5 / 연속동작예측기 (OpenVLA+ π^0 -스타일, 확산컨트롤러)	Open Interleaved X-Embodiment (21 만에피소드, 11 개데이터셋)	중단간인터리브이미지-텍스트지시자되따르는방식; 보고된 2-3x 스케치와새로운멀티모달프롬프트를통한영역외이득과제로샷실행.
Long-VLA [64]	장기과제작업 / 위상인식일력마스킹 + 트랜스포머정책 (하위작업단계분할) 을위한중단간 VLA	장기과제다단계로봇조작시연 (작업시퀀스)	'이동' 단계와' 상호작용' 단계를명시적으로분리하여장기작업수행을목표하여, 하위작업호환성과장기작업범위에걸쳐견고함을향상시킵니다.
RetoVLA [123]	VLM 기반정책/재사용레지스터토큰을공간적맥락으로사용하여행동예측을지원합니다	7 도 DoF 팔에서의실제로봇조작 + 작업별시연	레지스터토큰을재활용하여경량공간추론업그레이드; 최소한의아키텍처오버헤드로복잡한조작에서상당한성공률향상을보고했습니다.
Vlaser [256]	VLM-to-VLA 파이프라인 / 시너지적체성추론 + 정책학습 (추론인지 VLA 미세조정)	Vlaser-6M 내체 추론 데이터셋 + VLA 미세조정데이터 (로봇시연)	Bridge 는추론과통제를구현했으며, 체화된기반/QA/기획전반에서강력한성과를보여주며, 도메인전환시정책학습으로의전환을개선했습니다.
Discrete Diffusion VLA [138]	단일트랜스포머 VLA / 이산화된액션청크 + 이산확산정책 (CE 교육; 재마스킹)	LIBERO + SimplerEnv(프랙탈/Bridge) 스타일의 VLA 벤치마크	분산스타일의정책을이산토큰인터페이스로통합합니다: 적응형디코딩순서, 리마스킹을통한오류수정, 그리고강한벤치마크성공률로자기회귀병목현상감소.
Being-H0 [152]	인간영상에서사전학습된민첩한 VLA / 명시적인손동작모델링 + 행동을위한 VL 그라운딩	대규모인간조작영상 (자기중심/3인칭) + 로봇제어로의전환	다양한인간비디오데이터를활용하여능숙한조작을확장합니다; 작은텔레옴스전용로봇데이터셋과비교해새로운작업/장면에대한일반화를향상시킵니다.
EgoVLA [258]	자기중심적인간조작 / 통합인간-로봇행동공간 + 로봇미세조정예관한 VLM 사전학습	대규모자기중심적인간영상 + 소규모로봇시연세트	확장가능한사전학습을위해풍부한자기중심적영상을활용한후, 통합행동공간을통해구현체를정렬하여제한된로봇데이터로실용적인로봇전송을가능하게합니다.
StereoVLA [47]	스테레오강화 VLA / 기하학-의미론적융합스테레오쌍 (+ 보조깊이단서) 을통한행동디코딩	스테레오로봇조작시연 (언어조건 동작이포함된스테레오 RGB)	스테레오기하학을명시적으로활용하여공간정밀도 (깊이민감잡기/조작) 와시점/카메라변화에대한견고함을향상시킵니다.
EfficientVLA [262]	훈련없는 VLA 가속/압축/구조화된중복성 제거 (VLA 파이프라인전반에걸친토큰/인터페이스최적화)	기존 VLA 에적용됨 (표준조작벤치마크평가; 예: SIMPLER 스타일 설정)	실용적배포가능성: 대규모 VLA 정책의연산/메모리를최소화하면서정확도저하를최소화하여, 제한된하드웨어에서더근실한추론을가능하게합니다.

3.2. 비전-언어-액션모델에서의훈련효율성발전

VLA 모델은멀티모달입력을조화시키고, 계산요구량을줄이며, 실시간제어를가능하게하는훈련및최적화기법에서빠르게발전했습니다. 주요발전분야는다음과같습니다:

• 데이터효율적인학습.

- 코-미세튜닝방대한비전-언어말뭉치 (예: LAION-5B) 와로봇계적집합 (예: Open X-Embodiment) 에서의미이해를운동능력과일치시킵니다. Open-VLA(7B 파라미터) 는 55B 매개변수 RT-2 변형보다 16.5% 더높은성공률을보이며, 공동파인튜닝이 적은매개변수로강력한일반화를가능하게함을보여줍니다 [194, 227, 122] 모델크기확장만비교했을 때와비교했을때입니다.
- 합성데이터생성 Via UniSim 은오클루전과동적조명을포함한사진사실적인장면을생성하여드문예외상황을보강하며, 복잡한환경에서모델의강인성을 20% 이상향상시킵니다 [264, 218].
- 자기감독사전훈련 CLIP 과같은대조적목적을채택하여동작미세조정전에시각적텍스트임베딩을 학습하여작업별라벨에대한의존도를줄입니다. Qwen2-VL 은자체감독정렬을활용하여하위 pick-and-place 과제의수렴을 12% 가속화합니다 [177, 98].

- **매개변수효율적인적응.** 저랭크적응 (LoRA) 은경량어댑터매트릭스를동결된트랜스포머층에삽입하여성능을 유지하면서도훈련가능한가중치를최대 70% 까지줄입니다 [93]. Pi-0 Fast 변형은정적백본위에 10M 어댑터매개변수만사용해정확도손실이거의없는 200Hz 연속제어를제공합니다 [171].

• 추론가속.

- 압축액션토큰 (빠르게) 그리고 병렬디코딩이중 시스템프레임워크 (예: GR00T N1) 에서는최대 2.5 를달성합니다.x 정책추론속도증가, 단계당지연시간을 5 미만으로줄입니다 실시간조작및휴먼노이드제어벤치마크를기준으로평가한단일정책 VLA 에서의표준자기회귀디코딩과 MS 의비교. 이가속은계적의부드러움에약간의대가를치르며, 액션이산화오차가약간더크고고주파제어하에서 미세한운동연속성이감소하는방식으로나타납니다 [14, 209].

이방법들은 VLA 를언어조건화되고시각요도작업을역동적이고현실세계에서처리할수있는실용적인에이전트로탈바꿈시켰습니다.

3.3. VLA 모델에서의매개변수효율적인방법과가속기법

데이터효율적인학습전략을넘어, VLA 모델에대한보완적인연구는모델크기, 메모리사용량, 추론지연을줄여제한된계산및전력자원을가진실제로봇플랫폼에배포할수있도록하는데중점을두고있습니다. 앞서는의한훈련중심효율성방법과달리, 이하위색선의기법들은주로대상 적응시점의매개변수효율성그리고 정책추론중런타임가속.

1. **저랭크모델을통한매개변수효율적인적응:** 완전한미세조정대신, 많은 VLA 는훈련효율성맥락에서이전에도입된 Low-Rank Adaptation(LoRA) 과같은매개변수효율적인적응메커니즘을채택합니다. 이절에서는그역할에중점을둡니다 배포중유�효매개변수발자국감소. 예를들어, OpenVLA 는경량 LoRA 어댑터 (약 20 개) 를사용합니다.M 파라미터) 위에얼어붙은 7 위에있습니다B-파라미터백본으로, 최소한의메모리오버헤드로작업적응을가능하게하고작업간전체모델가중치의중복을방지합니다 [93, 122]. 이설계는자원이제한된시스템에서여러작업특화된정책이공존하면서도사전학습과정에서학습된범용시각적-의미론적표현을유지할수있게합니다.
2. **엣지배포를위한양자화:** 모델양자화는추론처리량과메모리효율성을개선하기위해수치정밀도를낮춥니다. OpenVLA 실험결과, NVIDIA Jetson Orin 과같은임베디드플랫폼에서의 INT8 양자화는 pick-and-place 벤치마크에서완전정밀도대비약 97% 의작업성공률을보존했으며완전한정밀도로작업성공을거두었으며, 세밀하고민첩한조작에서는약간의저하만있었습니다 [194, 122]. 학습후양자화와채널별보정은고동적범위센서입력시정확도손실을더욱완화합니다 [164]. 이러한최적화는최대 30 개의지속제어주파수를가능하게합니다 이동식로봇의엄격한전력에산내에서 Hz 를사용합니다.
3. **모델가지치기와아키텍처슬림:** 구조적가지치기는어텐션헤드나피드포워드서브레이어와같은중복된아키텍처구성요소를제거하여메모리및컴퓨팅요구량을줄입니다. VLA 에서는독립적인시각또는언어모델에비해덜탐구되지만, 확산기반시력운동정책에관한초기연구들은합성곱시각인코더를최대 20% 까지줄여도잡기안정성저하가거의없음을보여줍니다 [40]. 트랜스포머기반 VLA 에적용되는유사한가지치기전략 (예: RDT-1B) 은메모리사용량을약 25% 줄이면서도과제성공률감소를 2% 미만으로유지할수있습니다, Sub-4 허용 GB 배포 [144, 131].
4. **압축액션토큰화:** 장기과제제어시퀀스에서발생하는추론병목현상을해결하기위해압축액션표현이제안되었습니다. FAST 는연속동작경로를컴팩트한주파수영역토큰으로재구성하여디코딩길이를크게단축합니다. Pi-0 Fast 변형은최대 15 개까지달성할수있습니다 x 1000 을 압축하여더빠른추론 MS 액션창은 16 개의개별토큰으로나뉘어최대 200 의제어속도를지원합니다 데스크톱 GPU 에서의 Hz [171]. 이방법은최소한의계적세분성을 대가로대가로큰속도향상을제공하여, 고주파의반응적인조작작업에적합합니다.
5. **병렬디코딩과액션청킹:** 표준자기회귀 VLA 는동작을 순차적으로디코딩하여누적지연시간을발생시킵니다. GR00T N1 과같은이중시스템아키텍처에서사용되는병렬디코딩전략은시공간액션토큰그룹을동시에생성하여약 2.5 개의액션토큰을달성합니다.x 100 도에서작동하는 7-DoF 로봇팔의중단간추론지연감소 Hz [14, 209]. 액션청킹은다단계루틴 (예: 픽앤플레이스) 단일고수준토큰으로변환하여추론단계를최대 40% 까지줄입니다 주방위크플로우와같은장기과제조작작업에서 [112].
6. **하드웨어인식컴파일및런타임최적화:** 마지막으로, 하드웨어인식최적화는컴파일러수준의그래프재작성, 커널융합, 가속기특화된시요소를활용하여처리량을극대화합니다. TensorRT-LLM 과같은프레임워크는텐서코

어, 퓨즈드어텐셔널커널, 파이프라인메모리전송을활용하여트랜스포머추론과확산샘플링을모두가속화합니다. OpenVLA-OFT 에서는이러한최적화가추론지연시간을약 30% 줄이고추론당에너지소비를 25% 줄입니다 RTX 급 GPU 와표준 PyTorch 실행비교 [121]. 이러한시스템수준의최적화는엄격한전력제약이있는이동식로봇, 항공플랫폼, 휴머노이드시스템에서실시간 VLA 를배포하는데매우중요합니다.

토론: 매개변수효율적인적응및추론가속기법은 VLA 배포를함께민주화합니다:

- LoRA 와양자화는소규모연구소와연구개발프로그램이소비자급하드웨어에서수십억파라미터의 VLA 를미세조정하고작동할수있게하여로봇의최첨단의미론적이해를열어줍니다 [93, 122].
- 가치치기와 FAST 토큰화는모델및액션표현을압축하여 4GB 이하, 5ms 미만의제어루프를가능하게하며, 민첩한작업에서정밀도를희생하지않습니다 [144, 171].
- 병렬디코딩과액션청킹은자기회귀정책의순차적병목현상을극복하여민첩한조작과다리이동에필요한 100–200 Hz 의사결정률을지원합니다 [14, 209].
- 하이브리드 RL-SL 훈련은복잡한환경에서탐색을안정화하며, 하드웨어인식컴파일은엣지가속기에서실시간성능을보장합니다 [159, 121].

이러한발전들은산업용조작기, 보조드론, 소비자로봇전반에 VLA 모델을내장하는것이실용적이어서, 연구프로토타입에서실제자율성으로의간극을메우는데기여하고있습니다.

3.4. 비전-언어-액션모델의응용

VLA 모델은지각, 자연어이해, 운동제어를통합한통합아키텍처내에서체화된지능의기초구성요소로빠르게부상하고있습니다. 시각적및언어적양식을공유된의미공간에인코딩하고맥락에기반한행동을생성함으로써, VLA 모델은에이전트와그환경간의원활한상호작용을가능하게합니다 [131, 293]. 시각적및언어적양상을공유된의미공간에인코딩하고맥락에기반한행동을생성함으로써, VLA 모델은에이전트와그환경간의원활한상호작용을가능하게합니다. 이러한멀티모달역량은 VLA 를다양한실제응용분야에서변혁적에이전트로자리매김하게했습니다. 휴머노이드로봇공학에서는 Helix 와 RoboNurse-VLA 와같은시스템이시각, 언어, 그리고능숙한조작을결합하여가정업무와외과수술을지원하며, 실시간추론과안전인식제어를보여줍니다 [132, 242]. 자율주행차에서는 OpenDriveVLA 와 ORION 과같은모델이동적시각스트림과자연어지시를처리하여복잡한도시환경에서투명하고적응적인운전결정을내립니다 [69, 293]. 산업배프는고정밀조립, 검사및협업제조를위해 VLA 아키텍처를활용합니다 [131]. 농업분야에서는 VLA 기반로봇시스템이시각유도과일수확, 식물모니터링, 이상탐지를가능하게하여노동의존도를줄이고지속가능성을높일수있습니다. 더 나아가, 최근인터랙티브증강현실시스템의발전은 VLA 모델을활용해실시간언어조건공간내비게이션을제공하며, 음성이나시각적단서를바탕으로실내외환경에서사용자를안내합니다 [191, 74]. 이들영역전반에걸쳐 VLA 는건고하고적응력

있으며의미적으로정렬된작업실행을위한통합프레임워크를제공하며, 이는체화된범용에이전트로의중대한전환을의미합니다.

표 3 최근 VLA 모델의방법론, 응용분야, 주요혁신을요약하여보여줍니다.

다음하위섹션들은그림에나타난대로적용분야를시간순서대로심층적으로탐구합니다 11.

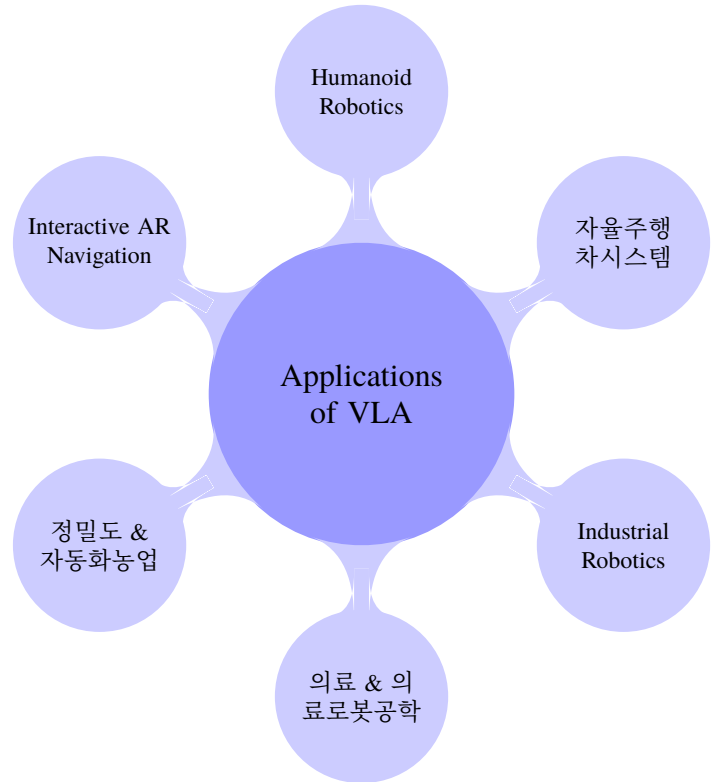


Figure 11: 비전-언어-액션모델의응용영역을마인드맵으로구성하며, 휴머노이드로봇공학이맨위에위치하고나머지영역들은이섹션의논의순서에맞춰시계방향으로배열되어있습니다.

3.4.1. 휴머노이드로봇

인간의신체형태와기능을모방하도록설계된휴머노이드로봇은 VLA 모델배포에있어가장까다롭지만영향력있는분야중하나를대표합니다. 이플랫폼들은복잡한환경을원활하게인지하고, 말하거나쓰여진자연어를이해하며, 인간수준의민첩성으로복잡한신체작업을수행해야합니다 [184, 23]. VLA 모델의핵심강점은지각, 인지, 제어를단일중단간훈련가능한프레임워크로통합할수있다는점에있습니다. 이를통해휴머노이드로봇이시각적입력(예: 복잡한장면의 RGB-D 이미지)을해석하고, 언어적지시(예: "순가락을서랍에넣기")를이해하며, 정밀한운동경로를생성할수있습니다 [151, 296].

최근의발전은휴머노이드로봇분야에서 VLA 의도입을크게가속화했습니다. 예를들어, Helix(Helix)²Figure AI 가개발한휴머노이드로봇은완전통합 VLA 모델을활용해전신조작을고주파로수행하며, 팔, 손, 몸통, 심지어미세한손가락움직

²<https://www.figure.ai/news/helix>

Table 3: 대표적인비전-언어-액션 (VLA) 방법론, 적용분야및주요혁신비교.

참고문헌 (연도)	VLA 방법론	응용분야	강점 / 핵심혁신
CogACT [131] (2024)	확산트랜스포머를기반으로한특수액션모델을갖춘모듈형 VLA 입니다.	산업로봇공학; 언어유도조작.	빠른적응과강력한일반화를갖춘견고한행동모델링을통해다양한 산업환경에서과제성공률을향상시킵니다.
VLA-Test [237] (2024)	VLA 조작평가를위한자동화된대규모테스트프레임워크.	로봇조작벤치마킹 (견고성/신뢰성).	체계적인다중장면및다중작업평가를통해강인성격차를드러내고 목표 VLA 개선을지원합니다.
NaVILA [38] (2024)	2 단계 VLA: 고수준시각언어는중간수준의행법명령을생성합니다; RL 로코모션은실행됩니다.	어지러운현실세계장면에서자연어를통한다리내비게이션.	모듈식고·저수준분할방식으로, 강한일반화에도전적인지형에서의높은실제성공률을갖추고있습니다.
RoboNurse-VLA [132] (2024)	실시간음성투어액션파이프라인을지원하는비전모델 (SAM2) + 언어모델 (Llama2).	수술보조 (기구잡기및인수).	정확한실시간지원, 보이지않는도구에대한견고함, 그리고동적인수술실조건에서의안정적인성능.
Mobility VLA [42] (2024)	목표위치및위상그래프탐색을위한긴컨텍스트 VLM 을가진게층적 VLA 입니다.	시연투어가포함된멀티모달교수법내비게이션 (MINT).	대규모공간에서의확장가능한내비게이션을위한시연을활용합니다; 새로운언어 + 이미지쿼리에견고합니다.
CoVLA [5] (2025)	CLIP 기반비전, Llama-2 언어, 행동을위한게적예측.	자율주행 (데이터셋 + VLA 훈련).	해석가능한장면이해와견고한경로계획을가능하게하는대규모풍부하게주석이달린데이터셋.
OpenDriveVLA [293] (2025)	2D/3D 시각토큰과언어임베딩의게층적정렬; 자기회귀에이전트-환경-자아모델링.	중단간자율주행.	복잡한교통환경에서계획및질의응답성능을향상시키는통합의미공간및상호작용모델링.
ORION [69] (2025)	QT-Former 는과거문맥, LLM 추론, 생성계획기획기 (generative planner) 를사용해과거문맥, LLM 추론, 계획예측을위한도구입니다.	종합적인자율주행시스템.	추론및행동공간을 VQA 및계획의통합최적화와일치시켜더강력한폐쇄루프성능을제공합니다.
QUAR-VLA [54] (2025)	QUART 기반시각과언어의융합을통해액션생성을지원합니다.	네발달린내비게이션, 조작, 전진작업.	레그드플랫폼에대한 VL-A 결합과명령어정렬, 강력한시물레이션-실물일반화.
TinyVLA [242] (2025)	Diffusion Policy 디코더가포함된컴팩트멀티모달백본.	빠르고데이터효율적인조작제어.	강한일반화로추론비용감소; 많은사전훈련없이효율성을향상시킵니다.
UAV-VLA [191] (2025)	모듈식파이프라인: 목표추출을위한 GPT, 객체탐색을위한 VLM, 액션생성을위한 GPT.	언어 + 위성이미지를기반으로한 UAV 임무계획.	최소한의사전훈련으로효율적인비행및작전계획; 직관적인인간-UAV 상호작용및벤치마킹을지원합니다.
Bi-VLA [74] (2025)	시각, 언어, 양손행동모델을연결하는멀티모달트랜스포머.	양손가정조작.	명시적인액션-모델통합을통한높은적응력과견고한실제양손실행.
ChatVLA [295] (2025)	VL-A 통합을위한 Mixture-of-Experts 와의단계별정렬교육.	통합된멀티모달이해와로봇제어.	잊어버림/간절을줄여 VQA 와조작모두를효율적인전문화로향상시킵니다.
RoboMamba [143] (2025)	Mamba/SSM 기반 VLA 와공동학습된비전인코더와 SE(3) 동작모델링을포함합니다.	효율적인로봇추론과조작.	선형복잡도추론과최소한의미세조정은빠른포즈인식추론과효율적인제어를가능하게합니다.
OTTER [97] (2025)	고정된사전학습된 VLM 을이용한텍스트인식특정추출.	제로샷일반화를이용한조작.	VLM 의의미적정렬을 VLM 의미세조정없이도과제관련특징선택을통해강력한제로샷전송을보장합니다.
PointVLA [125] (2025)	모듈식스킵블록을통해동결된 VLA 에 3D 포인트클라우드특징을주입합니다.	공간적추론; 소수의샷과장거리지평선조작.	2D 지식을유지하면서 3D 기하학을추가하여완전한재훈련없이도공간적기반을개선합니다.
VLA-Cache [252] (2025)	정적시각토큰을선택적으로재사용하는적응형토큰캐싱.	실시간효율적인조작추론.	레이어별토큰재사용은정확도손실을최소화하면서큰속도향상을제공하며, 로봇내배포에실용적입니다.
CombatVLA [34] (2025)	빠른추론을위한단축된 AoT 를적용한비디오액션 AoT 훈련.	3D 게임에서의실시간전투의사결정.	대규모추론가속과향상된전술적추론, 실시간상호작용환경에서의높은성공률.
HybridVLA [142] (2025)	협업확산및자기회귀액션정책이포함된통합 LLM.	시물레이션과실제작업모두에서단일또는이중중조작.	적응행동양상블은복잡한조작에대한견고성과일반화를향상시킵니다.
NORA [103] (2025)	Qwen-2.5-VL-3B 백본과 FAST+ 토큰라이저를사용하는 3B 매개변수 VLA 입니다.	범용물체로봇공학 (시물레이터 + 실제).	낮은오버헤드와강한시각적추론, 빠른액션디코딩; 더큰 VLA 과경쟁할수있습니다.
SpatialVLA [175] (2025)	공간인식 VLA 를위한 Ego3D 위치인코딩및적응형동작격자.	크로스로봇, 멀티태스킹, 제로샷조작.	명시적 3 차원공간통합과적응적이산화는전이와일반화를개선합니다; 오픈소스입니다.
MoLe-VLA [283] (2025)	라우터와증류를통한동적레이어건너뛰기를통한레이어혼합.	RLBench + 실제로봇에서의효율적인조작.	선택적충활성화는인지를유지하면서성공률을높이면서강력한속도향상을제공합니다.
JARVIS-VLA [130] (2025)	VL 가이드인스카키보드/마우스제어용액션헤드가있는대형 VLM 을후처리한모델입니다.	오픈월드비주얼게임 (예: 마인크래프트), 1k+ 과제등이있습니다.	자기지도후교육은강력한일반화와세계지식을제공합니다; 오픈소스작업스위트와모델.
UP-VLA [279] (2025)	공동멀티모달이해와미래예측목표를가진통합 VLA.	정밀한공간추론을통한신체적조작.	공동의미론 + 동역학학습은세밀한공간제어가필요한장기과제의수행능력을향상시킵니다.
Shake-VLA [120] (2025)	비전, 음성변환, RAG, 이상감지, 양손팔을갖춘모듈시스템.	양손캐테일잡음과잡음을섞는것.	신뢰할수있는재료취급, 레시피적응, 실제작업원료를포함한견고한중단간배포.
MoRE [285] (2025)	LoRA 모듈과 RL 기반 Q-함수훈련을활용한최소한 MoE.	네발로다기능이동, 항법, 조작을모두수행했습니다.	혼합품질데이터에대한확장가능한미세조정과시물과실에서강력한다중작업및 OOD 일반화를제공합니다.
DexGraspVLA [291] (2025)	게층적계획기 (사전학습된 VL) + 확산저수준컨트롤러.	다양한조건에서의일반적인손잡이능력.	도메인불변표현은객체, 조명, 배경전반에걸쳐강력한제로샷일반화를지원합니다.
DexVLA [241] (2025)	구현된커리큘럼학습을갖춘플러그인확산행동전문가.	일반로봇조작: 단일팔, 양손, 손재주.	신체간행동모델링과빠른적응은과제별튜닝없이도강력한장기성적성능을가능하게합니다.

임까지 실시간으로 제어합니다. 아키텍처는 이중시스템설계를 따르며, 멀티모달 트랜스포머는 언어 명령과 비전 스트림 같은 입력을 처리하고, 실시간 모터 정책은 200Hz 의 밀집된 동작 벡터를 출력합니다. 이로 인해 Helix 는 이전에 못한 객체와 작업에 걸쳐 일반화할 수 있으며, 작업별 재학습 없이 변화하는 환경에 유연하게 적응할 수 있습니다.

휴머노이드 시스템에서 VLA 의 주요 장점은 공유 표현을 통해 다양한 작업에 확장할 수 있다는 점입니다 [8]. 작업별 프로그래밍이나 모듈식 파이프라인에 의존하는 전통적인 로봇 시스템과 달리, VLA 기반 휴머노이드는 통합 토큰 기반 프레임워크 하에서 작동합니다. 시각 입력은 DINOv2 나 SigLIP 와 같은 사전 학습된 시각 언어 모델로 인코딩되며, 명령어는 LLaMA 나 GPT 스타일 인코더와 같은 LLM 을 사용해서 처리됩니다. 이 표현들은 장면과 작업의 전체 맥락을 포착하는 접두사 토큰으로 융합됩니다. 액션 토큰은 언어 디코딩과 유사한 자동 회귀 방식으로 생성되지만, 로봇의 관절과 최종 이펙터에 대한 운동 명령을 나타냅니다.

이 능력 덕분에 휴머노이드 로봇은 가정, 병원, 소매 환경 등 인간 중심 공간에서 효과적으로 작동할 수 있습니다. 가정 환경에서는 VLA 기반 로봇이 음성 명령을 해석하는 것만으로도 표면을 청소하거나 간단한 식사를 준비하거나 물건을 정리할 수 있습니다 [151, 296]. 의료 분야에서는 RoboNurse-VLA 와 같은 시스템이 있습니다 [132] 실시간 음성 및 시각적 신호를 이용해 의사와 제정밀한 기구 인수를 수행할 수 있는 능력을 입증했습니다. 소매업에서는 VLA 가장착된 휴머노이드 플랫폼이 고객 문의, 진열대 재입고, 매장 배치 탐색을 명시적 인사 전 프로그래밍 없이 지원할 수 있습니다 [8].

현대 Humanoid-VLA 를 구별하는 점은 임베디드 저전력 하드웨어에서 동작할 수 있어 실제 배포가 가능하다는 점입니다. 예를 들어, TinyVLA 와 같은 시스템들이 있습니다 [242] 그리고 MoManipVLA [246] Jetson 급 GPU 에서 실행되는 효율적인 추론 파이프라인을 보여주어 성능 저하 없이 모바일 배포를 가능하게 합니다. 이 모델은 확산 기반 정책, LoRA 기반 미세 조정, 동적 토큰 캐싱과 같은 기법을 활용하여 계산 비용을 최소화 하면서도 높은 정밀도와 일반화를 유지합니다.

물류 및 제조 분야에서는 VLA 지원 휴머노이드가 이미 상업적 영향력을 발휘하고 있습니다. 그림 01 과 같은 로봇은 창고에서 인간 작업자와 함께 피킹, 분류, 선반 정리와 같은 반복적이고 육체적으로 집약적인 작업을 수행합니다. 새로운 객체 범주와 역동적으로 변화하는 장면을 다루는 능력은 지속적인 학습과 견고한 멀티모달 기반 덕분에 강화됩니다 [252, 131].

VLA 모델이 다양한 액션 생성, 공간 추론, 실시간 적응 능력을 계속 발전함에 따라, 휴머노이드 로봇은 가정, 산업 환경, 공공 장소에서 매우 유능한 조수로 부상하고 있습니다. 이들의 강점은 토큰 기반 아키텍처를 통해 시각, 언어 이해, 운동 제어를 통합하는 능력에 있으며, 이를 통해 비구조화된 인간 환경에서 맵핑과 맥락 인식 있는 행동을 가능하게 합니다.

예를 들어, 그림에 묘사된 것처럼 12 차세대 VLA 모델을 갖춘 최첨단 휴머노이드 로봇 'Helix' 를 생각해 보십시오. "냉장고에서 물병을 꺼내주세요" 라는 구두 지시를 받으면 Helix 는 통합 지각 시스템을 활성화하며, 기초 시각 언어 모델 (예: SigLIP 또는 DINOv2) 이 시각 장면을 분할하여 냉장고, 손잡이, 병을 식별합니다. 언어 입력은 LLaMA-4 와 같은 LLM 에 의해 처리되며, 명령어를 토큰화하고 시각적 맥락과 융합합니다. 이 융합된 표현은 계층적 컨트롤러로 전달됩니다: 상위 정책 담당자는 작업 시퀀스 (손잡이 위치, 문 닫기, 병 식별, 잡기 등) 를 계획하고, 중간 단계 계획자는 잡기 유형과 관절 궤적과 같

은 모터 원시 요소를 정의합니다. 저수준 VLA 컨트롤러는 종종 Diffusion Policy 네트워크를 기반으로 하며, 이동 작동을 초미만의 지연 시간으로 실행합니다. 변이 (예: 기울어진 병이나 미끄러운 그림) 를 만나면, Helix 의 에이전트 AI 모듈은 실시간으로 미세 정책을 정제하여 피드백에 따라 그림을 조정합니다. 이 예시는 VLA 가 적용된 휴머노이드의 변혁적 잠재력을 보여줍니다. 주방부터 클리닉에 이르기까지, 이 시스템들은 복잡한 지시를 해석하고 신체 작업을 능숙하게 수행할 뿐만 아니라 환경의 예측 불가능성에도 적응합니다. 에이전트 추론과 안전 정렬 메커니즘을 내장함으로써, VLA 로 구동되는 현대 휴머노이드 로봇은 좁은 작업 수행자에서 범용적이고 신뢰할 수 있는 협력자로 전환하고 있습니다. TinyVLA 와 MoManipVLA 와 같은 에너지 효율 모델이 성숙함에 따라, 모바일의 저전력 플랫폼 배포가 점점 더 실용적이 되어 사회적으로 연계된 새로운 AI 시대 를 열고 있습니다.

3.4.2. 자율주행차 시스템

자율주행차 (AV), 자율주행차, 트럭, 공중 드론을 포함한 자율주행차는 VLA 모델의 최전선 적용 분야를 대표하며, 안전에 중요한 사실적이지 않거나, 의미 이해, 실시간 액션 생성이 긴밀하게 결합되어야 합니다. 인식, 계획, 제어를 명시적으로 분리하는 전통적인 모듈형 AV 파이프라인과 달리, VLA 프레임워크는 시각적 관찰, 고수준의 의미론적 단서, 내부 상태 표현을 통합된 모델 내에서 공동으로 처리하여 보다 긴밀한 아키텍처 결합을 탐구합니다. 이러한 엔드 투 엔드 공식화는 시뮬레이션과 명령 조건 내비게이션 및 추론의 제어 벤치마크에서 유망한 결과를 보여주었지만, 대규모 상업용 시스템 (예: 테슬라 오토파일럿) 은 Source Link) 는 모듈형 또는 하이브리드 파이프라인에 계속 의존하며, 최근 산업계에서는 안전 중요한 전에서 VLA 스타일의 액션 생성 전면 도입보다는 비전 언어 추론 구성 요소 통합에 집중하고 있습니다.

VLA 모델은 AV 가 픽셀 수준 객체 인식을 넘어서는 복잡한 환경을 이해할 수 있도록 지원합니다. 예를 들어, 도시 환경을 이동하는 자율주행차는 교통 표지판을 감지하고, 보행자 행동을 이해하며, "주유소 다음 두 번 재우회전" 같은 내비게이션 명령을 해석해야 합니다. 이 작업들은 시각 및 언어 신호를 융합하여 공간적 관계를 이해하고, 의도를 예측하며, 맥락 인식 주행 행동을 생성하는 것을 포함합니다. VLA 는 이 정보를 토큰 기반 표현을 통해 인코딩하며, 시각적 인코더 (예: ViT, CLIP), 언어 모델 (예: LLaMA-4), 궤적 디코더가 일관된 의미 공간에서 작동하여 고수준 목표를 추론하고 이를 저수준 움직임으로 변환할 수 있게 합니다.

이 방향에서 주목할 만한 기여는 다음과 같습니다 CoVLA [5] 는 80 시간 이상의 실제 주행 영상과 동기화된 센서 스트림 (예: LiDAR, 주행 거리), 상세한 자연어 주석, 고해상도 주행 경로를 결합한 포괄적인 데이터셋을 제공합니다. 이 데이터셋은 VLA 모델을 학습시켜 지각 및 언어적 특징을 신체 행동과 일치시킬 수 있게 합니다. CoVLA 는 시각적 그라운드링에 CLIP, 명령어 임베딩을 위해 LLaMA-2, 그리고 움직임 예측을 위해 궤적 디코더를 사용합니다. 이 구성은 AV 가 음성 신호 (예: "구급차에 양보") 와 환경 조건 (예: 교통 혼잡) 을 해석하여 투명하고 안전한 운전 결정을 내릴 수 있게 합니다.

OpenDriveVLA [293] 는 2D/3D 다중 뷰 시각 토큰과 자연어 입력의 계층적 정렬을 통합하여 VLA 모델링 기술을 발전시킵니다. 이 아키텍처는 자기 중심적 공간 인식과 외부 장면이 해를 모두 활용하여 동적인 에이전트-환경-자아 상호 작용 모델을 구축합니다. 자기 회귀 디코딩을 통해 OpenDriveVLA 는 조

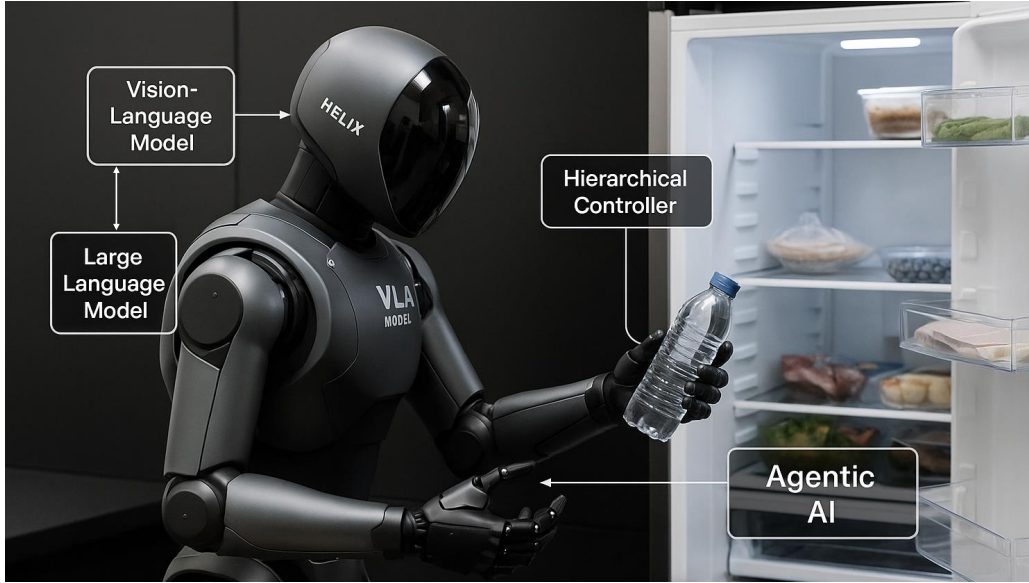


Figure 12: 이그림은 VLA 프레임워크를사용해집안일을수행하는차세대휴머노이드로봇 'Helix' 를보여줍니다. 구두명령을받으면 Helix 는시각언어모델(예: SigLIP) 과언어모델(예: LLaMA-4) 을통합하여환경을함께인지하고해석합니다. 계층적 VLA 컨트롤러는냉장고열기, 병잡기등하위작업을계획하고실행하며, 에이전트 AI 모듈은실시간으로동작을적응합니다. 이연구는동적과제적응과안전하고의미론적으로기반한조작을갖춘 VLA 기반범용로봇공학을시연합니다.

작계획(예: 조향각도, 가속도) 과인간이해석할수있는궤적시각화를생성합니다. 이종단간프레임워크는 nuScenes 와 Waymo Open Motion 데이터셋의계획및궤적예측과제뿐만아니라운전장면의비전-언어질문-응답벤치마크를포함한공자율주행벤치마크에서선도적인성능을발휘하며, 도시내비게이션과행동예측에서강력한강인성을 입증합니다.

또다른기념비적인모델, **ORION** [69] 은 QT-Former 를통합하여장기시각맥락을유지하고, LLM 을통해교통내러티브를추론하며, 생성궤적계획기틀도입하여폐쇄루프자율주행의경계를넓혔습니다. ORION 은시각언어모델의이산추론공간을 AV 모션의연속적제어공간과일치시키는데뛰어납니다. 이통합최적화는정확한시각적질문답변(VQA) 과궤적계획으로이어지며, 모호한인간지시나가려진장애물(예: "빨간트럭뒤출구로나가라") 이필요한상황에서매우중요합니다.

예를들어, 그림 13에묘사된것처럼, 차세대 VLA 아키텍처를사용하여밀집된도시환경에서운행하는자율운송차량인"AutoNav" 를고려해보십시오. AutoNav 가클라우드기반명령어로"빵집옆빨간차양근처에소포를내려놓고, 공사구역을피하며기지로돌아오라" 는명령을받으면, 온보드VLM(예: CLIP 또는 SigLIP) 이여러카메라의시각스트림을분석하여빵집간판, 빨간차양, 교통콘등등적랜드마크를식별합니다. 동시에, LLaMA-4 에그라운드된 LLM 모듈은명령어를해독하여 LiDAR, GPS, 관성주행측정등실시간감각맥락과융합합니다. 계층적제어스택은자기중심적관점과세계중심지도를통합하여적응경로를계획하는자기회귀 VLA 디코더를통해이러한멀티모달신호를처리합니다. 차량이배송위치에접근하면, 예상치못한보행자활동이강화학습기반정책개선루틴을사용해에이전트형서브모듈을촉발합니다. 동시에 AutoNav 는보행자에게경고를내고안전마진을유지하기위해속도를재조정합니다. 이러한의미적이해, 지각기반구조, 적응적제어의상호작용은 VLA 기반시스템이안전이중요한상황에서해석가능하고인간친화적인행동을달성하는데

마나강력한지를보여줍니다.

이시나리오는긴밀히통합된 VLA 아키텍처가종단간의미론적추론, 빠른모듈간적응, 해석가능한의사결정을가능하게하여전통적인인식-계획-제어파이프라인을능가할수있음을보여줍니다. 모듈식파이프라인에서는인식출력, 계획작업데이트, 제어조정이느슨하게결합된구성요소와의해제한된의미적피드백을받지만, VLA 기반시스템은언어의도, 시각적맥락, 체화된상태를함께추론하여동적으로경로를재계획하고, 안전관련의도를인간에게전달하며, 실시간으로제어정책을조정할수있습니다. 그결과, 시스템은더큰자율성, 인간이해석할수있는출력을통한투명성향상, 그리고안전이중요한환경에서의사결정의민첩성을향상시킵니다.

항공로봇공학에서 VLA 는드론이나 UAV 의전달및기타임무능력을향상시킵니다. 예를들어 **UAV-VLA** [191] 위성이미지, 자연어임무설명, 온보드감지를결합하여고수준명령(예: "파란방수포와함께옥상파드에전달") 을실행합니다. 이시스템들은모듈식 VLA 아키텍처를사용하여비전언어플래너가전지구적맥락을해석하고비행관제사가정확한웨이포인트를실행하여물류, 재난대응, 군사정찰분야의응용을지원합니다.

자율시스템이점점더비구조화환경에서작동함에따라, VLA 는전통적인파이프라인에대한확장가능하고해석가능하며데이터효율적인대안을제공합니다. 대규모멀티모달데이터셋에서학습하고의사결정을토론예측으로모델링함으로써, VLA 는인간수준의의미론을로봇움직임과일치시켜더안전하고스마트한자율주행및내비게이션기술의길을열어줍니다.

3.4.3. 산업로봇

산업로봇공학은 VLA 모델의통합으로패러다임전환을겪고있으며, 고수준추론, 유연한작업수행, 인간조작자와의자연스러운소통이가능한새로운세대의지능형로봇을가능하게하고있습니다 [32, 7]. 전통적인산업용로봇은일반적으로임

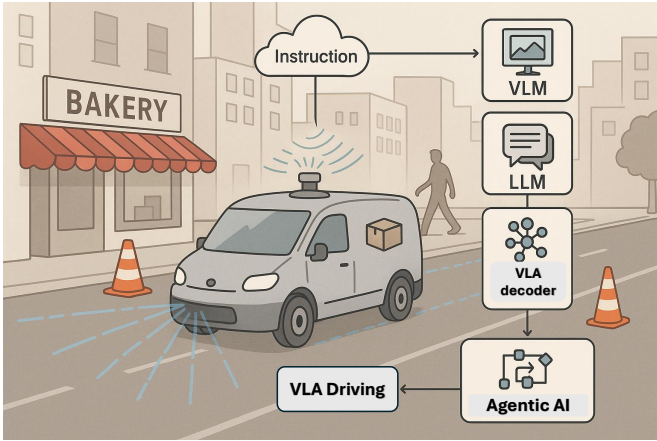


Figure 13: 이그림은 VLA 시스템으로구동되는자율운송차량을묘사하며, 시각적그라운드팅 VLM, 명령어파싱용 LLM, 경로계획용 VLA 디코더를통합합니다. 에이전트 AI 는동적환경에서적응형계획정제를가능하게하며, 멀티모달통합이실제내비게이션작업에서안전하고해석가능하며자율적인 의사결정을어떻게이끄는지를보여줍니다.

격한프로그래밍을사용하여고도로구조화된환경에서작동하며, 새로운조립라인이나제품변형에적응할때광범위한재구성과수동개입이필요합니다 [6, 182]. 이러한시스템은현대의동적제조환경에필요한의미론적기반과적응력이부족합니다.

반면 VLA 모델은더인간이해석하기쉽고일반화가능한프레임워크를제공합니다. 시각적입력 (예: 구성요소배치또는컨베이어벨트상태), 자연어명령어 (예: "빨간모듈의나사를조여라"), 로봇상태를공동삽입함으로써 VLA 는맥락을추론하고적절한제어명령을실시간으로실행할수있습니다 [135, 72, 156]. 비전트랜스포머 (예: ViT, DINOv2), LLM (예: LLaMA-4), 그리고자기회귀또는확산기반액션디코더가이시스템의중추를이루며, 로봇이멀티모달명령어를해석하고환경에기반한작업을수행할수있게합니다.

이분야에서가장중요한기여중하나 는 **CogACT** [131] 산업용로봇조작을위해명시적으로설계된모듈형 VLA 프레임워크입니다. 초기 VLA 가고정된언어-시각임베딩과직접액션양자화에의존하는것과달리, CogACT 는액션시퀀스를보다견고하고적응적으로모델링하는확산기반행동트랜스포머를도입합니다. 이시스템은시각언어인코더 (예: Prismatic-7B) 를사용하여고수준장면및명령어임베딩을추출하고, 이를확산트랜스포머 (DiT-Base) 로전달하여미세한모터동작을생성합니다. 이러한모듈식분리는보이지않는도구, 부품, 레이아웃에대한우수한일반화를가능하게하면서도실제제약조건하에서도해석가능성과견고성을유지합니다.

더나아가, CogACT 는 6-DoF 팔이나양손시스템등다양한로봇구현체에대한빠른적응을효율적인미세조정을통해입증하여, 이기종적인공장환경전반에배치할수있도록합니다 [131]. 실증평가에따르면 CogACT 는 OpenVLA 와같은이전모델보다 28% 이상우수한성능을보입니다실제과제성공률에서 [122] 특히다단계조립, 나사고정, 부품분류와같은복잡하고고정밀작업에서그렇습니다 [131, 262].

제조업이인더스트리 4.0 패러다임으로전환함에따라, VLA 는프로그래밍오버헤드를줄이고, 음성명령로봇프로그래밍을지원하며, 혼합이니셔티브작업에서실시간인간-로봇

협업을촉진할것을약속합니다. 실행정밀도, 안전보장, 지연최적화분야는여전히활발히연구되고있지만, VLA 모델의활용은자율적이고지능적이며적응력있는로봇으로나아가공장을변화시키는중요한진전입니다.

3.4.4. 헬스케어및의료로봇

의료및의료로봇공학은정밀함, 안전성, 적응력이 VLA 모델이점점더잘제공하는중요한특성인고위험분야를대표합니다 [132, 193]. 전통적인의료로봇시스템은원격조작이나미리프로그래밍된행동에크게의존합니다 [166, 204] 이는역동적인외과또는돌봄환경에서자율성과반응성을제한합니다. 반면, VLA 모델은실시간시각지각, 언어이해, 세밀한운동제어를통합하는유연한프레임워크를제공하여의료로봇이고수준지시를이해하고복잡한절차나보조작업을자율적으로수행할수있게합니다 [131, 54, 225].

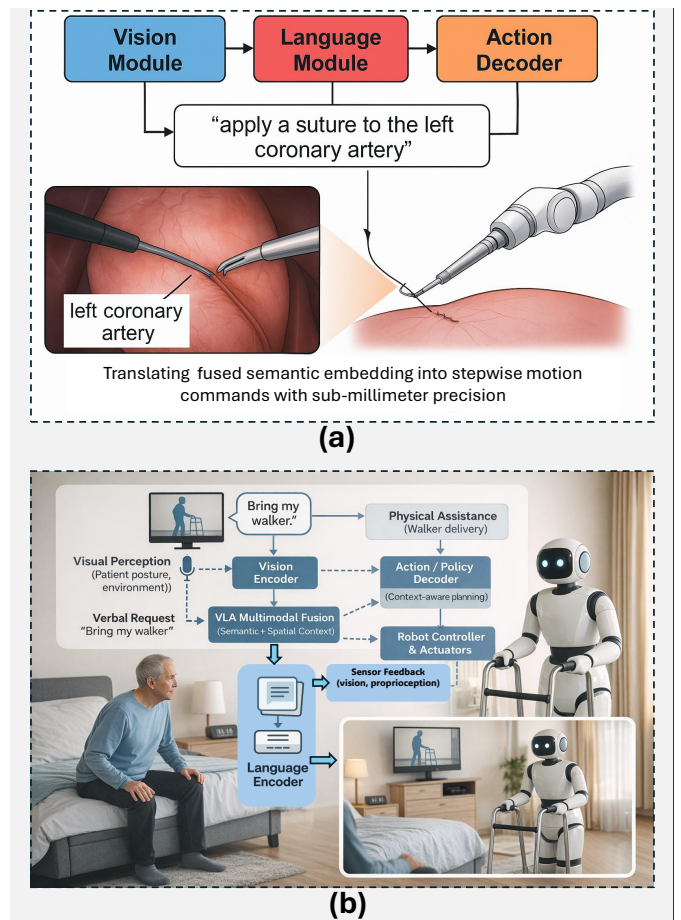


Figure 14: a) 이그림은 VLA 수술시스템이" 왼쪽관상동맥에봉합사를삽입" 하는작업을수행하고있음을보여줍니다. 비전모듈은해부학적표적을식별하고, 언어모델은지시를해석하며, 액션디코더는정밀한운동명령을생성하여적응형도구제어, 실시간피드백, 안전한자율작동을가능하게합니다. b) VLA 기반보조로봇은환자행동을인지하고, 구두요청 (예: "내워커를가져와줘") 을처리하며, 상황인식동작계획을자율적으로실행하여노인돌봄, 재활, 병원물류분야에서사전정의된스크립트나수동감독에의존하지않고실시간지원을가능하게합니다.

수술로봇공학에서 VLA 는최소침습수술의능력을획기적으로향상시킬수있습니다 [50, 230]. 이시스템들은복강경비디오피드를융합할수있습니다 [127], 해부학지도 [147, 50] 음성명령은비전인코더 (예: ViT, SAM-2) 와언어모델 (예:

LLaMA, T5) 을 사용하여 통합된 토큰화된 표현으로 변환합니다 [236]. 예를 들어, 그림에 나타난 바와 같이 14A, "왼쪽관상동맥에 봉합사를 붙이기"와 같은 작업에서는 시각모듈이 해부학적 대상을 식별하고, 언어모듈은 지시를 맥락화합니다. 액션 디코더는 융합된의미 임베딩을 서브밀리미터 정밀도로 단계별 모션 명령으로 변환합니다. 시각적인식, 언어기반의도, 행동수준 제어의 이러한 폐쇄루프 융합은 로봇이 도구를 적응적으로 재배치하고, 동적 포스피드백을 적용하며, 중요한 해부학적 구조를 회피함으로써 의사의 세세한 관리가 필요없이 줄어 들고 인간의 실수 위험을 최소화할 수 있게 합니다.

수술실 외에도, VLA 모델은 노인 돌봄, 재활, 병원 물류 분야에서 새로운 세대의 환자 보조 로봇을 구동하고 있습니다. 이 시스템들은 환자 행동을 자율적으로 인지하고, 말이나 몸짓 입력을 이해하며, 약물 수거, 이동 보조기 구안, 응급 상황 시 돌봄 제공자 알림과 같은 반복적인 작업을 수행할 수 있습니다. 예를 들어, 그림 14에 묘사된 것처럼, VLA 지원 로봇은 환자가 침대에서 일어나려는 것을 시각적으로 감지하고, "내워커를 가져오라"와 같은 구두 요청을 해석하며, 사전 정해진 대본이나 지속적인 감독 없이도 상황에 맞는 동작 계획을 생성할 수 있습니다.

RoboNurse-VLA와 같은 최근 VLA 프레임워크 [132]이 접근법의 현실적 실현 가능성을 강조합니다. RoboNurse는 의미론적 장면 분할을 위해 SAM-2를, 명령 이해를 위해 LLaMA-2를 사용하며, 이를 실시간 음성-투-액션 파이프라인에 통합하여 수술실 내 수술기 구인수인계를 로봇이 지원할 수 있게 합니다 [132]. 이 시스템은 다양한 도구, 다양한 조명 조건, 소음 환경에 견고함을 입증하며, 이는 임상 환경에서 흔히 발생하는 도전 과제입니다.

또한, VLA 아키텍처는 규제 의료 분야에서 매우 중요한 설명 가능성과 감사 가능성 측면에서 장점을 제공합니다 [224, 146]. 장면 지상 및 궤적 예측은 사후에 시각화하고 검토할 수 있습니다 [278] 임상 신뢰를 촉진하고 FDA 식검증 파이프라인을 가능하게 할 수 있습니다. LoRA 기반 미세 조정은 최소한의 데이터와 계산 인프라로 특정 병원 환경이나 절차적 워크플로우에 적응할 수 있게 해줍니다 [9, 229, 147].

중요하게도, VLA 모델의 멀티모달 기반은 도메인 간 전환성을 가능하게 합니다: 수술 도구 조작에 대해 훈련된 동일한 모델이 적당 한 재교육만으로도 환자 이동성 작업에 적응할 수 있습니다 [56]. 이러한 모듈화는 작업별 자동화 시스템에 비해 개발 시간과 비용을 크게 줄여줍니다 [94]. 의료 로봇 공학이 원격 조종 보조에서 반자율 및 협업 시스템으로 전환함에 따라, VLA 모델은 이 변화의 핵심에 자리 잡고 있습니다.

앞서 다른 응용 분야에서는 논의되었듯이, VLA가 고수준의 미론적 이해와 저수준 제어를 결합하는 능력은 확장 가능하고 인간 친화적이며 적응형로봇의료를 위한 통합 솔루션을 제공하는 데 필수적입니다 [249, 295, 279]. 의료 시스템이 수요와 인력 부족 증가에 직면함에 따라, VLA 기반 로봇 공학은 의료 정밀도, 운영 효율성, 환자 중심 치료를 향상시키는 데 중요한 역할을 할 것입니다.

3.4.5. 정밀 및 자동화 농업

그림 15에 나타난 바와 같이, VLA 모델은 다양한 농업 환경에서 노동 집약적 작업에 지능적이고 적응력 있는 솔루션을 제공하는 정밀 및 자동화 농업 분야에서 혁신적 도구로 부상하고 있습니다 [70, 191]. 전통적인 농업 자동화 시스템은 각 작업이나 환경 변화에 따라 수동으로 재프로그래밍해야 하는 경직되고 센서 기반 파이프라인에 의존합니다 [220, 108], VLA는 멀티모달

지각, 자연어 이해, 실시간 액션 생성을 통합된 프레임워크 내에 통합합니다 [168, 84]. 이 통합된 멀티모달 통합은 자율 지상 로봇과 드론이 복잡한 밭 장면을 해석하고, 음성 또는 텍스트 기반 농업 지시를 따르며, 선택적 과일 수확이나 적응적 관개 같은 상황 인식 행동을 생성할 수 있게 합니다. VLA는 폐색, 지형 불규칙성, 조명 변동성, 다양한 작물 유형에 동적으로 적응할 수 있는 능력과 합성되고 사실적인 데이터셋을 이용한 학습을 결합하여 작물 유형, 지리, 계절 전반에 걸쳐 일반화할 수 있습니다. 액션 토큰화를 활용함으로써 [245], 트랜스포머 기반 정책 생성 [12, 85], 그리고 LoRA 미세 조정과 같은 기법들 [93] 이 시스템들은 지속 가능하고 정밀한 농업에 위한 농업 로봇의 확장성과 지능을 재정의하고 있습니다.

현대 과수원과 기타 작물 밭에서는 VLA가 RGB-D 카메라, 다중 분광 센서, 드론 등에서 오는 시각적 입력을 처리하여 식물 성장을 모니터링하고, 질병을 감지하며, 영양 결핍을 식별할 수 있습니다. 시각 트랜스포머 (예: ConvNeXt, DINOv2)는 시각적 장면에서 공간적·의미론적 정보를 인코딩하는 반면, LLM (예: T5, LLaMA)은 "동쪽 구역에 흰가루병 검사"나 "관개도랑 근처에서 익은 사과 수확하기"와 같은 자연어 명령을 구문 분석합니다. 토큰 융합을 통해 이러한 양식들은 공유 표현 공간에 정렬되어, 로봇이 정밀하고 맥락 인식 있는 동작을 수행할 수 있게 합니다.

예를 들어, 그림 15의 과일 따기 과제에서 볼 수 있듯이, VLA가 장착된 지상 로봇은 이미지 기반 속성 신호를 사용해 익은 농산물을 식별하고, "A 등급 과일만 따라"와 같은 사용자가 지정한 기준을 해석하며, 최종 이펙터를 제어하는 액션 토큰을 통해 동작 시퀀스를 실행할 수 있습니다. 이 방법은 작물 피해를 최소화하고, 수확률을 최적화하며, 가림이나 지형 이동과 같은 예상치 못한 변수에 실시간으로 적응할 수 있게 합니다. 관개 관리에서는 VLA 모델에 의해 유도된 드론이 현장 지도와 구두 지시를 해석하여 수분 부담 지역을 선택적으로 적용하여 물 사용량을 줄일 수 있습니다.

즉각적인 작업 실행을 넘어, VLA 모델은 폐쇄 루프 피드백 메커니즘을 통한 동적 재구성 과 평생 학습을 지원할 것으로 기대됩니다. 특히, 실행 결과, 감각 관찰, 배치 중 수집된 작업 성공 신호는 기록되어 정기적으로 VLA 정책의 오프라인 또는 점진적 업데이트에 반영될 수 있습니다. 작물 환경의 사진 사실 시뮬레이션 (예: 3D 과수원 렌더링)에서 생성된 합성 학습 데이터와 결합하면, 이러한 피드백 기반 적응은 광범위한 수동 주석 없이도 새로운 작물 품종, 해충 상태, 계절 변화에 대한 견고성을 점진적으로 향상시킬 수 있게 할 수 있습니다. LoRA 어댑터와 확산 기반 정책 정제와 같은 매개 변수 효율적인 기법이 계산 오버헤드를 최소화하면서 이러한 지속적인 적응을 촉진하는 데 핵심적인 역할을 할 것으로 기대됩니다.

전반적으로 VLA 모델을 농업 작업 흐름에 통합하는 것은 숙련된 작업의 의존도 감소, 목표 기반 개입을 통한 수확량 향상, 최적화된 투입 사용을 통한 환경 지속 가능성 증대 등 여러 장기적 이점을 제공할 것으로 기대됩니다. 전세계 식품 시스템이 점점 더 기후 변동성과 자원 제약에 직면함에 따라, VLA 기반 농업 기술은 실제 세계의 복잡성을 더 잘 수용하는 확장 가능하고 지능적이며 맥락 인식 있는 농업 관행에 기여할 것으로 기대됩니다.

3.4.6. 비전-언어-액션 모델을 이용한 인터랙티브 AR 내비게이션

인터랙티브 증강 현실 (AR) 내비게이션은 VLA 모델이 실시간으로 지능적이고 맥락 인식적인 안내를 제공하여 인간과 환경 상호 작용을 크게 향상시킬 수 있는 최전선을 나타냅니다

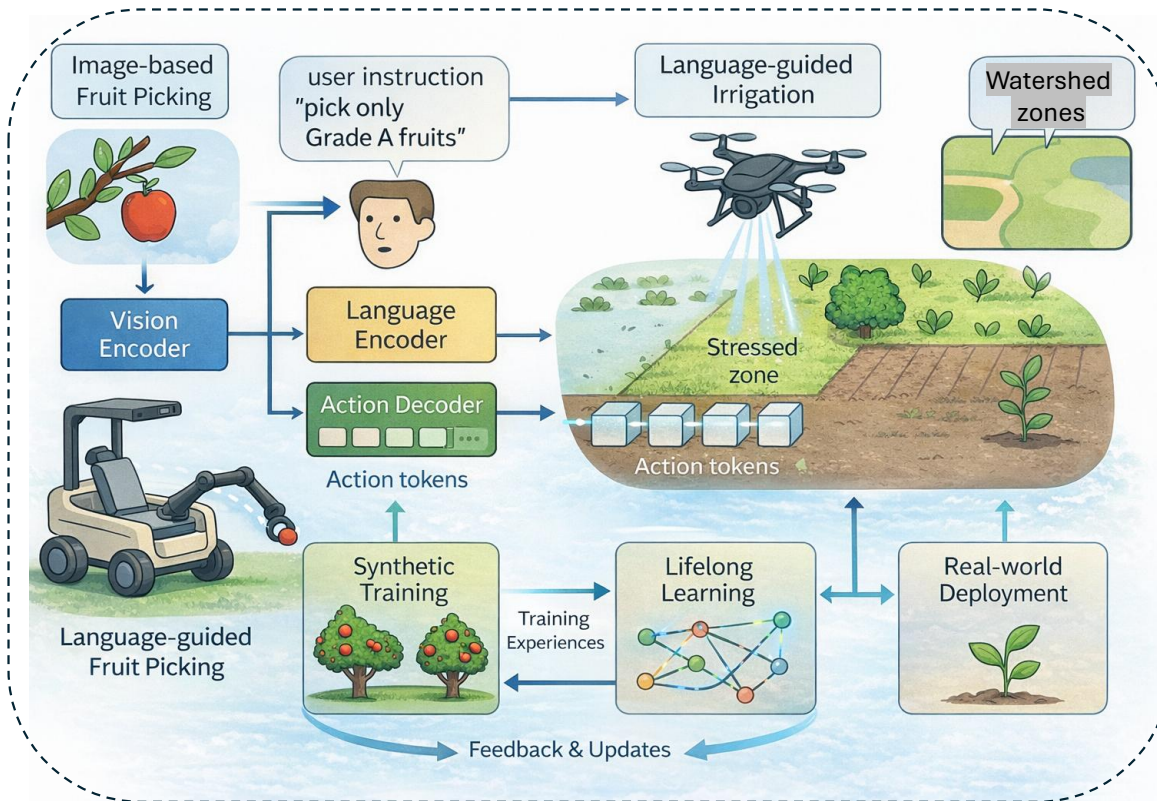


Figure 15: 정밀및자동화농업에서의 VLA 모델의개념적설명. 지상로봇은시각인코더와언어명령 (예: "A 등급과일만따라") 을결합해무손상수확을위한액션 토큰을생성하며, 항공로봇은언어유도관개를위해 VLA 추론을사용합니다. 합성교육, LoRA 기반적응, 평생피드백을통해작물, 환경, 지리적전반에걸친일 반화가능하며, 지속가능하고데이터기반농업자동화를지원합니다.

다 [31, 104, 255]. 이 패러다임에서 VLA 는스마트글라스나 스마트폰과같은 AR 지원장치에서발생하는시각적데이터의 연속적인스트림을자연어쿼리와함께처리하여, 사용자의물리적세계시점에직접접촉놓은동적내비게이션단서를생성합니다. 전통적인 GPS 기반시스템과달리, 엄격한지도와제한된사용자입력에의존하는것과다릅니다 [29, 207], VLA 기반 AR 요원들은교차로, 실내복도, 표지판등복잡한시각적장면을해석하고, " 휠체어경사로가있는가장가까운약국으로데려다주세요" 또는" 회의실로가는가장조용한경로를안내해주세요" 와같은자유로운형태의지시에응답합니다.

기술적으로이모델은 RGB 카메라프레임에서장면표현을추출하는비전인코더 (예: ViT, DINOv2), 사용자프롬프트나 음성명령을처리하는언어인코더 (예: T5 또는 LLaMA), 그리고방향오버레이, 웨이포인트, 음성지시와같은토큰화된내비게이션신호를예측하는액션디코더를통합합니다. 트랜스포머기반아키텍처는이러한양식을융합하여공간적배치와의미적의도를모두추론하여, AR 에이전트가사용자의시야내에서경로, 랜드마크, 위험요소를적응적으로하이라이트할수있게합니다 [211, 165]. 예를들어, 그림 16 에나타난바와같이 혼잡한공항에서는 VLA 상담원이에스컬레이터, 게이트, 수하물찾는곳을시각적으로식별하고, " 계단없이 22 번게이트에어떻게도달하나요?" 와같은질문을이해하고, 실시간인원상황과장애물에맞춰경로를조정할수있습니다.

VLA 는사용자가먼저고수준명령어 (예: ' 약국으로이동') 를내린후, ' 봄비는구역피하기' 나' 경치좋은길로이동하기' 와같은추가제약으로이를세정할수있는대화형명령어루프를

지원할것으로기대됩니다. 맥락인식피드백과반복적인명확화를통해이러한상호작용패러다임은시각장애인이나인지장애인의접근성과사용성을향상시킬수있습니다. 물류및실내내비게이션에서는 IoT 센서와디지털트윈과통합되어창고작업자, 유지보수팀또는배달로봇이복잡한환경에서안내할수있습니다. 또한, VLA 모델이사용자의선호도와지역공간배치를시간에따라학습하는지속적인미세조정을통해개인화된내비게이션이가능합니다.

AR 하드웨어가점점더저렴해지고일상생활에통합됨에따라, VLA 기반내비게이션시스템은공공, 산업, 보존환경에서원활한공간이해, 멀티모달상호작용, 자율안내를가능하게하여인간이물리적공간을인식하고탐험하며상호작용하는방식을재정의할것입니다.

4. 비전-언어-액션모델의도전과한계

VLA 모델은연구프로토타입에서견고한실제시스템으로의전환을방해하는상호연관된다양한도전과제에직면해있습니다. 첫째, 실시간자원인지추론을달성하는것은여전히어렵습니다: DeeR-VLA 와같은모델은동적조기종료아키텍처를활용해조작벤치마크에서 5-6x 계산을줄이면서도정확도를유지하지만, 복잡한상황에서는그효과가줄어듭니다 [269]. 마찬가지로 Uni-NaVid 는 5Hz 내비게이션을위해자기중심적인비디오토큰을압축하지만, 매우모호한지시와긴지평선아래에서는여전히어려움을겪습니다 [280]. 더불어, 이러한효율성중심설계는계산속도와표현커버리지사이의절충관계

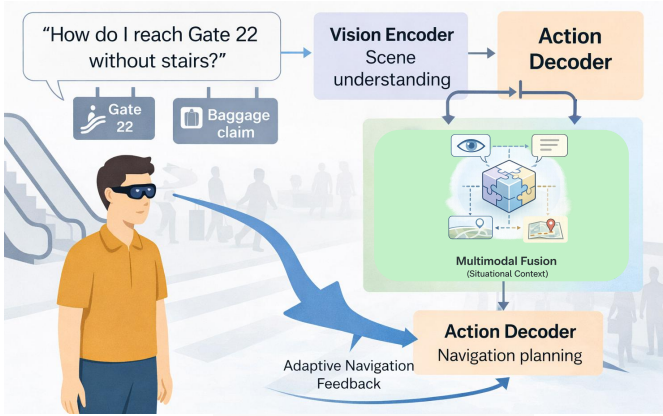


Figure 16: VLA 모델이 실시간 시각 인식, 언어 이해, 행동 계획을 융합하여 인터랙티브 AR 내비게이션을 가능하게 하는 방식을 보여줍니다. 공항과 같은 동적 환경에서는 VLA가 “22번 게이트로 계단을 피하라”와 같은 사용자 쿼리를 해석하고, 시각적 장면(예: 에스컬레이터 감지)을 분석하며, 이에 따라 내비게이션 경로를 조정하여 개인화되고 접근 가능하며 상황 인식 있는 이동 안내를 지원합니다.

를 자주 드러냅니다. 공격적인 압축이나 조기 종료 제약 하에서 작동할 때, 고급 하이브리드 비전-언어 기반 기법조차도 제한된 객체 일반화를 보입니다; 예를 들어, ObjectVLA는 새로운 객체의 64%에 대해서만 일반화되며, 이는 실시간 최적화가 오픈월드 강인성의 격차를 얼마나 크게 확대할 수 있는지를 보여줍니다 [297].

둘째, 최소한의 감독으로 VLA 모델을 적응시키고, 드물고 잡음이 많은 데이터에서 안정적인 정책 업데이트를 보장하는 것은 쉽지 않습니다. ConRFT는 행동 복제와 Q-학습을 인간 개입 미세 조정 (Human-in-the-loop fine-tuning) 과 결합하여, 접촉이 많은 8개 과제에서 96.3%의 성공률을 빠르게 수렴했지만, 전문가의 개입과 보상 형성에 크게 의존합니다 [36]. Hi Robot과 같은 계층적 프레임워크는 고수준 추론과 저수준 실행을 분리하여 명령어 충실도를 높이지만, 이러한 모듈을 조정하고 모호한 피드백을 안정시키는 것은 여전히 어렵습니다 [199]. 마찬가지로, 촉각 스트림과 언어 명령어를 융합한 촉각-언어-액션 모델은 보이지 않는 peg-in-hole 작업에서 85% 이상의 성공률을 달성했지만, 데이터셋의 폭과 실시간 단계 디코딩은 여전히 더 넓은 일반화에 한계가 있습니다 [88].

더 나아가, 동적 환경에서 안전성, 일반화 및 중단 간 신뢰성을 보장하기 위해서는 새로운 모델링 및 평가 기준이 필요합니다. OccLLaMA와 같은 점유-언어-행동 모델은 3D 장면 이해와 행동 계획을 통합하지만, 양식 전반에 걸쳐 더 풍부한 장면 역학과의 미적일 관성을 위해 확장되어야 합니다 [239]. RaceVLA는 양자화된 반복 제어 루프를 통해 고속 드론 내비게이션을 발전시킵니다; 하지만 더 큰 VLA와 전용 추론 모델에 비해 시각적-물리적 일반화가 제한적이며 보이지 않거나 빠르게 변화하는 환경에서의 안전 우려가 제기됩니다 [195]. ReVLA의 모델 병합 전략은 영역 외 시각적 강인성을 잃었던 것을 회복하여 방향 대상 감지 (OOD) 감지 성공률을 최대 77%까지 향상시킵니다 하지만 추가적인 계산 과부하를 도입합니다 [48]. 마지막으로, SafeVLA는 제약 마르코프의 사결정 과정을 통해 안전하지 않은 행동을 80% 이상 줄이기 위한 제약 조건을 공식화하지만, 다양한 실제 업무에 대한 포괄적이고 제한 없는 안전 규칙을 정의하는 것은 여전히 미해결 문제입니다 [274]. 이러한 교차하는 한계를 해결하는 것은 VLA 모델이 실제로 보

공학의 전체 복잡성 하에서 신뢰할 수 있고 자율적으로 작동할 수 있도록 매우 중요합니다.

위에서 언급한 중요한 한계를 바탕으로, 각 도전 과제를 목표 완화 전략과 매핑하고 시스템 수준의 영향을 평가하는 것이 필수적입니다. 표 4는 핵심 한계를 식별하고, 최근 발전에서도 출현 잠재적 기술적 해결책이며, 실제 VLA 배포에 대한 기대되는 이점을 명확히 설명합니다. 예를 들어, 실시간 추론 제약 문제를 해결하는 방법은 병렬 디코딩과 양자화된 트랜스포머 파이프라인을 하드웨어 가속 (예: TensorRT) 으로 활용하여 드론과 매니퓰레이터에서 제어 루프 속도를 유지합니다 [129, 122, 75, 142]. 하이브리드 확산-자기회귀 정책을 통한 멀티모달 액션 표현은 복잡한 작업에 대해 다양하고 맥락 민감한 동명령을 생성하는 모델을 풍부하게 만듭니다 [171, 156]. 오픈월드의 안전을 보장하기 위해 동적 위험 평가 모듈과 적응형 계획 계층을 통합하여 예측 불가능한 환경에서 강력한 비상 정지 행동을 보장할 수 있습니다 [183, 233, 113]. 마찬가지로, 데이터셋 편향과 그라운드링은 선별된 편향 제거 방법과 고급 대조 미세 조정을 통해 최소화할 수 있으며, 새로운 객체와 장면의 일반화 할 때 공정성과 의미적 정확성을 강화할 수 있습니다 [185, 17, 175]. 이러한 전략들과 시뮬레이션-실제 (Sim-to-Real) 전이, 촉각 통합, 에너지 효율 아키텍처를 아우르는 다양한 접근법이 결합되어 VLA 연구를 신뢰할 수 있고 확장 가능한 자율성으로 전환하는 포괄적인 로드맵을 형성합니다.

이 섹션의 나머지 부분은 다섯 개의 집중된 하위 섹션으로 구성되어 있으며, 각 섹션은 문헌에서 확인된 VLA 문제의 뚜렷한 군집을 살펴봅니다. 먼저, 실시간 추론 제약과 이를 해결하기 위한 신흥 방법을 분석합니다. 다음으로, 오픈월드 환경에서의 안전성 보장과 함께 멀티모달 액션 표현을 탐구합니다. 그 후 데이터셋 편향, 기반 전략, 보이지 않는 작업에 대한 일반화에 대해 논의하고, 이어서 시스템 통합 복잡도와 계산 요구량에 대해 탐구합니다. 마지막으로, 우리는 견고성과 실제 응용 분야에서 VLA 배포의 윤리적 함의를 고려합니다.

4.1. 실시간 추론 제약

최근 진전에도 불구하고, 지연이 중요한 환경에서 VLA 모델을 배치하는 것은 특히 로봇 조작, 자율 주행, 항공 제어와 같은 응용 분야에서 실시간 추론 요구에 의해 제약받고 있습니다. VLA는 일반적으로 이전에 예측을 바탕으로 순차적으로 액션 토큰을 생성하는 자기회귀 디코딩 전략에 의존합니다. 많은 작업에 효과적이지만, 이 자기회귀 복호화 패러다임은 추론 속도를 크게 제한하여 보통 3-5 에 불과합니다 표준 GPU 기반 연구 플랫폼 (예: 단일 고급 소비자용 또는 데이터 센터 GPU) 에서 중단 간 VLA 추론을 위해 배포 할 때 Hz [67]. 이 속도는 반응적이고 안정적인 로봇 작동에 일반적으로 요구되는 제어 주파수보다 상당히 낮으며, 고수준 계획은 수십 헤르츠, 저수준 피드백 제어는 더 높은 업데이트 속도까지 다양하며, 이는 작업 및 하드웨어 플랫폼에 따라 다릅니다. 예를 들어, 로봇 팔이 섬세한 물체를 조작 할 때는 정확성을 유지하고 손상을 방지하기 위해 자주 위치 정보를 업데이트 하는 것이 필수적입니다. OpenVLA와 같은 모델 [122] 그리고 Pi-0 [15] 이 순차적 토큰 생성 방식은 본질적인 어려움에 직면해 동적 환경에서의 효율성이 제한됩니다.

병렬 디코딩과 같은 신흥 솔루션으로, NVIDIA의 GR00T N1 모델이 대표적입니다 [14]. 여러 토큰을 동시에 예측하여 추론을 가속화하는 것을 목표로 합니다. GR00T N1은 전통적인 디코딩 방식보다 약 2.52x 속도 향상을 달성합니다; 하지만 이러한 병렬성은 종종 계층적 부드러움에 상충을 초래하여 로봇의 움직임이 최적이지 않습니다. 이러한 움직임은 정밀함, 적응

Table 4: 도전과제, 잠재적해결책, 그리고 VLA 모델의예상영향.

과제 / 한계	잠재적해결책	기대효과
실시간추론제약	병렬 디코딩, 양자트랜스포머, 하드웨어가속 (예: TensorRT) [129, 122]; 자기회귀오버헤드감소 [75, 142].	자연이중요한영역에서실시간제어및배포를가능하게합니다 [251, 191] (예: UAV, 조작기).
멀티모달액션표현	확산과자기회귀정책을결합한하이브리드토큰화 [171]; 다양한시연과멀티모달산출물에대해훈련합니다 [156].	복잡하고동적인조작에서다중유효한해법모드를통해성능을향상시킵니다 [74].
오픈월드에서의안전보장	동적위험평가모듈 [183, 233]; 저지연비상정지및적응형계획계층 [113].	가정, 공장, 의료등에측불가능한환경에서신뢰성과안전성을향상시킵니다; 사용자수용도를향상시킵니다.
데이터셋편향과그라운드링	다양하고편향없는데이터셋을선별하세요 [185]; 더강한그라운드링 (예: 하드네거티브를이용한 CLIP 미세조정) [282, 17].	공정성과의미정확성을향상시킵니다 [109], 그리고새로운실제입력에대한일반화를강화합니다 [227, 288, 175].
제한된 3D 지각과추론	깊이/LiDAR 통합; 3D 인식아키텍처를개발하고; 포인트클라우드와비전—언어특징을융합한다.	복잡한환경에서조작과내비게이션을위한더강한공간적추론을가능하게합니다 [129].
신체간일반화	다양한형태를거치며훈련합니다; 형태중립적행동추상화를배우고; 도메인간적응적용 [266].	로봇플랫폼및구성간정책이전을용이하게합니다 [279, 122].
주석복잡성및비용	수작업라벨링감소를위한약한감독, 능동적학습, 합성데이터생성 [148].	개발비용을줄이고새로운작업/도메인으로의확장을가속화합니다 [233, 286].
시물레이션-실제 (Sim-to-Real) 전이갭	도메인적응, 시물레이션-실-미세조정, 그리고실제세계보정 [210, 134].	시물레이션이후로배포할때신뢰성과일관성을향상시킵니다 [4, 66].
물리적지식의통합	물리사전, 시물레이션환경, 그리고학습파이프라인에서의동역학모델링 [53].	실제물리적제약조건하에서예측및계획을개선합니다 [106].
멀티모달통합 (촉각, 오디오)	촉각/오디오와시각및언어를융합합니다 [114]; 멀티모달트랜스포머퓨전을확장하세요.	폐쇄/모호성하에서도견고성을높이고작업범위를확장합니다 [76, 136, 91].
장기과제다단계과제	계층적정책, 메모리증강네트워크, 그리고계획계획모듈 [136].	순차적계획, 기억, 구성적실행을향상시킵니다 [131, 286, 227, 175].
시스템통합복잡도	통합트랜스포머백본 [289]; 시간적정렬과시물레이션-실제 (Sim-to-Real) 전송전략 [163, 277].	더엄격한계획수립가능—제어조정및물리적로봇으로의보다견고한이송 [187, 208].
에너지및컴퓨팅수요	가지치기, LoRA, 양자화인식학습, 저전력가속기.	효율적인임베디드/모바일배포가능하게합니다 [252, 283, 125, 246].
보이지않는작업으로의일반화	구성일반화, few-shot 메타러닝, 그리고작업에구애받지않는사전학습 [173, 146].	과적합을줄이고제로/few-shot 적응을강화합니다 [97, 242, 286].
환경변동성에대한견고성	도메인무작위화, 센서융합, 광범위한훈련데이터셋, 온라인재보정 [156].	변화하는조명, 잡동사니, 장면역학상황에서의안정성향상 [285, 242].
윤리적·사회적합의	기기내처리/익명화를 통한프라이버시 [149, 198, 252, 34]; 공정성감사; 규제및신뢰프레임워크.	사회, 의료, 노동분야에서공정하고신뢰할수있는채택을촉진합니다 [160, 180, 217, 172].

성, 유연성이가장중요한수술용로봇과같은민감한응용분야에서는바람직하지않습니다. 따라서출력품질을저해하지않고빠른추론을달성하는것은여전히열린도전과제로남아있습니다.

게다가하드웨어의한계는실시간추론제약을더욱악화시킵니다. 예를들어, 일반적으로 512 차원 400 개이상의비전토큰을포함하는고차원시각임베딩처리는약 1.2 GB/s 메모리대역폭을필요로합니다. 이수요는현재임베디드시스템이나 NVIDIA Jetson 플랫폼과같은엣지 AI 하드웨어의용량을상당히초과하여실질적인배포를제한합니다 [82, 275]. 메모리제약을완화하기위해부동소수점연산의정밀도를낮추는효율적인양자화기법을사용하더라도, 특히양손로봇조작이나의료로봇공학과같이서브밀리미터정밀도가요구되는작업에서는정확도가자주저해됩니다.

4.2. 멀티모달액션표현및안전보장

멀티모달액션표현: 현재 VLA 모델의중요한한계중하나 는특히연속적이고세밀한제어가필요한시나리오에서멀티모달동작을정확히표현하는것입니다 [62, 46]. 256 개의뚜렷한빈으로행동을나누는전통적인이산토큰화방식은본질적으로정밀도가부족하여세밀한로봇잡기나복잡한수술절차와같은세밀한작업에서상당한오류를초래합니다 [171]. 예를들어, 조립작업에서정밀한로봇조작과정에서이산표현은정렬이맞

지않거나부정확한동작을초래하여성능과신뢰성을저해할수 있습니다. 반면, 연속다층퍼셉트론 (MLP) 기반접근법은모드붕괴위험에직면해있습니다 [162, 232] 모델이여러실행가능한경로가있음에도불구하고단일동작계획에조기수렴하는경우입니다. 이로인해매우역동적인환경에서적응적의사결정을위한유연성이줄어듭니다. Pi-Zero 와 RDT-1B 모델로대표되는확산기반정책의신흥 [144] 는다양한액션가능성을포착할수있는더풍부한멀티모달액션표현을제공합니다. 하지만기존트랜스포머기반디코더의약세배에달하는상당한계산오버헤드때문에실시간배포에는실용적이지않습니다. 따라서 VLA 모델은현재밀집된공간에서의로봇내비게이션이나정교한양손조작과같은복잡한동작작업에어려움을겪고있습니다 [74, 247], 여기서여러전략적행동이동등하게유효하고상황에따라달라질수있습니다.

오픈월드의안전보장: VLA 가직면한또다른중요한과제는실제상황에서특징적인역동적이고예측불가능한환경에서견고한안전성을보장하는것입니다 [39, 274]. 현재많은구현은미리정의된하드코딩된힘과토큰계값에크게의존하여, 예기치못한장애물이나갑작스러운환경변화와같은새로운조건에대비한적응력을크게제한합니다 [156]. 충돌에측에서사용되는모델은일반적으로복잡하고역동적인공간에서약 82%의정확도에그치는경우가많은안전마진이가려없는참고물류나가정용로봇같은응용분야에서심각한위험을초래합니다

다 [288, 122]. 더불어, 긴급정지와같은필수안전메커니즘은 포괄적인안전검증으로인해 200에서 500 밀리초사이의상당한지연시간을포함합니다 [170, 122]. 지연은겉보기에는 미미해보일수있지만, 고속운행이나자동화운전, 긴급로봇대응과같은중요한개입상황에서는위험할수있습니다.

4.3. 데이터셋편향, 그라운드, 그리고보이지않는과제로의일반화

VLA 모델의효과를제한하는주요장애물은데이터셋편향과그라운드결합의만연한존재입니다. 현재학습데이터셋은 주로웹크롤저장소에서수집되며, 종종내재된편향을보입니다 [214, 119]. 연구에따르면약 17% 표준데이터셋내연관성은 ‘의사’ 와같은용어를남성인물과불균형적으로연관짓는 등고정관념적해석에치우쳐있습니다 [222, 124]. 이러한편향은훈련을통해전파되어, 다양한환경에서사용될때의미적으로어긋나거나매락상부적절한반응을내는 VLA 를초래합니다. 예를들어, OpenVLA 와같은모델은새로운환경에서객체참조어의약 23% 를누락하는것으로보고되어, 명령어의정확한해석이중요한실제응용분야에서는실용성을크게제한합니다 [122]. 이그라운드문제는조합일반화에서도확장되는데, VLA 는훈련말뭉치에서“yellow horse” 같은구절을해석하는 등드물거나비전통적인조합을만날때종종실패하는경우가 많습니다. 이러한한계는신중하게선별되고균형잡힌포괄적도메인특화데이터셋의긴급한필요성을강조하며, 다양한맥락에서편향을완화하고의미적정렬을강화하기위한고급그라운드알고리즘과결합되어야합니다.

데이터셋편향이제기하는도전과더불어, VLA 의실질적배포에중요한장벽인보이지않는작업으로의일반화문제가더 넓게존재합니다. 기존모델은익숙한환경이나훈련시나리오와유사한작업에서속련도를보여도, 성능이크게저하되어최대 40% 까지저하되는경우가 많으며, 완전히새로운과제나낯선변형을마주칠때입니다. 예를들어, 가정용작업에특화된 VLA 는산업또는농업환경에도입될때주로객체유형, 환경역학, 운영제약등의차이로인해어려움을겪거나실패할수있습니다. 이한계는주로좁은범위의훈련분포에과적합되고다양한과제표현에대한노출이부족하기때문입니다. 따라서현재 VLA 는제로샷또는 few-shot 학습시나리오에서일반화가제한적이어서적응성과확장성을저해합니다.

4.4. 시스템통합복잡성과계산요구사항

고수준인지계획 (시스템 2) 과실시간물리제어 (시스템 1) 를결합한이중시스템아키텍처내에 VLA 모델을통합하는것은로봇응용에서상당한복잡성을야기합니다. 주요도전과제는이두시스템간의시간적불일치에서발생합니다. 일반적으로 System 2 는복잡한작업분해와전략적기획을위해 GPT 나 LLaMA-4 와같은 LLM 을활용합니다. 이러한모델은상당한계산요구량때문에종종다음과같은수준의추론지연을겪습니다 ~800 표준 GPU 기반추론플랫폼 (예: 단일고급소비자용 또는데이터센터 GPU) 에서실행될경우 MS 이상을기록하며, 이는 LLM 배포에흔히사용됩니다. 반대로, 시스템도 저수준 모터실행을담당하는 1 구성요소는일반적으로실시간 CPU, 마이크로컨트롤러또는전용로봇컨트롤러에구현된엄격하게제한된제어루프내에서실행되며, 플랫폼과작업에따라수밀리초단위의업데이트간격이흔히발생합니다. 이러한극명한 운영주기차이는동기화에어려움을초래하여지연과잠재적으로비효율적인실행경로를초래합니다. 예를들어, NVIDIA 의

GR00T N1 모델은이두시스템의효과적인통합을보여주지만, 비동기상호작용으로인해움직임중가끔씩끊김이발생해이본질적인도전과제를부각시킵니다.

더욱이, Vision Transformers(ViT) 와같은고차원비전인코더와저차원액션인코더간의특정공간불균형은통합복합성을악화시킵니다. 이러한이질적인임베딩을조화시키려할때, 각각이해와실행가능한명령간의일관성이크게저하될수있습니다. OpenVLA [122] 그리고 RoboMamba [143] 트랜스포머기반시각처리와후속액션인코딩을활용하는이통합문제는시물레이션환경에서물리적하드웨어배포로이식할때성능저하로 이어지는문제를보여줍니다. 이러한차이는주로시물레이션된동역학과실제센서노이즈또는보정문제간의불일치로인해성능저하로이어질수있습니다 [122, 49, 92].

에너지및컴퓨팅수요는특히자율드론, 모바일로봇, 웨어러블로봇시스템등է지컴퓨팅환경에서 VLA 도입에도다른중요한장벽을형성합니다. 고급 VLA 에서흔히볼수있는방대한매개변수수 (예: 70 억개이상의매개변수를가진모델) 는기본형태에서 28GB 를초과하는 VRAM 의연산자원을요구합니다. 이러한요구사항은대부분의최신է지지향프로세서와 GPU 의성능보다훨씬높아, 특수고자원환경외부에서정교한 VLA 의실용적적용을제한합니다.

4.5. VLA 배포의견고성과윤리적도전과제

VLA 모델의실제배치가가장큰장애물은환경변동성에대한견고성이제한적이라는점에있으며, 이는중요한윤리적및안전문제를제기합니다. 환경적견고성은동적으로변화하고부분적으로관찰가능한조건하에서시스템이신뢰할수있는인식, 추론, 액션생성을지속할수있는능력의의미합니다. 실제로실제환경은조명변화, 악천후, 센서잡음, 물체가림과같은요인으로상당한불확실성을초래합니다.

경험적증거는여러 VLA 구성요소에서이러한한계를강조합니다. 예를들어, OpenDriveVLA 와같은시스템에서사용되는비전모듈 [293] 정확도가약 20–30% 저하되는현상저대비또는그림자가지배적인장면에서현재시각인코더의민감한조명조건을반영합니다. 마찬가지로, CoVLA 와같은 VLA 에서의언어이해도마찬가지입니다 [5] 음향적으로잡음이크거나의미가모호한환경에서악화되며, 명령어오해가잘못된동작실행으로이어질수있습니다. 조작중심시나리오에서는 RoboMamba 와같은 VLA 지원로봇시스템이있습니다 [143] 어수선한환경에서어려움을겪으며, 부분적으로가려진물체의자세나방향을자주잘못판단하여작업성공률을떨어뜨립니다.

이러한견고성제한은안전에중요한배포에서직접적인윤리적영향을미치는데, 실제변동성하에서성능저하가의도치않은동작, 신뢰성저하, 사용자신뢰상실로이어질수있기때문입니다. 따라서견고성문제는단순한기술적도전이아니라인간중심환경에서 VLA 시스템을책임감있고윤리적으로배포하기위한전제조건이기도합니다.

5. 토론

그림에나와같이 17VLA 모델은알고리즘, 계산, 윤리적차원에걸쳐다면적인도전과제에직면해있습니다. 첫째, 자원이제한된하드웨어에서실시간추론을달성하는것은자기회귀인코더의순차적특성과멀티모달입력의고차원성때문에여전히 어렵습니다. 둘째, 비전, 언어, 행동을일관된정책으로융합

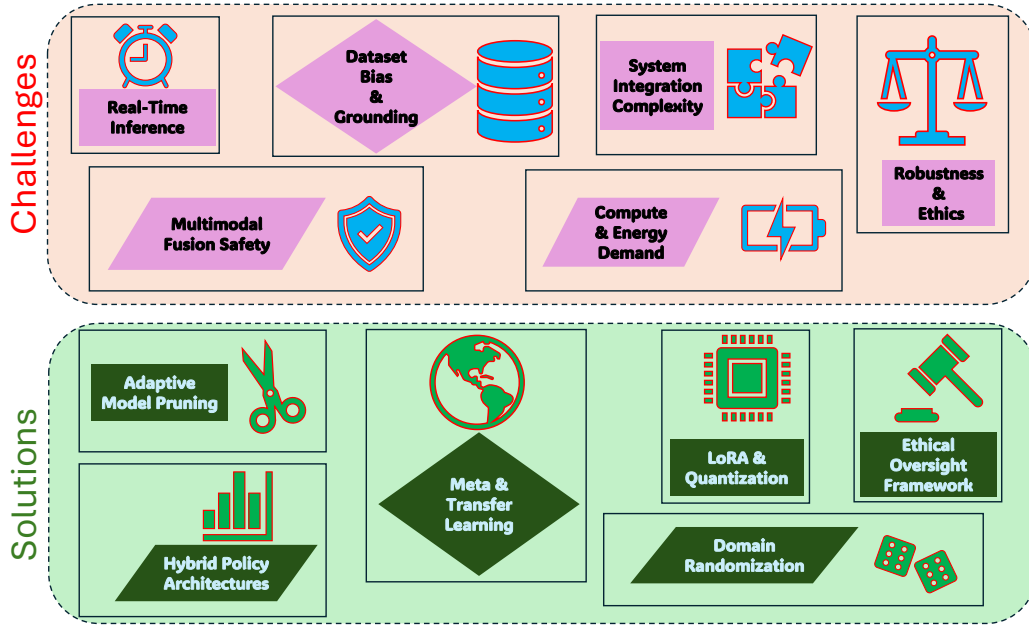


Figure 17: 그림은여섯가지핵심 VLA 도전과제를지도화합니다: 즉, 실시간추론, 멀티모달융합안전성, 데이터셋편향, 통합복잡도, 계산요구, 강인성/윤리성 등이 포함됩니다: 적응형가지치기, 하이브리드정책아키텍처, 메타/전이학습, LoRA/양자화, 도메인무작위화, 윤리적감독등 6 가지목표해결책에대항합니다. 이체계적인정렬은더넓은실제로봇영역에서건고하고효율적이며안전한 VLA 배포경로를명확히합니다.

하면예상치못한환경변화에직면했을때안전취약점을초래합니다. 셋째, 데이터셋편향과그라운드오류가일반화를저해하여종종분포외작업에서모델이실패하는원인이됩니다. 넷째, 다양한구성요소인인식, 추론, 제어를통합하면최적화및 유지가어려운복잡한아키텍처가만들어집니다. 다섯째, 대규모 VLA 시스템의에너지및컴퓨팅요구는임베디드또는모바일플랫폼에서의배포를방해합니다.

마지막으로, 환경변동성에대한견고성이제한되면안전하지않거나신뢰할수없는행동이발생할수있으며, 이는안전보장, 책임성, 프라이버시, 편향완화와관련된윤리적·규제적우려를야기합니다. 이러한한계들은실제로봇공학, 자율시스템, 인터랙티브응용분야에서 VLA 모델의실질적채택을제한합니다. 이러한도전에대한잠재적해결책은아래에서논의됩니다.

5.1. 잠재적해결책

- **실시간추론제약.** 향후연구는지연시간, 처리량, 작업특화의정확성을조화시키는 VLA 아키텍처를개발해야합니다. 유망한방향중하나인 FPGA 기반비전프로세서와희소행렬연산에최적화된텐서코어와같은특수하드웨어가속기의통합으로, 서브밀리초규모의합성곱및변환계층을실행하는것입니다 [122, 129]. 로우랭크적응 (LoRA) 과같은모델압축기법 [93] 지식증류는매개변수수를최대 90 까지줄일수있습니다% 메모리사용량과추론시간을모두줄이면서도 95 이상을유지합니다% 벤치마크작업에서의원래성과에대한평가를받았습니다. 혼합정밀도산술 (예: FP16/INT8) 과블록별보정을결합한점진적양자화전략은정확도손실을최소화하면서계산량을 2-4x 더줄일수있습니다 [121]. 입력복잡도에따라네트워크깊이또는폭을동적으로조절하는적응형추론아키텍처, DeeR-VLA 의조기종료분기과유사합니다 [269] 시각적장면이나언어명령이단순할때트랜스포머계층을

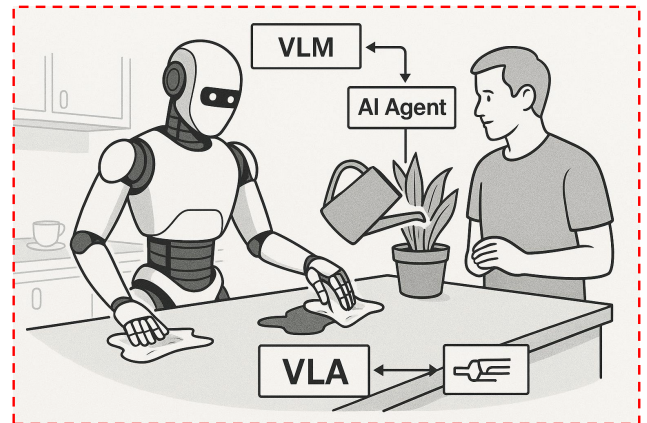


Figure 18: 이개념적일러스트는시각언어모델 (VLM), VLA 프레임워크, 에이전트 AI 시스템으로구동되는미래의휴머노이드어시스턴트”Eva” 를소개합니다. VLM 은의미론적장면이해와객체어포던스에측을가능하게하며, VLA 는언어기반명령어를계층적운동계획으로번역합니다. 에이전트 AI 모델은개방형환경에서적응형학습, 자기정제, 상호작용적의사결정을보장합니다. 이구성요소는지각, 언어이해, 계획, 안전한자율행동이현실세계의사회적인식과제에서융합되는로봇공학내인공일반지능 (AGI) 의기초청사진을제공합니다.

선택적으로우회하여평균계산량을줄일수있습니다. 마지막으로, 서브워드패치임베딩과동적어휘할당을활용한효율적인도큰화방식은시각적·언어적입력을압축하여의미론적풍부함을희생하지않으면서도큰수를최소화할수있습니다 [171]. 이러한혁신들은일반엡지 GPU 에서 50ms 미만의종단간추론을가능하게하여, 자율드론비행, 실시간원격조작, 협업제조등지연에민감한응용분야로발전할수있습니다.

- **멀티모달액션표현및안전보장.** 멀티모달액션표현과전

고한안전성을 해결하려면 엄격한 안전 제약 하에서 시각, 추론, 제어를 통합하는 중단간 프레임워크가 필요합니다. 저수준 모션원시요소를 위화환산 기반 샘플링을 결합한 하이브리드 정책 아키텍처 [40] 자기회귀적이고 수준 계획자와 함께 [242] 다양한 액션 계획을 압축된 확률적 표현으로 하여 동적 환경에서의 적응력을 향상시킵니다. 안전은 시각, 깊이, 고유수용성 데이터를 포함한다 중 센서 융합 스트림을 실시간으로 수집하여 충돌 확률과 관절 스트레스 임계값을 예측하고, 미리 정해진 안전한 계를 위반할 때 비상 정지 회로를 작동시키는 실시간 위험 평가 모듈을 통해 강화할 수 있습니다 [183, 233]. 제약 최적화가 추가된 강화 학습 알고리즘 (예: SafeVLA 의 라그랑지안 방법) [274] 안전 제약을 엄격히 준수하면서 작업 성공률을 극대화하는 정책을 학습할 수 있습니다. 규칙 기반 RL (GRPO) 과 직접 선호 최적화 (DPO) 와 같은 온라인 모델적응 기법은 새로운 환경 조건에서 행동 선택을 더욱 정교하게 조정하여 시나리오 전반에 걸쳐 일관된 안전 성능을 보장합니다 [113]. 무엇보다도, 실행 전에 계획 출력을 상징적으로 분석하는 형식적 검증 계층을 내장하면 신경망 기반 제어 장치에서도 안전 불변량 준수를 보장할 수 있습니다. 이러한 방법론을 통합하면 복잡하고 멀티모달 동작을 실행할 뿐만 아니라 비구조화된 실제 환경에서 안전성이 입증된 VLA 시스템을 만들어낼 것입니다.

- **데이터셋 편향, 그라운드, 그리고 보이지 않는 과제에 대한 일반화.** 견고한 일반화는 확장된 데이터 다양성과 고급 학습 패러다임을 모두 요구합니다. LAION-5B 와 같은 웹 규모 이미지-텍스트 쌍을 결합한 대규모 편향 없는 멀티모달 데이터셋을 큐레이션합니다 [194] Open X-Embodiment 와 같은 로봇 중심 객체 아카이브를 포함해 [227] 이는 공정한 의미론적 그라운드를 위한 토대를 마련합니다. 하드네거티브 샘플링과 비전-언어 골격의 대조적 미세 조정 (예: CLIP 변형) 은 허위 상관 관계를 완화하고 의미 충실도를 향상시킬 수 있습니다 [17, 282]. 메타러닝 프레임워크는 시각 언어로봇 내비게이션 모델에서 입증된 바와 같이 과제 공간에 공유된 사전 학습을 통해 새로운 과제에 빠르게 적응할 수 있도록 합니다 [175]. 재생 버퍼와 정규화 전략을 가진 연속 학습 알고리즘은 기존 지식을 보존하면서 새로운 개념을 통합하여 VLA 모델에서의 치명적인 망각 문제를 해결합니다 [48]. 3D 지각 영역에서의 학습 전이 (예: 3D-VLA 에서의 포인트 클라우드 추론) [288] 는 모델에 더 강한 공간 귀납 편향을 부여하여 분산 외 상황에 대한 견고성을 향상시킬 수 있습니다. 마지막으로, 시뮬레이션-투리얼 (sim2real) 미세 조정과 도메인 무작위화, 동적 조명, 텍스처, 물리 변화 같은 실제 보정을 통해 합성 환경에서 배운 정책이 물리적으로 보에 효과적으로 전달되도록 보장합니다 [4, 66]. 이러한 통합 전략들은 VLA 가 실제 배치에서 보이지 않는 물체, 장면, 과제에 대해 자신 있게 일반화할 수 있도록 힘을 실어줄 것입니다.
- **시스템 통합 복잡도와 계산 요구.** 좁은 컴퓨팅 예산 하에서 멀티모달 파이프라인의 복잡한 조정을 관리하기 위해, 연구자들은 모델 모듈화 와 하드웨어-소프트웨어 공동 설계를 수용해야 합니다. 저랭크 적응 (LoRA) 어댑터는 사전 학습된 트랜스포머 계층에 주입할 수 있어, 코어가 중지를 변경하지 않고도 작업별 미세 조정이 가능합니다 [93]. 대규모 '교사' VLA 에서 경량 '학생' 네트워크로 지식을 정제하면서, 학생이 교사의 중간 표상과 행동 분포를 맞추도록 유도하는 상호 정보 기반 목표에 따라 5-10 규모의 간결한 모델을 만듭니다 $\times$ 매개 변수는 적으면서도 90-95 를 유지한다 % 과제 수행에 관한 내용 [121]. 혼합 정밀도 양자화와 양자화 인식 훈련을 결합하면 가중치를 4-8 비트로 압축하여 메모리 대역폭과 에너지 소비를 60 비트 이상 줄일 수 있습니다 % [122]. VLA 워크로드에 맞춘 하드웨어 가속기는 희소 텐서 연산, 동적 토큰 라우팅, 융합된 비전-언어 커널을 지원하며, 20-30W 전력 범위 내에서 지속적 인 100+ TOPS 처리량을 제공함으로써 임베딩으로봇 플랫폼의 요구를 충족할 수 있습니다 [171, 242]. TensorRT-LLM 과 같은 툴체인 [129] TVM 은 특정 엣지 장치에 맞춰 중단간 VLA 그래프를 최적화하여 계층을 융합하고 정적 부분 그래프를 사전 계산할 수 있습니다. TinyVLA 와 같은 신흥 아키텍처는 1B 이하 매개 변수 VLA 가 실시간 추론을 통해 조작 벤치마크에서 최첨단 성능을 달성할 수 있음을 보여주며, 자원이 제한된 환경에서 광범위한 배포의 길을 개척합니다.

록유도하는 상호 정보 기반 목표에 따라 5-10 규모의 간결한 모델을 만듭니다 $\times$ 매개 변수는 적으면서도 90-95 를 유지한다 % 과제 수행에 관한 내용 [121]. 혼합 정밀도 양자화와 양자화 인식 훈련을 결합하면 가중치를 4-8 비트로 압축하여 메모리 대역폭과 에너지 소비를 60 비트 이상 줄일 수 있습니다 % [122]. VLA 워크로드에 맞춘 하드웨어 가속기는 희소 텐서 연산, 동적 토큰 라우팅, 융합된 비전-언어 커널을 지원하며, 20-30W 전력 범위 내에서 지속적 인 100+ TOPS 처리량을 제공함으로써 임베딩으로봇 플랫폼의 요구를 충족할 수 있습니다 [171, 242]. TensorRT-LLM 과 같은 툴체인 [129] TVM 은 특정 엣지 장치에 맞춰 중단간 VLA 그래프를 최적화하여 계층을 융합하고 정적 부분 그래프를 사전 계산할 수 있습니다. TinyVLA 와 같은 신흥 아키텍처는 1B 이하 매개 변수 VLA 가 실시간 추론을 통해 조작 벤치마크에서 최첨단 성능을 달성할 수 있음을 보여주며, 자원이 제한된 환경에서 광범위한 배포의 길을 개척합니다.

- **VLA 배포의 환경 변동성에 대한 견고성.** 실제 환경에서 견고한 VLA 성능을 보장하려면 환경 불확실성과 장기적 인 시스템 드리프트를 다루기 위한 목표지향적인 기술적 개선이 필요합니다. UniSim 의 폐쇄 루프 센서 시뮬레이션과 같은 도메인 무작위화 및 합성 데이터 증강 파이프라인은 조명, 폐쇄, 센서 노이즈에서 사실적인 변화를 생성하여 분포 변화에 대한 회복력을 향상시킵니다 [264]. 또한, 실시간 피드백을 기반으로 인지의 임계값을 동적으로 조정하고 이득을 제어하는 적응형 재보정 모듈은 센서 노후화나 작동 조건 변화로 인한 성능 저하를 완화할 수 있습니다. 이러한 접근법들은 다양하고 변화하는 배치 시나리오에서 VLA 시스템의 안정성과 신뢰성을 향상시키는 것을 목표로 합니다.
- **VLA 배포 시 윤리적, 개인정보 보호, 사회적 고려 사항.** 기술적 견고성을 넘어, VLA 시스템 도입은 거버넌스 지향적 해결책이 필요한 중요한 윤리적·사회적 도전을 제기합니다. 편향 감사 도구 가 학습 데이터에서 왜곡된 인구 통계학적 또는 의미적 분포를 식별하고, 그 다음에는 적대적 편향 해소 및 반 사실 데이터 증강과 같은 교정 전략이 필요합니다 [185, 282]. 기기 내 처리, 민감한 데이터 스트림에 대한 동동형 암호화, 교육 중차 등 프라이버시 보호 메커니즘은 의료 및 스마트 홈 등 분야에서 사용자 데이터를 보호하는데 매우 중요합니다 [160, 180]. 더 나아가, 투명한 영향 평가, 이해관계자 참여, 인력 역량 강화 이니셔티브는 사회적 제적 영향을 관리하는데 도움이 되며, 규제 체계와 업계 표준은 책임성과 책임 있는 VLA 채택을 보장하는데 필수적입니다.

5.2. 미래 로드맵

VLA 기반 시스템의 미래는 점점 더 강력한 멀티모달 기반, 에이전트 추론, 그리고 구현된 지속적 학습이 교차하는 지점에서 진화할 것으로 예상됩니다. 앞으로 10 년간 우리는 VLA 가 유능하지만 취약한 작업 전문가에서 신뢰할 수 있고 범용적인 로봇 지능으로 전환하는 여러 수렴 추세 가 나타날 것으로 예상합니다. 그러나 이 궤적은 앞서 강조된 지속적인 제약 및 한계에 의해 형성될 것입니다: (i) 폐쇄 루프 제어에서의 실시간 추론 병목, (ii) 불완전한 멀티모달 동작 표현과 약한 안전성 보장, (iii) 분포 이동 시 데이터셋 편향 및 그라운드 실패, (iv) 인지-메모리-추

론-제어전반에걸친통합복잡성, (v) 엣지배포를방해하는높은컴퓨팅및에너지요구, 그리고 (vi) 견고성과투명성, 그리고 오픈월드환경에서의윤리적문제들. 이문제들을전체적으로해결하기위해, Figure 19 시스템수준의연구로드맵을요약하며, 그림 18 VLM, VLA 아키텍처, 에이전트 AI 모듈이로봇공학에서구현된 AGI 로어떻게공진화할수있는지에대한직관적인개념적예시를제공합니다.

멀티모달기초모델은체화된지각의 ‘피질’ 으로보입니다.. 오늘날의 VLA 스택은중중비언어백본과작업별정책헤드에의존하여일반지식의재사용이제한되고도메인간재학습비용이증가합니다. 그럴듯한다음단계는웹규모의이미지, 비디오, 텍스트, 상호작용/어포던스흔적을기반으로훈련된통합멀티모달기초모델을구축하여정적의미론뿐아니라역학, 접촉사전, 상식적물리적지식까지인코딩하는공유 ‘피질’ 로가능하는것입니다 [295, 289, 272]. 이러한피질은언어를객체중심표현과지속적인장면구조에그라운딩함으로써피상적상관관계로인한고장모드를줄일수있습니다 [206]. 그림 18에서강조된바와같이, 이기초모델피질은로봇이환경을실행가능한개체 (객체, 영역, 어포던스) 로분할하고, 하위계획자와컨트롤러에게안정적인의미적앵커를제공할수있게할것입니다. 그러나과신과환각적그라운딩을방지하기위해, 이러한모델은보정된불확실성과증거연계추론을포함하여인식기반계획이가림, 혼잡, 모호한지시아래에서도검증가능하도록해야합니다 [145].

에이전트형, 자기감독적, 평생학습및지속적인적응.. 현재 VLA 의결정적인한계는정적인특성입니다: 한번훈련된정책은비고정환경에서작동하더라도변경되지않고배포됩니다. 미래의 VLA 는모델이탐색목표를제안하고, 결과를가설하며, 시뮬레이션및실제물아웃을통해자기수정하는에이전트학습루프를채택하여수개월또는수년에걸쳐지속적인기술향상을가능하게해야합니다 [43, 90]. 이방향은분포이동, 데이터셋편향, 장기과제취약성을완화하여모델이시간이지남에따라적응할수있도록할것으로기대됩니다; 하지만지속적인정책업데이트는새로운위험을초래합니다. 특히, 반복되는온라인또는점진적학습은이전에습득한역량을덜어쓰는치명적망각, 의도치않은행동회귀를유도하고, 학습된정책을손상시킬수있는소음적이고적대적이거나의도치않은환경피드백에취약해질수있습니다 [292, 263]. 따라서평생학습은재생및안전인식업데이트, 모듈형어댑터, 검증기반정책개정과함께이루어져야합니다 [121]. 그림 19의로드맵에서, 이에이전트형평생학습패러다임은자연스럽게 신뢰가능하고안전한지능과 통합시스템 & 거버넌스의교차점에위치하며, 지속적인학습을임시로미세조정하는단계가아니라통제되고감사가가능한생애주기과정으로다루는것입니다.

확장성과해석가능성을위한계층적이고신경상징적인계획.. 저수준모터프리미티브에서장기과제목표로확장하려면명시적인계층구조가필요합니다 [261, 265]. 차세대 VLA 시스템은언어기반플래너 (LLM 스타일모듈, 유연성과제약조건에맞게미세조정됨) 를사용할가능성이높으며, 목표를구조화된하위작업으로분해한뒤, 중간수준의기술정책과준수동작을보장하는저수준컨트롤러로연결합니다 [94, 253]. 이러한신경-상징적혼합은디버깅, 모니터링, 인증이더쉬운인터페이스를부여함으로써통합복잡성을연결하는데도움을줍니다 [201]. 중요한점은계층구조가선택적검증을가능

하게한다는것입니다: 상위계획은제약위반 (안전하지않은단계, 금지영역) 을확인할수있습니다 [192, 71] 저수준계층은제어배리어기능, MPC, 런타임안전모니터로차폐할수있습니다 [188, 196]. 이구성요소들은그림 19의 신뢰가능하고안전한지능축과일치합니다. 여기서안전성은사후평가가아닌계획시간제약과실행시간가드를통해모두강화됩니다 [101, 274].

세계모델과물리적/인과적추론을통한실시간적응.. 비구조화된환경에서견고한배포를위해서는 VLA 가객체, 접촉, 동적요소에대한내부예측모델을유지해야합니다. 단기상태전이와실패가능성을예측하는세계모델은반사실적평가 (‘여기서말면무엇이충돌하는가?’) 와현실이기대에서벗어날때 (예: 미끄러짐, 예상치못한마찰) 신속한교정조치를지원할수있습니다 [203]. 이기능은안전한조작, 항법, 인간-로봇 상호작용의핵심으로, 작은오류가빠르게누적되는부분입니다 [156]. 하지만세계모델은답재시충분히효율적이어야하며, 다중센서증거와일치해야합니다. 따라서하드웨어인식과메모리효율적인예측모델링 (예: 시간토큰압축) 이주요미래방향입니다 [267, 216], 이벤트기반상태업데이트 [250, 226], 그리고미분가능한물리시뮬레이터와학습된동역학을결합하여제어관련업데이트속도를가능하게하는하이브리드물리학습모델입니다 [102, 52]. 그림 19에서이러한필요성은 효율적인배포축 (실시간제약) 과 신뢰가능하고안전한지능축 (그라운딩된의사결정을위한물리-인과추론) 에함께제시됩니다.

효율성과확장성: 일반성과엣지배포의다리.. VLA 채택의중심장벽은대형멀티모달백본의계산용량과폐쇄루프제어의지연/에너지제약간의불일치입니다 [268, 238]. 향후 VLA 는일반화를유지하면서추론비용을줄이면서매개변수효율적인설계 (구조적희소성, 저랭크적응, 모듈러전문가) 를우선시해야합니다 [65, 197]. 훈련시간효율성을넘어, 언제/조기종료정책은계산을적응적으로할당하여안전중요단계가높은충실도를유지하는반면, 일상단계는더저렴한경로를사용할수있습니다 [41, 274]. 동등하게중요한것은액션공간효율성입니다: 컴팩트한액션토큰화확정된제어표현은자기회귀구간을단축하여시간적부드러움을희생하지않으면서도더높은제어속도를가능하게합니다 [171]. 이러한모델수준의선택은 GPU/NPU/엣지가속기전반에걸친하드웨어인식컴파일과함께이루어져야하며, 양자화인식스케줄링과메모리최적화주의커널등이포함됩니다 [65, 167]. 컴퓨팅인식개성과에피소드메모리는중복순방향패스를줄이고장기과제작업에서반응성을향상시킵니다 [27, 133]. 이러한방향은그림 19의 효율적인배포축을구체화하며앞서강조한컴퓨팅/에너지한계를직접해결할수있습니다.

신체간전이및형태중립적기술표현.. 각로봇형태에대해별도의 VLA 를훈련시키는방식은확장성이낮을가능성이큽니다. 미래의핵심주제는형태중립적정책학습으로, 기술이추상적행동공간 (예: 접촉목표, affordance-point 조작, 작업공간제약) 에서표현되어바퀴달린플랫폼, 네발보행, 휴머노이드간에전이됩니다 [275, 118]. 메타러닝과 few-shot 보정은몇주간의훈련이아닌몇분의데이터로신형로봇을빠르게부트스트랩할수있게해줍니다 [35]. 이방향은구현체와환경간에불변성을강제함으로써데이터셋편향을완화하지만, 표준화된표현, 공통인터페이스, 재현가능한평가프로토콜을요구합니다 [118, 275]. 그림 19에서이형태중립적기술학습패러다임

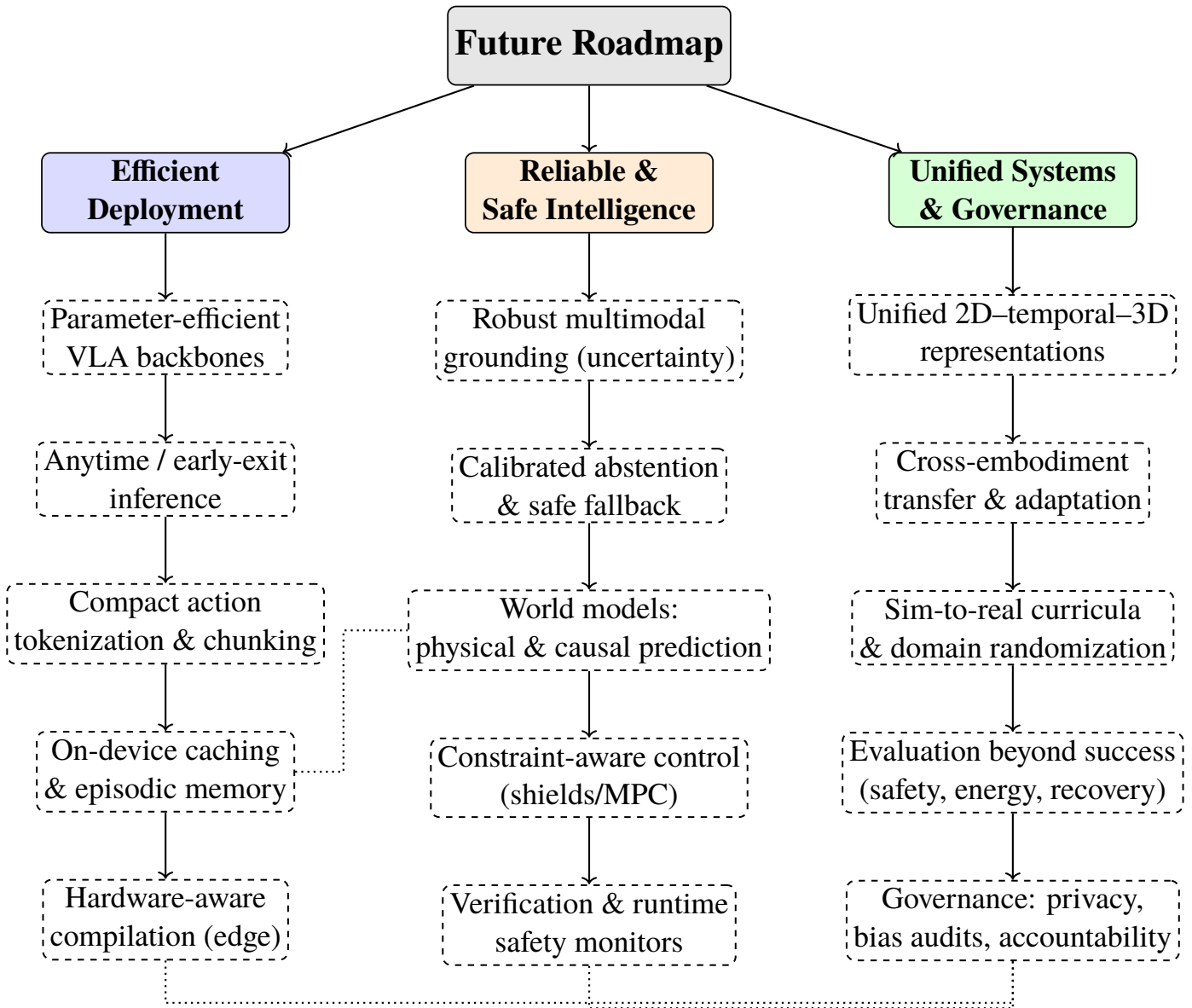


Figure 19: 비전-언어-액션 (VLA) 모델의향후연구로드맵 (그림 19). 점선상자는이작업에서확인된 VLA 모델의핵심한계를해결하기위한우선순위방향을나타냅니다: (i) 효율적인배포 (지연, 계산, 에너지), (ii) 신뢰가능하고안전한지능 (그라운드링, 불확실성, 안전보장), 그리고 (iii) 통합시스템 & 거버넌스 (2D-시간-3D 통합, 이전, 평가및책임있는배포).

은 통합시스템 & 거버넌스에속하며, 아키텍처통합과원칙적전이및벤치마킹을연결합니다.

과제성공을넘어선평가: 안전, 회복, 자원인지지표: VLA의발전은배포현실을반영하는측정이필요합니다. 과제성공만으로는실패의심각성, 시간적불일치, 위험한근접실패, 에너지비효율을명확히나타내지않습니다. 향후벤치마크는인간의제약하에서안전위반, 불확실성보정, 회복행동, 시간적일관성, 에너지소비, 그리고다운스트림효율을정량화해야합니다 [274, 78]. 더불어, 평가는공정한비교와편향기반성과를진단하기위해컴퓨팅예산, 데이터셋구성, 배포조건을보고해야합니다 [215, 273, 118]. 이러한측정은단순한과학적위생이아니라감사와거버넌스의기초이며, 시스템이대규모로책임감있게배포될수있는지를결정합니다 [274, 275]. 이동기는그림평가분기의근간입니다 19 그리고그림에나오는개념

적배치서사를보완합니다 18.

안전, 윤리, 인간중심의정렬을일류설계목표로삼는것: VLA 가자율성을얻으면서내장된안전성과가치정렬이매우중요해집니다. 향후시스템에는고위험행동을실행하기전에잠재적피해를평가하는실시간위험추정기를통합하고, 모호성있는자연어확인을요청하며, 책임성을위한투명한로그를유지해야합니다 [274, 118]. 특히보조로봇과안전에중요한자율성에있어프라이버시인식감지, 편향감사, 그리고인간개입 (Human-in-the-loop) 감독이생명주기에통합되어야합니다 [284, 235]. 규제부합의평가프로토콜과표준화노력이 VLA 의진전을신뢰할수있는실제시스템에적용하는데필수적입니다 [169, 294, 234]. 이거버넌스관점은그림에서명확히담겨있습니다 19 그리고이는그림에묘사된사회적인식이강한인간형배경에암묵적으로반영된다 18.

교차주제: 지속적인학습, 고장복구, 상호작용, 그리고제어 충실도.: 그림에나오는모든축 19 여러교차주제들이향후 10 년간 VLA 연구를형성할것으로예상됩니다. 첫째, 지속적이고평생학습은안전하고감사가가능하며미치지않아야하며, 배포된시스템을불안정하게하지않으면서장기적인적응을가능하게해야합니다 [290]. 둘째, 실패감지및복구는내성적모니터링, 불확실성인식, 실행결과와다르게발생할때구조화된복구행동을포함하는일류역량으로다뤄져야합니다 [116, 259]. 셋째, 액션생성의정밀도와신뢰성향상이여전히중요합니다. VLA 기반플래너는유연하고언어조건화된사결정을가능하게하지만, 궤적정확도와제어안정성은현재모델예측제어, 샘플링기반모션계획, 피드백선형제어등기존의분석접근법에비해뒤쳐져있습니다. 그결과, VLA 기반고수준계획과고전적또는학습된저수준제어기를결합한하이브리드아키텍처는의미론적유연성과제어수준의정밀도를모두달성하는데중심적인역할을할것으로기대됩니다. 마지막으로, 인간의정령과상호작용은의도명확화, 공유자율성, 설명가능한행동그라운딩을위한메커니즘을필요로하며, 이는다양한실제환경에서신뢰와사용성을지원합니다 [111, 110].

요약하자면, 그림 19 실험실시연과견고한실제배치간의격차를해소하려면효율성, 안전성, 데이터및일반화, 시스템통합, 평가및거버넌스에서조정된발전이필요함을강조합니다. 보완적으로그림 18 이러한발전이모바일로봇, 조작기, 보조시스템, 휴머노이드등다중플랫폼에서멀티모달지각, 계층적계획, 지속적인적응, 인간친화적안전이통합된정보스택내에통합된범용체화된에이전트로수렴할수있음을보여줍니다. 이방향들을함께해결하면 VLA 가유망한연구프로토타입에서단일로봇형태에국한된솔루션이아닌신뢰할수있고폭넓게적용가능한내장시스템으로변화할것으로기대됩니다.

6. 결론

이종합리뷰에서우리는지난 3 년간발표된비전-언어-액션(VLA) 모델의최근발전, 방법론및적용을체계적으로평가했습니다. 우리의분석은 VLA 의기초개념에서시작하여, 물리적또는시뮬레이션환경에서시각지각, 자연어이해, 액션생성을통합하는멀티모달시스템으로서이들의역할을정의했습니다. 우리는이들의진화와연대를추적하며, 고립된지각-행동모듈에서완전히통합된지시를따르는로봇에이전트로의전환을나타내는주요이정표들을상세히설명했습니다. 우리는멀티모달통합이느슨하게결합된파이프라인에서양식간원활한조정을가능하게하는트랜스포머기반아키텍처로성숙해졌다는점을강조했습니다.

다음으로, VLA 가행동원시및공간의를포함한시각적및언어적정보를어떻게인코딩하는지에초점을맞춰토론화와표현기법을살펴보았습니다. 학습패러다임을탐구하며, VLA 성능을형성한학습대상학습, 모방학습부터강화학습, 멀티모달사전학습에이르기까지다양한데이터셋과훈련전략을상세히설명했습니다. ‘적응형제어및실시간실행’ 섹션에서는현대 VLA 가동적환경에최적화되고, 지연시간에민감한작업을지원하는정책을분석했습니다. 그후주요아키텍처혁신을분류하여 50 개이상의최근 VLA 모델을조사했습니다. 이논의에는모델설계, 메모리시스템, 상호작용충실도의발전이포함되었습니다. 또한 LoRA, 양자화, 모델가치치기와같은매개변수효율적인방법과병렬디코딩, 하드웨어인식추론과같은가속기법을포함한훈련효율성향상전략을연구했습니다. 실제응용분석은휴머노이드로봇, 자율주행차, 산업자동화, 의

료, 농업, 증강현실(AR) 내비게이션등 6 개분야에서 VLA 모델의가능성과현재의한계를모두강조했습니다. 이러한환경에서 VLA 는고수준의미론적추론, 지시-따르기, 과제일반화에서특히구조화된또는부분적으로통제된환경에서강력한능력을보여주었습니다. 그러나실시간추론지연, 환경변동성에서의견고성제한, 그리고기존의분석계획및제어파이프라인에비해장기또는안전에중요한제어의정밀도가떨어지는점이효과적으로제약을받았습니다. 더불어, 안정적인성능을달성하기위해애플리케이션별적응과광범위한데이터선별이자주필요했는데, 이는확장성과배포의어려움을부각시켰습니다. 이결과들은 VLA 가의미론적의사결정과유연한작업명세에적합하지만, VLA 추론과고전적또는학습된저수준제어기를통합한하이브리드아키텍처가실용적이고현실적인작동에여전히필수적임을시사합니다.

도전과한계를해결하기위해우리는실시간추론, 멀티모달동작표현및안전성, 편향과일반화, 시스템통합및계산제약, 윤리적배포라는다섯가지핵심영역에집중했습니다. 우리는모델압축, 교차모달그라운딩, 도메인적응, 에이전트학습프레임워크등최신문헌에서잠재적해결책을제안했습니다. 마지막으로, 우리의논의와향후로드맵에서는 VLM, VLA 아키텍처, 에이전트 AI 시스템의융합이로봇공학을인공일반지능(AGI) 으로이끌고있음을명확히했습니다. 이검토는 VLA 발전에대한통합된이해를제공하고, 해결되지않은도전과제를식별하며, 미래에지능적이고체체적이며인간친화적에이전트를개발하기위한체계적인경로를제시합니다.

자금지원선언

이연구는미국국립과학재단(NSF) 과미국농무부(USDA), 국립식품농업연구소(NIFA) 의 ‘인공지능(AI) 농업연구소’ 프로그램을통해 ‘로봇소프트매니플레이터를이용한로봇꽃숙음’ 프로젝트의 AWD003473 및 AWD004595 및 USDA-NIFA 등록번호 1029004 지원을받았습니다. 추가지원은 ‘ExPanding UCF AI 연구를새로운농업공학응용으로 확장(PARTNER)’ 프로젝트하에 USDA/NIFA 보조금번호 2024-67022-41788, 등록번호 1031712 를통해제공되었습니다.

선언문

저자들은이해상충이없음을선언합니다.

AI 글쓰기지원에관한성명서

ChatGPT 와 Perplexity 는문법정확성을높이고문장구조를다듬기위해활용되었으며; 모든 AI 생성수정본은관련성을위해철저히검토되고편집되었습니다.

References

- [1] Achiam, J., Adler, S., Agarwal, S., Ahmad, L., Akkaya, I., Aleman, F.L., Almeida, D., Altenschmidt, J., Altman, S., Anadkat, S., et al., 2023. Gpt-4 technical report. arXiv preprint arXiv:2303.08774 .

- [2] Agarwal, L., Verma, B., 2024. From methods to datasets: A survey on image-caption generators. *Multimedia Tools and Applications* 83, 28077–28123.
- [3] Alayrac, J.B., Donahue, J., Luc, P., Miech, A., Barr, I., Hasson, Y., Lenc, K., Mensch, A., Millican, K., Reynolds, M., et al., 2022. Flamingo: a visual language model for few-shot learning. *Advances in neural information processing systems* 35, 23716–23736.
- [4] Anderson, P., Shrivastava, A., Truong, J., Majumdar, A., Parikh, D., Batra, D., Lee, S., 2021. Sim-to-real transfer for vision-and-language navigation, in: *Conference on Robot Learning*, PMLR. pp. 671–681.
- [5] Arai, H., Miwa, K., Sasaki, K., Watanabe, K., Yamaguchi, Y., Aoki, S., Yamamoto, I., 2025. Covla: Comprehensive vision-language-action dataset for autonomous driving, in: *2025 IEEE/CVF Winter Conference on Applications of Computer Vision (WACV)*, IEEE. pp. 1933–1943.
- [6] Asif, S., Bueno, M., Ferreira, P., Anandan, P., Zhang, Z., Yao, Y., Ragunathan, G., Tinkler, L., Sotoodeh-Bahraini, M., Lohse, N., et al., 2025. Rapid and automated configuration of robot manufacturing cells. *Robotics and Computer-Integrated Manufacturing* 92, 102862.
- [7] Assres, G., Bhandari, G., Shalaginov, A., Gronli, T.M., Ghinea, G., 2025. State-of-the-art and challenges of engineering ml-enabled software systems in the deep learning era. *ACM Computing Surveys*.
- [8] Asuzu, K., Singh, H., Idrissi, M., 2025. Human–robot interaction through joint robot planning with large language models. *Intelligent Service Robotics*, 1–17.
- [9] Ayaz, M., Khan, M., Saqib, M., Khelifi, A., Sajjad, M., Elsaddik, A., 2024. Medvlm: Medical vision-language model for consumer devices. *IEEE Consumer Electronics Magazine*.
- [10] Bai, S., Chen, K., Liu, X., Wang, J., Ge, W., Song, S., Dang, K., Wang, P., Wang, S., Tang, J., Zhong, H., Zhu, Y., Yang, M., Li, Z., Wan, J., Wang, P., Ding, W., Fu, Z., Xu, Y., Ye, J., Zhang, X., Xie, T., Cheng, Z., Zhang, H., Yang, Z., Xu, H., Lin, J., 2025. Qwen2.5-vl technical report. *arXiv preprint arXiv:2502.13923*.
- [11] Bartoccioni, F., Ramzi, E., Besnier, V., Venkataramanan, S., Vu, T.H., Xu, Y., Chambon, L., Gidaris, S., Odabas, S., Hurych, D., et al., 2025. Vavim and vavam: Autonomous driving through video generative modeling. *arXiv preprint arXiv:2502.15672*.
- [12] Bathula, N.V., Paleti, I., Pagidi, S., Akkumahanthi, S.S., Guduru, N.T., 2024. Policy learning-based image captioning with vision transformer, in: *2024 IEEE International Students' Conference on Electrical, Electronics and Computer Science (SCEECS)*, IEEE. pp. 1–6.
- [13] Belkhale, S., Ding, T., Xiao, T., Sermanet, P., Vuong, Q., Tompson, J., Chebotar, Y., Dwibedi, D., Sadigh, D., 2024. Rt-h: Action hierarchies using language. *arXiv preprint arXiv:2403.01823*.
- [14] Bjorck, J., Castañeda, F., Cherniadev, N., Da, X., Ding, R., Fan, L., Fang, Y., Fox, D., Hu, F., Huang, S., et al., 2025. Gr00t n1: An open foundation model for generalist humanoid robots. *arXiv preprint arXiv:2503.14734*.
- [15] Black, K., Brown, N., Driess, D., Esmail, A., Equi, M., Finn, C., Fusai, N., Groom, L., Hausman, K., Ichter, B., et al., 2024. Pi-0: A vision-language-action flow model for general robot control. *arXiv preprint arXiv:2410.24164*.
- [16] Bolya, D., Huang, P.Y., Sun, P., Cho, J.H., Madotto, A., Wei, C., Ma, T., Zhi, J., Rajasegaran, J., Rasheed, H., et al., 2025. Perception encoder: The best visual embeddings are not at the output of the network. *arXiv preprint arXiv:2504.13181*.
- [17] Bordes, F., Pang, R.Y., Ajay, A., Li, A.C., Bardes, A., Petryk, S., Mañas, O., Lin, Z., Mahmoud, A., Jayaraman, B., et al., 2024. An introduction to vision-language modeling. *arXiv preprint arXiv:2405.17247*.
- [18] Brohan, A., Brown, N., Carbajal, J., Chebotar, Y., Chen, X., Choromanski, K., Ding, T., Driess, D., Dubey, A., Finn, C., et al., 2023. Rt-2: Vision-language-action models transfer web knowledge to robotic control. *arXiv preprint arXiv:2307.15818*.
- [19] Brohan, A., Brown, N., Carbajal, J., Chebotar, Y., Dabis, J., Finn, C., Gopalakrishnan, K., Hausman, K., Herzog, A., Hsu, J., et al., 2022. Rt-1: Robotics transformer for real-world control at scale. *arXiv preprint arXiv:2212.06817*.
- [20] Budzianowski, P., Maa, W., Freed, M., Mo, J., Xie, A., Tipnis, V., Bolte, B., 2024. Edgevla: Efficient vision-language-action models. *environments* 20, 3.
- [21] Cangelosi, A., Metta, G., Sagerer, G., Nolfi, S., Nehaniv, C., Fischer, K., Tani, J., Belpaeme, T., Sandini, G., Nori, F., et al., 2010. Integration of action and language knowledge: A roadmap for developmental robotics. *IEEE Transactions on Autonomous Mental Development* 2, 167–195.
- [22] Cao, J., Gan, Z., Cheng, Y., Yu, L., Chen, Y.C., Liu, J., 2020. Behind the scene: Revealing the secrets of pre-trained vision-and-language models, in: *Computer Vision—ECCV 2020: 16th European Conference, Glasgow, UK, August 23–28, 2020, Proceedings, Part VI* 16, Springer. pp. 565–580.
- [23] Cao, L., 2024. Ai robots and humanoid ai: Review, perspectives and directions. *arXiv preprint arXiv:2405.15775*.

- [24] Cao, Y., Ju, Y., Xu, D., 2024a. 3dgs-det: Empower 3d gaussian splatting with boundary guidance and box-focused sampling for 3d object detection. *arXiv preprint arXiv:2410.01647*.
- [25] Cao, Y., Zeng, Y., Xu, H., Xu, D., 2023. Coda: Collaborative novel box discovery and cross-modal alignment for open-vocabulary 3d object detection, in: *NeurIPS*.
- [26] Cao, Y., Zeng, Y., Xu, H., Xu, D., 2024b. Collaborative novel object discovery and box-guided cross-modal alignment for open-vocabulary 3d object detection. *arXiv preprint arXiv:2406.00830*.
- [27] Chang, C., Shi, Y., Cao, D., Yang, W., Hwang, J., Wang, H., Pang, J., Wang, W., Liu, Y., Peng, W.C., et al., 2025. A survey of reasoning and agentic systems in time series with large language models. *arXiv preprint arXiv:2509.11575*.
- [28] Chang, Y., Wang, X., Wang, J., Wu, Y., Yang, L., Zhu, K., Chen, H., Yi, X., Wang, C., Wang, Y., et al., 2024. A survey on evaluation of large language models. *ACM transactions on intelligent systems and technology* 15, 1–45.
- [29] Chatzopoulos, D., Bermejo, C., Huang, Z., Hui, P., 2017. Mobile augmented reality survey: From where we are to where we go. *Ieee Access* 5, 6917–6950.
- [30] Chen, B., Xu, Z., Kirmani, S., Ichter, B., Sadigh, D., Guibas, L., Xia, F., 2024a. Spatialvlm: Endowing vision-language models with spatial reasoning capabilities, in: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 14455–14465.
- [31] Chen, H., Hou, L., Wu, S., Zhang, G., Zou, Y., Moon, S., Bhuiyan, M., 2024b. Augmented reality, deep learning and vision-language query system for construction worker safety. *Automation in Construction* 157, 105158.
- [32] Chen, H., Li, S., Fan, J., Duan, A., Yang, C., Navarro-Alarcon, D., Zheng, P., 2025a. Human-in-the-loop robot learning for smart manufacturing: A human-centric perspective. *IEEE Transactions on Automation Science and Engineering*.
- [33] Chen, H., Liu, B., Wang, S., Wang, X., Han, W., Zhu, Y., Wang, X., Bi, Y., 2025b. Language modulates vision: Evidence from neural networks and human brain-lesion models. *arXiv preprint arXiv:2501.13628*.
- [34] Chen, P., Bu, P., Wang, Y., Wang, X., Wang, Z., Guo, J., Zhao, Y., Zhu, Q., Song, J., Yang, S., et al., 2025c. Combatvla: An efficient vision-language-action model for combat tasks in 3d action role-playing games. *arXiv preprint arXiv:2503.09527*.
- [35] Chen, X., Xu, W., Kan, S., Zhang, L., Jin, Y., Cen, Y., Li, Y., 2025d. Vision-semantics-label: A new two-step paradigm for action recognition with large language model. *IEEE Transactions on Circuits and Systems for Video Technology*.
- [36] Chen, Y., Tian, S., Liu, S., Zhou, Y., Li, H., Zhao, D., 2025e. Conrft: A reinforced fine-tuning method for vla models via consistency policy. *arXiv preprint arXiv:2502.05450*.
- [37] Chen, Z., Wu, J., Wang, W., Su, W., Chen, G., Xing, S., Zhong, M., Zhang, Q., Zhu, X., Lu, L., et al., 2024c. Internvl: Scaling up vision foundation models and aligning for generic visual-linguistic tasks, in: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 24185–24198.
- [38] Cheng, A.C., Ji, Y., Yang, Z., Gongye, Z., Zou, X., Kautz, J., Bıyık, E., Yin, H., Liu, S., Wang, X., 2024a. Navila: Legged robot vision-language-action model for navigation. *arXiv preprint arXiv:2412.04453*.
- [39] Cheng, H., Xiao, E., Yu, C., Yao, Z., Cao, J., Zhang, Q., Wang, J., Sun, M., Xu, K., Gu, J., et al., 2024b. Manipulation facing threats: Evaluating physical vulnerabilities in end-to-end vision language action models. *arXiv preprint arXiv:2409.13174*.
- [40] Chi, C., Xu, Z., Feng, S., Cousineau, E., Du, Y., Burchfiel, B., Tedrake, R., Song, S., 2023. Diffusion policy: Visuomotor policy learning via action diffusion. *The International Journal of Robotics Research*, 02783649241273668.
- [41] Chi, H., Gao, H.a., Liu, Z., Liu, J., Liu, C., Li, J., Yang, K., Yu, Y., Wang, Z., Li, W., et al., 2025. Impromptu vla: Open weights and open data for driving vision-language-action models. *arXiv preprint arXiv:2505.23757*.
- [42] Chiang, H.T.L., Xu, Z., Fu, Z., Jacob, M.G., Zhang, T., Lee, T.W.E., Yu, W., Schenck, C., Rendleman, D., Shah, D., et al., 2024. Mobility vla: Multimodal instruction navigation with long-context vlms and topological graphs. *arXiv preprint arXiv:2407.07775*.
- [43] Chow, S.S., Alvi, R., Rahman, S.S., Rahman, M.A., Raiaa, M.A.K., Islam, M.R., Hussain, M., Azam, S., 2026. From language to action: a review of large language models as autonomous agents and tool users. *Artificial Intelligence Review*.
- [44] Dang, R., Yuan, Y., Zhang, W., Xin, Y., Zhang, B., Li, L., Wang, L., Zeng, Q., Li, X., Bing, L., 2025. Ecbench: Can multi-modal foundation models understand the egocentric world? a holistic embodied cognition benchmark. *arXiv preprint arXiv:2501.05031*.
- [45] Dasari, S., Ebert, F., Tian, S., Nair, S., Bucher, B., Schmeckpeper, K., Singh, S., Levine, S., Finn, C., 2019. Robonet: Large-scale multi-robot learning. *arXiv preprint arXiv:1910.11215*.

- [46] Deng, S., Yan, M., Wei, S., Ma, H., Yang, Y., Chen, J., Zhang, Z., Yang, T., Zhang, X., Cui, H., Zhang, Z., Wang, H., 2025a. Graspvla: a grasping foundation model pre-trained on billion-scale synthetic action data. URL: <https://arxiv.org/abs/2505.03233>, arXiv:2505.03233.
- [47] Deng, S., Yan, M., Zheng, Y., Su, J., Zhang, W., Zhao, X., Cui, H., Zhang, Z., Wang, H., 2025b. Stereovla: Enhancing vision-language-action models with stereo vision. arXiv preprint arXiv:2512.21970.
- [48] Dey, S., Zaech, J.N., Nikolov, N., Van Gool, L., Paudel, D.P., 2024. Revla: Reverting visual domain limitation of robotic foundation models. arXiv preprint arXiv:2409.15250.
- [49] Din, M.U., Akram, W., Saoud, L.S., Rosell, J., Husain, I., 2025. Vision language action models in robotic manipulation: A systematic review. arXiv preprint arXiv:2507.10672.
- [50] Ding, D., Yao, T., Luo, R., Sun, X., 2025a. Visual question answering in robotic surgery: A comprehensive review. IEEE Access.
- [51] Ding, J., Zhang, Y., Shang, Y., Zhang, Y., Zong, Z., Feng, J., Yuan, Y., Su, H., Li, N., Sukiennik, N., et al., 2024a. Understanding world or predicting future? a comprehensive survey of world models. arXiv preprint arXiv:2411.14499.
- [52] Ding, M., Chen, Z., Du, T., Luo, P., Tenenbaum, J., Gan, C., 2021. Dynamic visual reasoning by learning differentiable physics models from video and language. Advances in Neural Information Processing Systems 34, 887–899.
- [53] Ding, P., Ma, J., Tong, X., Zou, B., Luo, X., Fan, Y., Wang, T., Lu, H., Mo, P., Liu, J., et al., 2025b. Humanoid-vla: Towards universal humanoid control with visual integration. arXiv preprint arXiv:2502.14795.
- [54] Ding, P., Zhao, H., Zhang, W., Song, W., Zhang, M., Huang, S., Yang, N., Wang, D., 2024b. Quar-vla: Vision-language-action model for quadruped robots, in: European Conference on Computer Vision, Springer. pp. 352–367.
- [55] Donahue, J., Anne Hendricks, L., Guadarrama, S., Rohrbach, M., Venugopalan, S., Saenko, K., Darrell, T., 2015. Long-term recurrent convolutional networks for visual recognition and description, in: Proceedings of the IEEE conference on computer vision and pattern recognition, pp. 2625–2634.
- [56] Dong, H., Liu, M., Zhou, K., Chatzi, E., Kannala, J., Stachniss, C., Fink, O., 2025. Advances in multimodal adaptation and generalization: From traditional approaches to foundation models. arXiv preprint arXiv:2501.18592.
- [57] Doveh, S., Arbelle, A., Harary, S., Schwartz, E., Herzig, R., Giryas, R., Feris, R., Panda, R., Ullman, S., Karlinsky, L., 2023. Teaching structured vision & language concepts to vision & language models, in: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, pp. 2657–2668.
- [58] Driess, D., Xia, F., Sajjadi, M.S., Lynch, C., Chowdhery, A., Wahid, A., Tompson, J., Vuong, Q., Yu, T., Huang, W., et al., 2023. Palm-e: An embodied multimodal language model. Openreview.
- [59] Duan, J., Pumacay, W., Kumar, N., Wang, Y.R., Tian, S., Yuan, W., Krishna, R., Fox, D., Mandlekar, A., Guo, Y., 2024. Aha: A vision-language-model for detecting and reasoning over failures in robotic manipulation. arXiv preprint arXiv:2410.00371.
- [60] Duarte, N.F., Raković, M., Tasevski, J., Coco, M.I., Billard, A., Santos-Victor, J., 2018. Action anticipation: Reading the intentions of humans and robots. IEEE Robotics and Automation Letters 3, 4132–4139.
- [61] Ebert, F., Yang, Y., Schmeckpeper, K., Bucher, B., Georgakis, G., Daniilidis, K., Finn, C., Levine, S., 2021. Bridge data: Boosting generalization of robotic skills with cross-domain datasets. arXiv preprint arXiv:2109.13396.
- [62] Fan, C., Jia, X., Sun, Y., Wang, Y., Wei, J., Gong, Z., Zhao, X., Tomizuka, M., Yang, X., Yan, J., Ding, M., 2025a. Interleave-vla: Enhancing robot manipulation with interleaved image-text instructions. URL: <https://arxiv.org/abs/2505.02152>, arXiv:2505.02152.
- [63] Fan, L., Chen, K., Xu, Z., Yuan, M., Huang, P., Huang, W., 2024. Language reasoning in vision-language-action model for robotic grasping, in: 2024 China Automation Congress (CAC), IEEE. pp. 6656–6661.
- [64] Fan, Y., Ding, P., Bai, S., Tong, X., Zhu, Y., Lu, H., Dai, F., Zhao, W., Liu, Y., Huang, S., et al., 2025b. Long-vla: Unleashing long-horizon capability of vision language action model for robot manipulation. arXiv preprint arXiv:2508.19958.
- [65] Fang, H., Liu, Y., Du, Y., Du, L., Yang, H., 2025a. Sqap-vla: A synergistic quantization-aware pruning framework for high-performance vision-language-action models. arXiv preprint arXiv:2509.09090.
- [66] Fang, Y., Yang, Y., Zhu, X., Zheng, K., Bertasius, G., Szafir, D., Ding, M., 2025b. Rebot: Scaling robot learning with real-to-sim-to-real robotic video synthesis. arXiv preprint arXiv:2503.14526.
- [67] Firoozi, R., Tucker, J., Tian, S., Majumdar, A., Sun, J., Liu, W., Zhu, Y., Song, S., Kapoor, A., Hausman, K.,

- et al., 2023. Foundation models in robotics: Applications, challenges, and the future. *The International Journal of Robotics Research*, 02783649241281508.
- [68] Foster, D.J., Block, A., Misra, D., 2024. Is behavior cloning all you need? understanding horizon in imitation learning. *arXiv preprint arXiv:2407.15007*.
- [69] Fu, H., Zhang, D., Zhao, Z., Cui, J., Liang, D., Zhang, C., Zhang, D., Xie, H., Wang, B., Bai, X., 2025. Orion: A holistic end-to-end autonomous driving framework by vision-language instructed action generation. *arXiv preprint arXiv:2503.19755*.
- [70] Gao, B., Liu, Y., Li, Y., Li, H., Li, M., He, W., 2025a. A vision-language model for predicting potential distribution land of soybean double cropping. *Frontiers in Environmental Science* 12, 1515752.
- [71] Gao, C., Liu, Z., Chi, Z., Huang, J., Fei, X., Hou, Y., Zhang, Y., Lin, Y., Fang, Z., Jiang, Z., et al., 2025b. Vlasos: Structuring and dissecting planning representations and paradigms in vision-language-action models. *arXiv preprint arXiv:2506.17561*.
- [72] Gao, J., Belkhale, S., Dasari, S., Balakrishna, A., Shah, D., Sadigh, D., 2025c. A taxonomy for evaluating generalist robot policies. *arXiv preprint arXiv:2503.01238*.
- [73] Gao, S.H., Cheng, M.M., Zhao, K., Zhang, X.Y., Yang, M.H., Torr, P., 2019. Res2net: A new multi-scale backbone architecture. *IEEE TPAMI*.
- [74] Gbagbe, K.F., Cabrera, M.A., Alabbas, A., Alyunes, O., Lykov, A., Tsetserukou, D., 2024. Bi-vla: Vision-language-action model-based system for bimanual robotic dexterous manipulations, in: *2024 IEEE International Conference on Systems, Man, and Cybernetics (SMC)*, IEEE. pp. 2864–2869.
- [75] Geens, R., 2024. Bringing generative ai to edge devices through interoperable compute cores, in: *Flanders AI Research Day*, Location: Ghent.
- [76] Ghosh, A., Acharya, A., Saha, S., Jain, V., Chadha, A., 2024. Exploring the frontier of vision-language models: A survey of current methodologies and future directions. *arXiv preprint arXiv:2404.07214*.
- [77] Gu, J., Wang, Z., Kuen, J., Ma, L., Shahroudy, A., Shuai, B., Liu, T., Wang, X., Wang, G., Cai, J., et al., 2018. Recent advances in convolutional neural networks. *Pattern recognition* 77, 354–377.
- [78] Gu, Q., Ju, Y., Sun, S., Gilitschenski, I., Nishimura, H., Itkina, M., Shkurti, F., 2025a. Safe: Multitask failure detection for vision-language-action models. *arXiv preprint arXiv:2506.09937*.
- [79] Gu, Q., Su, J., Yuan, L., 2021. Visual affordance detection using an efficient attention convolutional neural network. *Neurocomputing* 440, 36–44. doi:<https://doi.org/10.1016/j.neucom.2021.01.018>.
- [80] Gu, X., Wen, C., Ye, W., Song, J., Gao, Y., 2023. Seer: Language instructed video prediction with latent diffusion models. *arXiv preprint arXiv:2303.14897*.
- [81] Gu, Z., Li, J., Shen, W., Yu, W., Xie, Z., McCrory, S., Cheng, X., Shamsah, A., Griffin, R., Liu, C.K., et al., 2025b. Humanoid locomotion and manipulation: Current progress and challenges in control, planning, and learning. *arXiv preprint arXiv:2501.02116*.
- [82] Guan, W., Hu, Q., Li, A., Cheng, J., 2025. Efficient vision-language-action models for embodied manipulation: A systematic survey. *arXiv preprint arXiv:2510.17111*.
- [83] Guo, Y., Zhang, J., Chen, X., Ji, X., Wang, Y.J., Hu, Y., Chen, J., 2025. Improving vision-language-action model with online reinforcement learning. *arXiv preprint arXiv:2501.16664*.
- [84] Guruprasad, P., Sikka, H., Song, J., Wang, Y., Liang, P.P., 2024. Benchmarking vision, language, & action models on robotic learning tasks. *arXiv preprint arXiv:2411.05821*.
- [85] Haldar, S., Peng, Z., Pinto, L., 2024. Baku: An efficient transformer for multi-task policy learning. *arXiv preprint arXiv:2406.07539*.
- [86] Han, S., Wang, M., Zhang, J., Li, D., Duan, J., 2024. A review of large language models: Fundamental architectures, key technological evolutions, interdisciplinary technologies integration, optimization and compression techniques, applications, and challenges. *Electronics* 13, 5040.
- [87] Hanson, A., Riseman, E., 2014. The visions image-understanding system, in: *Advances in Computer Vision*. Psychology Press, pp. 1–114.
- [88] Hao, P., Zhang, C., Li, D., Cao, X., Hao, X., Cui, S., Wang, S., 2025. Tla: Tactile-language-action model for contact-rich manipulation. *arXiv preprint arXiv:2503.08548*.
- [89] He, K., Zhang, X., Ren, S., Sun, J., 2016. Deep residual learning for image recognition, pp. 770–778.
- [90] Holldack, F., Banh, L., Strobel, G., 2026. Agentic information systems. *Electronic Markets* 36, 5.
- [91] Hong, Y., 2025. Building 3D Foundation Models for the Embodied Minds. Ph.D. thesis. University of California, Los Angeles.

- [92] Hou, Z., Zhang, T., Xiong, Y., Duan, H., Pu, H., Tong, R., Zhao, C., Zhu, X., Qiao, Y., Dai, J., et al., 2025. Dita: Scaling diffusion transformer for generalist vision-language-action policy. *arXiv preprint arXiv:2503.19757*.
- [93] Hu, E.J., Shen, Y., Wallis, P., Allen-Zhu, Z., Li, Y., Wang, S., Wang, L., Chen, W., et al., 2022. Lora: Low-rank adaptation of large language models. *ICLR* 1, 3.
- [94] Hu, Y., Tang, J., Gong, X., Zhou, Z., Zhang, S., Elvitigala, D.S., Mueller, F., Hu, W., Quigley, A.J., 2025. Vision-based multimodal interfaces: A survey and taxonomy for enhanced context-aware system design. *arXiv preprint arXiv:2501.13443*.
- [95] Huang, G., Liu, Z., Van Der Maaten, L., Weinberger, K.Q., 2017. Densely connected convolutional networks, in: *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 4700–4708.
- [96] Huang, H., Liu, F., Fu, L., Wu, T., Mukadam, M., Malik, J., Goldberg, K., Abbeel, P., 2024. Early fusion helps vision language action models generalize better, in: *1st Workshop on X-Embodiment Robot Learning*.
- [97] Huang, H., Liu, F., Fu, L., Wu, T., Mukadam, M., Malik, J., Goldberg, K., Abbeel, P., 2025a. Otter: A vision-language-action model with text-aware visual feature extraction. *arXiv preprint arXiv:2503.03734*.
- [98] Huang, S., Dong, L., Wang, W., Hao, Y., Singhal, S., Ma, S., Lv, T., Cui, L., Mohammed, O.K., Patra, B., et al., 2023a. Language is not all you need: Aligning perception with language models. *Advances in Neural Information Processing Systems* 36, 72096–72109.
- [99] Huang, W., Gu, Q., Ye, N., 2025b. Decision spike-former: Spike-driven transformer for decision making. *arXiv preprint arXiv:2504.03800*.
- [100] Huang, W., Wang, C., Zhang, R., Li, Y., Wu, J., Fei-Fei, L., 2023b. Voxposer: Composable 3d value maps for robotic manipulation with language models. *arXiv preprint arXiv:2307.05973*.
- [101] Huang, Y., Hua, H., Zhou, Y., Jing, P., Nagireddy, M., Padhi, I., Dolcetti, G., Xu, Z., Chaudhury, S., Rawat, A., et al., 2025c. Building a foundational guardrail for general agentic systems via synthetic data. *arXiv preprint arXiv:2510.09781*.
- [102] Huang, Z., Chen, F., Pu, Y., Lin, C., Su, H., Gan, C., 2023c. Diffvl: Scaling up soft body manipulation using vision-language driven differentiable physics. *Advances in Neural Information Processing Systems* 36, 29875–29900.
- [103] Hung, C.Y., Sun, Q., Hong, P., Zadeh, A., Li, C., Tan, U., Majumder, N., Poria, S., et al., 2025. Nora: A small open-sourced generalist vision language action model for embodied tasks. *arXiv preprint arXiv:2504.19854*.
- [104] Ikeda, B., Gramopadhye, M., Nekervis, L., Szafir, D., 2025. Marcer: Multimodal augmented reality for composing and executing robot tasks, in: *2025 20th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*, IEEE. pp. 529–539.
- [105] Imran, A., Gopalakrishnan, K., 2025. Foundation models in robotics, in: *AI for Robotics: Toward Embodied and General Intelligence in the Physical World*. Springer, pp. 139–210.
- [106] Intelligence, P., Black, K., Brown, N., Darpinian, J., Dhabalia, K., Driess, D., Esmail, A., Equi, M., Finn, C., Fusai, N., et al., 2025. pi0.5: a vision-language-action model with open-world generalization. *arXiv preprint arXiv:2504.16054*.
- [107] Jeong, H., Lee, H., Kim, C., Shin, S., 2024. A survey of robot intelligence with large language models. *Applied Sciences* 14, 8868.
- [108] Jha, K., Doshi, A., Patel, P., Shah, M., 2019. A comprehensive review on automation in agriculture using artificial intelligence. *Artificial Intelligence in Agriculture* 2, 1–12.
- [109] Jiang, J., Xiao, W., Lin, Z., Zhang, H., Ren, T., Gao, Y., Lin, Z., Cai, Z., Yang, L., Liu, Z., 2024. Solami: Social vision-language-action modeling for immersive interaction with 3d autonomous characters. *arXiv preprint arXiv:2412.00174*.
- [110] Jiang, J., Xiao, W., Lin, Z., Zhang, H., Ren, T., Gao, Y., Lin, Z., Cai, Z., Yang, L., Liu, Z., 2025a. Solami: Social vision-language-action modeling for immersive interaction with 3d autonomous characters, in: *Proceedings of the Computer Vision and Pattern Recognition Conference*, pp. 26887–26898.
- [111] Jiang, S., Huang, Z., Qian, K., Luo, Z., Zhu, T., Zhong, Y., Tang, Y., Kong, M., Wang, Y., Jiao, S., et al., 2025b. A survey on vision-language-action models for autonomous driving. *arXiv preprint arXiv:2506.24044*.
- [112] Jiang, Y., Gupta, A., Zhang, Z., Wang, G., Dou, Y., Chen, Y., Fei-Fei, L., Anandkumar, A., Zhu, Y., Fan, L., 2022. Vima: General robot manipulation with multimodal prompts. *arXiv preprint arXiv:2210.03094* 2, 6.
- [113] Jiang, Y., Zhang, R., Wong, J., Wang, C., Ze, Y., Yin, H., Gokmen, C., Song, S., Wu, J., Fei-Fei, L., 2025c. Behavior robot suite: Streamlining real-world whole-body manipulation for everyday household activities. *arXiv preprint arXiv:2503.05652*.
- [114] Jones, J., Mees, O., Sferrazza, C., Stachowicz, K., Abbeel, P., Levine, S., 2025. Beyond sight: Finetuning generalist robot policies with heterogeneous sensors via language grounding. *arXiv preprint arXiv:2501.04693*.

- [115] Karamcheti, S., Zhai, A.J., Losey, D.P., Sadigh, D., 2021. Learning visually guided latent actions for assistive teleoperation, in: *Learning for dynamics and control*, PMLR. pp. 1230–1241.
- [116] Karli, U.B., Kurumisawa, T., Fitzgerald, T., 2025. Ask before you act: Token-level uncertainty for intervention in vision-language-action models, in: *Second Workshop on Out-of-Distribution Generalization in Robotics at RSS* 2025.
- [117] Katiyar, N., 2023. A Model-Driven Framework for Domain-Specific Adaptation of Time Series Forecasting Pipeline. McGill University (Canada).
- [118] Kawaharazuka, K., Oh, J., Yamada, J., Posner, I., Zhu, Y., 2025. Vision-language-action models for robotics: A review towards real-world applications. *IEEE Access* .
- [119] Kelly, C., Hu, L., Yang, B., Tian, Y., Yang, D., Yang, C., Huang, Z., Li, Z., Hu, J., Zou, Y., 2024. Visiongpt: Vision-language understanding agent using generalized multimodal framework. *arXiv preprint arXiv:2403.09027* .
- [120] Khan, M.H., Asfaw, S., Iarchuk, D., Cabrera, M.A., Moreno, L., Tokmurziyev, I., Tsetserukou, D., 2025. Shake-vla: Vision-language-action model-based system for bimanual robotic manipulations and liquid mixing. *arXiv preprint arXiv:2501.06919* .
- [121] Kim, M.J., Finn, C., Liang, P., 2025. Fine-tuning vision-language-action models: Optimizing speed and success. *arXiv preprint arXiv:2502.19645* .
- [122] Kim, M.J., Pertsch, K., Karamcheti, S., Xiao, T., Balakrishna, A., Nair, S., Rafailov, R., Foster, E., Lam, G., Sanketi, P., et al., 2024. Openvla: An open-source vision-language-action model. *arXiv preprint arXiv:2406.09246* .
- [123] Koo, J., Cho, T., Kang, H., Pyo, E., Oh, T.G., Kim, T., Choi, A.J., 2025. Retovla: Reusing register tokens for spatial reasoning in vision-language-action models. *arXiv preprint arXiv:2509.21243* .
- [124] Lee, N., Bang, Y., Lovenia, H., Cahyawijaya, S., Dai, W., Fung, P., 2023. Survey of social bias in vision-language models. *arXiv preprint arXiv:2309.14381* .
- [125] Li, C., Wen, J., Peng, Y., Peng, Y., Feng, F., Zhu, Y., 2025a. Pointvla: Injecting the 3d world into vision-language-action models. *arXiv preprint arXiv:2503.07511* .
- [126] Li, D., Jin, Y., Sun, Y., Yu, H., Shi, J., Hao, X., Hao, P., Liu, H., Sun, F., Zhang, J., et al., 2024a. What foundation models can bring for robot learning in manipulation: A survey. *arXiv preprint arXiv:2404.18201* .
- [127] Li, J., Skinner, G., Yang, G., Quaranto, B.R., Schwaitzberg, S.D., Kim, P.C., Xiong, J., 2024b. Llava-surg: towards multimodal surgical assistant via structured surgical video learning. *arXiv preprint arXiv:2408.07981* .
- [128] Li, J., Wei, P., Han, W., Fan, L., 2023. Intentqa: Context-aware video intent reasoning, in: *Proceedings of the IEEE/CVF international conference on computer vision*, pp. 11963–11974.
- [129] Li, J., Zhu, Y., Tang, Z., Wen, J., Zhu, M., Liu, X., Li, C., Cheng, R., Peng, Y., Feng, F., 2024c. Improving vision-language-action models via chain-of-affordance. *arXiv preprint arXiv:2412.20451* .
- [130] Li, M., Wang, Z., He, K., Ma, X., Liang, Y., 2025b. Jarvis-vla: Post-training large-scale vision language models to play visual games with keyboards and mouse. *arXiv preprint arXiv:2503.16365* .
- [131] Li, Q., Liang, Y., Wang, Z., Luo, L., Chen, X., Liao, M., Wei, F., Deng, Y., Xu, S., Zhang, Y., et al., 2024d. Cogact: A foundational vision-language-action model for synergizing cognition and action in robotic manipulation. *arXiv preprint arXiv:2411.19650* .
- [132] Li, S., Wang, J., Dai, R., Ma, W., Ng, W.Y., Hu, Y., Li, Z., 2024e. Robonurse-vla: Robotic scrub nurse system based on vision-language-action model. *arXiv preprint arXiv:2409.19590* .
- [133] Li, S., Wu, H., Shao, J., Ma, Y., Gan, Y., Luo, Y., Wang, Y., Nie, D., Wang, L., Wu, W., et al., 2025c. Seeing the forest and the trees: Query-aware tokenizer for long-video multimodal language models. *arXiv preprint arXiv:2511.11910* .
- [134] Li, Y., Deng, Y., Zhang, J., Jang, J., Memmel, M., Yu, R., Garrett, C.R., Ramos, F., Fox, D., Li, A., et al., 2025d. Hamster: Hierarchical action models for open-world robot manipulation. *arXiv preprint arXiv:2502.05485* .
- [135] Li, Y., Gong, Z., Li, H., Huang, X., Kang, H., Bai, G., Ma, X., 2025e. Robotic visual instruction. *arXiv preprint arXiv:2505.00693* .
- [136] Li, Y., Lai, Z., Bao, W., Tan, Z., Dao, A., Sui, K., Shen, J., Liu, D., Liu, H., Kong, Y., 2025f. Visual large language models for generalized and specialized applications. *arXiv preprint arXiv:2501.02765* .
- [137] Li, Z., Wu, X., Du, H., Nghiem, H., Shi, G., 2025g. Benchmark evaluations, applications, and challenges of large vision language models: A survey. *arXiv preprint arXiv:2501.02189* 1.
- [138] Liang, Z., Li, Y., Yang, T., Wu, C., Mao, S., Nian, T., Pei, L., Zhou, S., Yang, X., Pang, J., et al., 2025. Discrete diffusion vla: Bringing discrete diffusion to action decoding in vision-language-action policies. *arXiv preprint arXiv:2508.20072* .

- [139] Lin, K.Q., Li, L., Gao, D., Yang, Z., Wu, S., Bai, Z., Lei, W., Wang, L., Shou, M.Z., 2024. Showui: One vision-language-action model for gui visual agent. arXiv preprint arXiv:2411.17465 .
- [140] Lin, Y., Zhou, H., Chen, M., Min, H., 2019. Automatic sorting system for industrial robot with 3d visual perception and natural language interaction. *Measurement and Control* 52, 100–115.
- [141] Liu, H., Yao, R., Liu, W., Huang, Z., Shen, S., Ma, J., 2025a. Codrivevlm: Vlm-enhanced urban cooperative dispatching and motion planning for future autonomous mobility on demand systems. URL: <https://arxiv.org/abs/2501.06132>, arXiv:2501.06132.
- [142] Liu, J., Chen, H., An, P., Liu, Z., Zhang, R., Gu, C., Li, X., Guo, Z., Chen, S., Liu, M., et al., 2025b. Hybridvla: Collaborative diffusion and autoregression in a unified vision-language-action model. arXiv preprint arXiv:2503.10631 .
- [143] Liu, J., Liu, M., Wang, Z., An, P., Li, X., Zhou, K., Yang, S., Zhang, R., Guo, Y., Zhang, S., 2024a. Robomamba: Efficient vision-language-action model for robotic reasoning and manipulation. *Advances in Neural Information Processing Systems* 37, 40085–40110.
- [144] Liu, S., Wu, L., Li, B., Tan, H., Chen, H., Wang, Z., Xu, K., Su, H., Zhu, J., 2024b. Rdt-1b: a diffusion foundation model for bimanual manipulation. arXiv preprint arXiv:2410.07864 .
- [145] Liu, S., Yang, S., Fang, D., Jia, S., Tang, Y., Su, L., Peng, R., Yan, Y., Zou, X., Hu, X., 2026. Vision-language introspection: Mitigating overconfident hallucinations in mllms via interpretable bi-causal steering. arXiv preprint arXiv:2601.05159 .
- [146] Liu, Y., Cao, X., Chen, T., Jiang, Y., You, J., Wu, M., Wang, X., Feng, M., Jin, Y., Chen, J., 2025c. From screens to scenes: A survey of embodied ai in healthcare. arXiv preprint arXiv:2501.07468 .
- [147] Liu, Y., Cao, X., Chen, T., Jiang, Y., You, J., Wu, M., Wang, X., Feng, M., Jin, Y., Chen, J., 2025d. A survey of embodied ai in healthcare: Techniques, applications, and opportunities. arXiv preprint arXiv:2501.07468 .
- [148] Liu, Z., Liang, H., Huang, X., Xiong, W., Yu, Q., Sun, L., Chen, C., He, C., Cui, B., Zhang, W., 2024c. Synthvlm: High-efficiency and high-quality synthetic data for vision language models. arXiv preprint arXiv:2407.20756 .
- [149] Lu, H., Li, H., Shahani, P.S., Herbers, S., Scheutz, M., 2025. Probing a vision-language-action model for symbolic states and integration into a cognitive architecture. arXiv preprint arXiv:2502.04558 .
- [150] Lu, J., Clark, C., Lee, S., Zhang, Z., Khosla, S., Marten, R., Hoiem, D., Kembhavi, A., 2024. Unified-io 2: Scaling autoregressive multimodal models with vision language audio and action, in: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 26439–26455.
- [151] Lu, Y., Liao, Z., 2023. Towards happy housework: Scenario-based experience design for a household cleaning robotic system. *EAI Endorsed Transactions on Scalable Information Systems* 10.
- [152] Luo, H., Feng, Y., Zhang, W., Zheng, S., Wang, Y., Yuan, H., Liu, J., Xu, C., Jin, Q., Lu, Z., 2025. Being-h0: vision-language-action pretraining from large-scale human videos. arXiv preprint arXiv:2507.15597 .
- [153] Luo, J., Xu, C., Wu, J., Levine, S., 2024. Precise and dexterous robotic manipulation via human-in-the-loop reinforcement learning. arXiv preprint arXiv:2410.21845 .
- [154] Lyre, H., 2024. "understanding ai": Semantic grounding in large language models. URL: <https://arxiv.org/abs/2402.10992>, arXiv:2402.10992.
- [155] Lyu, J., Li, Z., Shi, X., Xu, C., Wang, Y., Wang, H., 2025. Dywa: Dynamics-adaptive world action model for generalizable non-prehensile manipulation. arXiv preprint arXiv:2503.16806 .
- [156] Ma, Y., Song, Z., Zhuang, Y., Hao, J., King, I., 2024. A survey on vision-language-action models for embodied ai. arXiv preprint arXiv:2405.14093 .
- [157] Misra, I., Girdhar, R., Joulin, A., 2021. An end-to-end transformer model for 3d object detection, in: *ICCV*.
- [158] Mohammed, M.Q., Chung, K.L., Chyi, C.S., 2020. Review of deep reinforcement learning-based object grasping: Techniques, open challenges, and recommendations. *Ieee Access* 8, 178450–178481.
- [159] Moroncelli, A., Soni, V., Shahid, A.A., Maccarini, M., Forgiione, M., Piga, D., Spahiu, B., Roveda, L., 2024. Integrating reinforcement learning with foundation models for autonomous robotics: Methods and perspectives. arXiv preprint arXiv:2410.16411 .
- [160] Mumuni, A., Mumuni, F., 2025. Large language models for artificial general intelligence (agi): A survey of foundational principles and approaches. arXiv preprint arXiv:2501.03151 .
- [161] Ni, F., Hao, J., Wu, S., Kou, L., Yuan, Y., Dong, Z., Liu, J., Li, M., Zhuang, Y., Zheng, Y., 2024. Peria: Perceive, reason, imagine, act via holistic language and vision planning for manipulation. *Advances in Neural Information Processing Systems* 37, 17541–17571.

- [162] Nie, Y., Li, L., Gan, Z., Wang, S., Zhu, C., Zeng, M., Liu, Z., Bansal, M., Wang, L., 2021. Mlp architectures for vision-and-language modeling: An empirical study. *arXiv preprint arXiv:2112.04453* .
- [163] Noorani, E., Serlin, Z., Price, B., Velasquez, A., 2025. From abstraction to reality: Darpa’s vision for robust sim-to-real autonomy. *arXiv preprint arXiv:2503.11007* .
- [164] Oquab, M., Darcet, T., Moutakanni, T., Vo, H., Szafraniec, M., Khalidov, V., Fernandez, P., Haziza, D., Massa, F., El-Nouby, A., et al., 2023. Dinov2: Learning robust visual features without supervision. *arXiv preprint arXiv:2304.07193* .
- [165] Pang, J., Zheng, P., Fan, J., Liu, T., 2025. Towards cognition-augmented human-centric assembly: A visual computation perspective. *Robotics and Computer-Integrated Manufacturing* 91, 102852.
- [166] Pantalone, D., Faini, G.S., Cialdai, F., Sereni, E., Bacci, S., Bani, D., Bernini, M., Pratesi, C., Stefano, P., Orzalesi, L., et al., 2021. Robot-assisted surgery in space: pros and cons. a review from the surgeon’s point of view. *npj Microgravity* 7, 56.
- [167] Park, S., Kim, H., Kim, S., Jeon, W., Yang, J., Jeon, B., Oh, Y., Choi, J., 2025. Saliency-aware quantized imitation learning for efficient robotic control, in: *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pp. 13140–13150.
- [168] Park, S.M., Kim, Y.G., 2023. Visual language integration: A survey and open challenges. *Computer Science Review* 48, 100548.
- [169] Pasas-Farmer, S., Jain, R., 2025. From discovery to delivery: Governance of ai in the pharmaceutical industry. *Green Analytical Chemistry* 13, 100268.
- [170] Patel, D., Eghbalzadeh, H., Kamra, N., Iuzzolino, M.L., Jain, U., Desai, R., 2023. Pretrained language models as visual planners for human assistance, in: *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pp. 15302–15314.
- [171] Pertsch, K., Stachowicz, K., Ichter, B., Driess, D., Nair, S., Vuong, Q., Mees, O., Finn, C., Levine, S., 2025. Fast: Efficient action tokenization for vision-language-action models. *arXiv preprint arXiv:2501.09747* .
- [172] Plaata, A., van Duijn, M., van Stein, N., Preuss, M., van der Putten, P., Batenburg, K.J., 2025. Agentic large language models, a survey. *arXiv preprint arXiv:2503.23037* .
- [173] Polubarov, A., Lyubaykin, N., Derevyagin, A., Zisman, I., Tarasov, D., Nikulin, A., Kurenkov, V., 2025. Vintix: Action model via in-context reinforcement learning. *arXiv preprint arXiv:2501.19400* .
- [174] Qi, C.R., Litany, O., He, K., Guibas, L.J., 2019. Deep hough voting for 3d object detection in point clouds, in: *ICCV*.
- [175] Qu, D., Song, H., Chen, Q., Yao, Y., Ye, X., Ding, Y., Wang, Z., Gu, J., Zhao, B., Wang, D., et al., 2025. Spatialvla: Exploring spatial representations for visual-language-action model. *arXiv preprint arXiv:2501.15830* .
- [176] Radford, A., Kim, J.W., Hallacy, C., Ramesh, A., Goh, G., Agarwal, S., Sastry, G., Askell, A., Mishkin, P., Clark, J., et al., 2021. Learning transferable visual models from natural language supervision, in: *ICML, PMLR*.
- [177] Radford, A., Narasimhan, K., Salimans, T., Sutskever, I., et al., 2018. Improving language understanding by generative pre-training .
- [178] Rawal, P.K., 2025. An Intelligent Versatile Pipeline for 6D Localization of Industrial Components in a Production Environment. Ph.D. thesis. Fraunhofer Verlag.
- [179] Ray, P.P., 2023. Chatgpt: A comprehensive review on background, applications, key challenges, bias, ethics, limitations and future scope. *Internet of Things and Cyber-Physical Systems* 3, 121–154.
- [180] Raza, S., Qureshi, R., Zahid, A., Fioresi, J., Sadak, F., Saeed, M., Sapkota, R., Jain, A., Zafar, A., Hassan, M.U., et al., 2025. Who is responsible? the data, models, users or regulations? responsible generative ai for a sustainable future. *arXiv preprint arXiv:2502.08650* .
- [181] Reed, S., Zolna, K., Parisotto, E., Colmenarejo, S.G., Novikov, A., Barth-Maron, G., Gimenez, M., Sulsky, Y., Kay, J., Springenberg, J.T., et al., 2022. A generalist agent. *arXiv preprint arXiv:2205.06175* .
- [182] Rodriguez-Guerra, D., Sorrosal, G., Cabanes, I., Calleja, C., 2021. Human-robot interaction review: Challenges and solutions for modern industrial environments. *Ieee Access* 9, 108557–108578.
- [183] Rodriguez-Juan, J., Ortiz-Perez, D., Garcia-Rodriguez, J., Tomás, D., Nalepa, G.J., 2025. Integrating advanced vision-language models for context recognition in risks assessment. *Neurocomputing* 618, 129131.
- [184] Roychoudhury, A., Khorshidi, S., Agrawal, S., Bennewitz, M., 2023. Perception for humanoid robots. *Current Robotics Reports* 4, 127–140.
- [185] Sahili, Z.A., Patras, I., Purver, M., 2025. Scaling for fairness? analyzing model size, data composition, and multilinguality in vision-language bias. *arXiv preprint arXiv:2501.13223* .
- [186] Sameni, S., Kafle, K., Tan, H., Jenni, S., 2024. Building vision-language models on solid foundations with masked distillation, in: *Proceedings of the IEEE/CVF*

- Conference on Computer Vision and Pattern Recognition, pp. 14216–14226.
- [187] Samson, M., Muraccioli, B., Kanehiro, F., 2025. Scalable, training-free visual language robotics: a modular multi-model framework for consumer-grade gpus, in: 2025 IEEE/SICE International Symposium on System Integration (SII), IEEE. pp. 193–198.
- [188] Sanyal, S., Roy, K., 2025. Asma: An adaptive safety margin algorithm for vision-language drone navigation via scene-aware control barrier functions. *IEEE Robotics and Automation Letters* .
- [189] Sapkota, R., Karkee, M., 2025. Object detection with multimodal large vision-language models: An in-depth review. Available at SSRN 5233953 .
- [190] Sapkota, R., Roumeliotis, K.I., Cheppally, R.H., Calero, M.F., Karkee, M., 2025. A review of 3d object detection with vision-language models. *arXiv preprint arXiv:2504.18738* .
- [191] Sautenkov, O., Yaqoot, Y., Lykov, A., Mustafa, M.A., Tadevosyan, G., Akhmetkazy, A., Cabrera, M.A., Martynov, M., Karaf, S., Tsetserukou, D., 2025. Uav-vla: Vision-language-action system for large scale aerial mission generation. *arXiv preprint arXiv:2501.05014* .
- [192] Schakkal, A., Zandonati, B., Yang, Z., Azizan, N., 2025. Hierarchical vision-language planning for multi-step humanoid manipulation. *arXiv preprint arXiv:2506.22827* .
- [193] Schmidgall, S., Cho, J., Zakka, C., Hiesinger, W., 2024. Gp-vls: A general-purpose vision language model for surgery. *arXiv preprint arXiv:2407.19305* .
- [194] Schuhmann, C., Beaumont, R., Vencu, R., Gordon, C., Wightman, R., Cherti, M., Coombes, T., Katta, A., Mullis, C., Wortsman, M., et al., 2022. Laion-5b: An open large-scale dataset for training next generation image-text models. *Advances in neural information processing systems* 35, 25278–25294.
- [195] Serpiva, V., Lykov, A., Myshlyayev, A., Khan, M.H., Abdulkarim, A.A., Sautenkov, O., Tsetserukou, D., 2025. Racevla: Vla-based racing drone navigation with human-like behaviour. *arXiv preprint arXiv:2503.02572* .
- [196] Shadab Siddiqui, M., Rabbi, M., Islam, M.J., Ahmed, R.U., 2025. Comparison of different controller architectures for autonomous driving and recommendations for robust and safe implementations. *Journal of Advanced Transportation* 2025, 9995539.
- [197] Shao, R., Li, W., Zhang, L., Zhang, R., Liu, Z., Chen, R., Nie, L., 2025. Large vlm-based vision-language-action models for robotic manipulation: A survey. *arXiv preprint arXiv:2508.13073* .
- [198] Sharshar, A., Khan, L.U., Ullah, W., Guizani, M., 2025. Vision-language models for edge networks: A comprehensive survey. *arXiv preprint arXiv:2502.07855* .
- [199] Shi, L.X., Ichter, B., Equi, M., Ke, L., Pertsch, K., Vuong, Q., Tanner, J., Walling, A., Wang, H., Fusai, N., et al., 2025. Hi robot: Open-ended instruction following with hierarchical vision-language-action models. *arXiv preprint arXiv:2502.19417* .
- [200] Shin, H.C., Roth, H.R., Gao, M., Lu, L., Xu, Z., Nogues, I., Yao, J., Mollura, D., Summers, R.M., 2016. Deep convolutional neural networks for computer-aided detection: Cnn architectures, dataset characteristics and transfer learning. *IEEE transactions on medical imaging* 35, 1285–1298.
- [201] Shivadekar, S., 2025. Artificial Intelligence for Cognitive Systems: Deep Learning, Neuro-symbolic Integration, and Human-Centric Intelligence. Deep Science Publishing.
- [202] Shridhar, M., Manuelli, L., Fox, D., 2022. Cliport: What and where pathways for robotic manipulation, in: Conference on robot learning, PMLR. pp. 894–906.
- [203] Shukor, M., Aubakirova, D., Capuano, F., Kooijmans, P., Palma, S., Zouitine, A., Aractingi, M., Pascal, C., Russi, M., Marafioti, A., et al., 2025. Smolvla: A vision-language-action model for affordable and efficient robotics. *arXiv preprint arXiv:2506.01844* .
- [204] Si, W., Wang, N., Yang, C., 2021. A review on manipulation skill acquisition through teleoperation-based learning from demonstration. *Cognitive Computation and Systems* 3, 1–16.
- [205] Simonyan, K., Zisserman, A., 2014. Very deep convolutional networks for large-scale image recognition. *arXiv preprint arXiv:1409.1556* .
- [206] Singh, G., 2025. Neural Object-Centric Scene Representation and Generation. Ph.D. thesis. Rutgers The State University of New Jersey, School of Graduate Studies.
- [207] Singh, S., Singh, J., Shah, B., Sehra, S.S., Ali, F., 2022. Augmented reality and gps-based resource efficient navigation system for outdoor environments: Integrating device camera, sensors, and storage. *Sustainability* 14, 12720.
- [208] Song, M., Deng, X., Zhou, Z., Wei, J., Guan, W., Nie, L., 2025a. A survey on diffusion policy for robotic manipulation: Taxonomy, analysis, and future directions. *Author Preprints* .
- [209] Song, W., Chen, J., Ding, P., Zhao, H., Zhao, W., Zhong, Z., Ge, Z., Ma, J., Li, H., 2025b. Accelerating vision-language-action model integrated with action chunking via parallel decoding. *arXiv preprint arXiv:2503.02310* .

- [210] Sun, H., Wang, H., Ma, C., Zhang, S., Ye, J., Chen, X., Lan, X., 2025a. Prism: Projection-based reward integration for scene-aware real-to-sim-to-real transfer with few demonstrations. *arXiv preprint arXiv:2504.20520* .
- [211] Sun, J., Mao, P., Kong, L., Wang, J., 2025b. A review of embodied grasping. *Sensors (Basel, Switzerland)* 25, 852.
- [212] Sun, L., Xie, B., Liu, Y., Shi, H., Wang, T., Cao, J., 2025c. Geovla: Empowering 3d representations in vision-language-action models. *arXiv preprint arXiv:2508.09071* .
- [213] Sutskever, I., Martens, J., Hinton, G.E., 2011. Generating text with recurrent neural networks, in: *Proceedings of the 28th international conference on machine learning (ICML-11)*, pp. 1017–1024.
- [214] Szot, A., Mazouze, B., Agrawal, H., Hjelm, R.D., Kira, Z., Toshev, A., 2024. Grounding multimodal large language models in actions. *Advances in Neural Information Processing Systems* 37, 20198–20224.
- [215] Taherin, A., Lin, J., Akbari, A., Akbari, A., Zhao, P., Chen, W., Kaeli, D., Wang, Y., 2025. Cross-platform scaling of vision-language-action models from edge to cloud gpus. *arXiv preprint arXiv:2509.11480* .
- [216] Tan, X., Yang, Y., Ye, P., Zheng, J., Bai, B., Wang, X., Hao, J., Chen, T., 2025. Think twice, act once: Token-aware compression and action reuse for efficient inference in vision-language-action models. *arXiv preprint arXiv:2505.21200* .
- [217] Team, G.R., Abeyruwan, S., Ainslie, J., Alayrac, J.B., Arenas, M.G., Armstrong, T., Balakrishna, A., Baruch, R., Bauza, M., Blokzijl, M., et al., 2025. Gemini robotics: Bringing ai into the physical world. *arXiv preprint arXiv:2503.20020* .
- [218] Team, O.M., Ghosh, D., Walke, H., Pertsch, K., Black, K., Mees, O., Dasari, S., Hejna, J., Kreiman, T., Xu, C., et al., 2024. Octo: An open-source generalist robot policy. *arXiv preprint arXiv:2405.12213* .
- [219] Tellex, S., Gopalan, N., Kress-Gazit, H., Matuszek, C., 2020. Robots that use language. *Annual Review of Control, Robotics, and Autonomous Systems* 3, 25–55.
- [220] Tian, H., Wang, T., Liu, Y., Qiao, X., Li, Y., 2020. Computer vision technology in agricultural automation—a review. *Information processing in agriculture* 7, 1–19.
- [221] Tian, K., Jiang, Y., Yuan, Z., Peng, B., Wang, L., 2024. Visual autoregressive modeling: Scalable image generation via next-scale prediction. *Advances in neural information processing systems* 37, 84839–84865.
- [222] Torres, N., Ulloa, C., Araya, I., Ayala, M., Jara, S., 2024. A comprehensive analysis of gender, racial, and prompt-induced biases in large language models. *International Journal of Data Science and Analytics* , 1–38.
- [223] Touvron, H., Martin, L., Stone, K., Albert, P., Almahairi, A., Babaei, Y., Bashlykov, N., Batra, S., Bhargava, P., Bhosale, S., et al., 2023. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288* .
- [224] Trivedi, C., Bhattacharya, P., Prasad, V.K., Patel, V., Singh, A., Tanwar, S., Sharma, R., Aluvala, S., Pau, G., Sharma, G., 2024. Explainable ai for industry 5.0: vision, architecture, and potential directions. *IEEE Open Journal of Industry Applications* .
- [225] Verbaan, L., 2024. Perception and control with large language models in robotic manipulation. *TU Delft Library* .
- [226] Vinod, K., Ramesh, P.J., Chakravarthi, B., et al., 2025. Sebvs: Synthetic event-based visual servoing for robot navigation and manipulation. *arXiv preprint arXiv:2508.17643* .
- [227] Vuong, Q., Levine, S., Walke, H.R., Pertsch, K., Singh, A., Doshi, R., Xu, C., Luo, J., Tan, L., Shah, D., et al., 2023. Open x-embodiment: Robotic learning datasets and rt-x models, in: *Towards Generalist Robots: Learning Paradigms for Scalable Skill Acquisition@ CoRL2023*.
- [228] Waite, J.R., Hasan, M.Z., Liu, Q., Jiang, Z., Hegde, C., Sarkar, S., 2025. Rls3: RL-based synthetic sample selection to enhance spatial reasoning in vision-language models for indoor autonomous perception, in: *Proceedings of the ACM/IEEE 16th International Conference on Cyber-Physical Systems (with CPS-IoT Week 2025)*, Association for Computing Machinery, New York, NY, USA. doi:10.1145/3716550.3722033.
- [229] Wang, G., Bai, L., Nah, W.J., Wang, J., Zhang, Z., Chen, Z., Wu, J., Islam, M., Liu, H., Ren, H., 2024a. Surgical-llm: Learning to adapt large vision-language model for grounded visual question answering in robotic surgery. *arXiv preprint arXiv:2405.10948* .
- [230] Wang, H., Xing, Z., Wu, W., Yang, Y., Tang, Q., Zhang, M., Xu, Y., Zhu, L., 2024b. Non-invasive to invasive: Enhancing ffa synthesis from cfp with a benchmark dataset and a novel network, in: *Proceedings of the 1st International Workshop on Multimedia Computing for Health and Medicine*, pp. 7–15.
- [231] Wang, J., Guo, D., Liu, H., 2025a. Where to learn: Embodied perception learning planned by vision-language models. *IEEE Transactions on Cognitive and Developmental Systems* .

- [232] Wang, S., 2025. Roboflamingo-plus: Fusion of depth and rgb perception with vision-language models for enhanced robotic manipulation. *arXiv preprint arXiv:2503.19510* .
- [233] Wang, T., Han, C., Liang, J.C., Yang, W., Liu, D., Zhang, L.X., Wang, Q., Luo, J., Tang, R., 2024c. Exploring the adversarial vulnerabilities of vision-language-action models in robotics. *arXiv preprint arXiv:2411.13587* .
- [234] Wang, Y., Liu, Q., Jiang, Z., Wang, T., Jiao, J., Chu, H., Gao, B., Chen, H., 2025b. Rad: Retrieval-augmented decision-making of meta-actions with vision-language models in autonomous driving, in: *Proceedings of the Computer Vision and Pattern Recognition Conference*, pp. 3838–3848.
- [235] Wang, Y., Niu, X., Ba, J., Su, Z., Du, L., 2025c. Navigating embodied intelligence: Enabling technologies, security and privacy, and emerging trends. *IEEE Internet of Things Journal* .
- [236] Wang, Y., Wu, S., Zhang, Y., Yan, S., Liu, Z., Luo, J., Fei, H., 2025d. Multimodal chain-of-thought reasoning: A comprehensive survey. *arXiv preprint arXiv:2503.12605* .
- [237] Wang, Z., Zhou, Z., Song, J., Huang, Y., Shu, Z., Ma, L., 2024d. Towards testing and evaluating vision-language-action models for robotic manipulation: An empirical study. *arXiv preprint arXiv:2409.12894* .
- [238] Wei, C., Guo, C., Zhang, J., Shan, H., Xu, Y., Zhang, Z., Liu, Y., Wang, Q., Zhou, C., Li, H., et al., 2025. Focus: A streaming concentration architecture for efficient vision-language models. *arXiv preprint arXiv:2512.14661* .
- [239] Wei, J., Yuan, S., Li, P., Hu, Q., Gan, Z., Ding, W., 2024. Occllama: An occupancy-language-action generative world model for autonomous driving. *arXiv preprint arXiv:2409.03272* .
- [240] Wen, J., Zhu, M., Zhu, Y., Tang, Z., Li, J., Zhou, Z., Li, C., Liu, X., Peng, Y., Shen, C., et al., 2024. Diffusion-vla: Scaling robot foundation models via unified diffusion and autoregression. *arXiv preprint arXiv:2412.03293* .
- [241] Wen, J., Zhu, Y., Li, J., Tang, Z., Shen, C., Feng, F., 2025a. Dexvla: Vision-language model with plug-in diffusion expert for general robot control. *arXiv preprint arXiv:2502.05855* .
- [242] Wen, J., Zhu, Y., Li, J., Zhu, M., Tang, Z., Wu, K., Xu, Z., Liu, N., Cheng, R., Shen, C., et al., 2025b. Tinyvla: Towards fast, data-efficient vision-language-action models for robotic manipulation. *IEEE Robotics and Automation Letters* .
- [243] Woo, S., Debnath, S., Hu, R., Chen, X., Liu, Z., Kweon, I.S., Xie, S., 2023. Convnext v2: Co-designing and scaling convnets with masked autoencoders, in: *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pp. 16133–16142.
- [244] Wu, J., Zhong, M., Xing, S., Lai, Z., Liu, Z., Chen, Z., Wang, W., Zhu, X., Lu, L., Lu, T., et al., 2024a. Visionllm v2: An end-to-end generalist multimodal large language model for hundreds of vision-language tasks. *Advances in Neural Information Processing Systems* 37, 69925–69975.
- [245] Wu, W., Feng, X., Gao, Z., Kan, Y., 2024b. Smart: scalable multi-agent real-time motion generation via next-token prediction. *Advances in Neural Information Processing Systems* 37, 114048–114071.
- [246] Wu, Z., Zhou, Y., Xu, X., Wang, Z., Yan, H., 2025. Momanipvla: Transferring vision-language-action models for general mobile manipulation. *arXiv preprint arXiv:2503.13446* .
- [247] Xiang, T.Y., Jin, A.Q., Zhou, X.H., Gui, M.J., Xie, X.L., Liu, S.Q., Wang, S.Y., Duang, S.B., Wang, S.C., Lei, Z., et al., 2025. Vla model-expert collaboration for bi-directional manipulation learning. *arXiv preprint arXiv:2503.04163* .
- [248] Xiong, J., Liu, G., Huang, L., Wu, C., Wu, T., Mu, Y., Yao, Y., Shen, H., Wan, Z., Huang, J., et al., 2024. Autoregressive models in vision: A survey. *arXiv preprint arXiv:2411.05902* .
- [249] Xu, D., Chen, Y., Wang, J., Huang, Y., Wang, H., Jin, Z., Wang, H., Yue, W., He, J., Li, H., et al., 2024a. Mlevlm: Improve multi-level progressive capabilities based on multimodal large language model for medical visual question answering, in: *Findings of the Association for Computational Linguistics ACL 2024*, pp. 4977–4997.
- [250] Xu, F., Zhai, G., Kong, X., Fu, T., Gordon, D.F., An, X., Busam, B., 2025a. Stare-vla: Progressive stage-aware reinforcement for fine-tuning vision-language-action models. *arXiv preprint arXiv:2512.05107* .
- [251] Xu, J., Sun, Q., Han, Q.L., Tang, Y., 2025b. When embodied ai meets industry 5.0: human-centered smart manufacturing. *IEEE/CAA Journal of Automatica Sinica* 12, 485–501.
- [252] Xu, S., Wang, Y., Xia, C., Zhu, D., Huang, T., Xu, C., 2025c. Vla-cache: Towards efficient vision-language-action model via adaptive token caching in robotic manipulation. *arXiv preprint arXiv:2502.02175* .
- [253] Xu, Y., Liu, G., Kompella, R.R., Hu, S., Huang, T., Ilhan, F., Tekin, S.F., Yahn, Z., Liu, L., 2025d. Language-vision planner and executor for text-to-visual reasoning. *arXiv preprint arXiv:2506.07778* .
- [254] Xu, Z., Wu, K., Wen, J., Li, J., Liu, N., Che, Z., Tang, J., 2024b. A survey on robotics with foundation models: toward embodied ai. *arXiv preprint arXiv:2402.02385* .

- [255] Xue, H., Ren, J., Chen, W., Zhang, G., Fang, Y., Gu, G., Xu, H., Lu, C., 2025. Reactive diffusion policy: Slow-fast visual-tactile policy learning for contact-rich manipulation. *arXiv preprint arXiv:2503.02881*.
- [256] Yang, G., Zhang, T., Hao, H., Wang, W., Liu, Y., Wang, D., Chen, G., Cai, Z., Chen, J., Su, W., et al., 2025a. Vlaser: Vision-language-action model with synergistic embodied reasoning. *arXiv preprint arXiv:2510.11027*.
- [257] Yang, R., Chen, G., Wen, C., Gao, Y., 2025b. Fp3: A 3d foundation policy for robotic manipulation. *arXiv preprint arXiv:2503.08950*.
- [258] Yang, R., Yu, Q., Wu, Y., Yan, R., Li, B., Cheng, A.C., Zou, X., Fang, Y., Cheng, X., Qiu, R.Z., et al., 2025c. Egovla: Learning vision-language-action models from egocentric human videos. *arXiv preprint arXiv:2507.12440*.
- [259] Yang, S., Li, Y.L., Wang, S., 2025d. Upl-net: Uncertainty-aware prompt learning network for semi-supervised action recognition. *Neurocomputing* 619, 129126.
- [260] Yang, Y., Huang, W., Wei, Y., Peng, H., Jiang, X., Jiang, H., Wei, F., Wang, Y., Hu, H., Qiu, L., et al., 2023a. Attentive mask clip, in: *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pp. 2771–2781.
- [261] Yang, Y., Sun, J., Kou, S., Wang, Y., Deng, Z., 2025e. Lohovla: A unified vision-language-action model for long-horizon embodied tasks. *arXiv preprint arXiv:2506.00411*.
- [262] Yang, Y., Wang, Y., Wen, Z., Zhongwei, L., Zou, C., Zhang, Z., Wen, C., Zhang, L., 2025f. Efficientvla: Training-free acceleration and compression for vision-language-action models. *arXiv preprint arXiv:2506.10100*.
- [263] Yang, Y., Zhou, J., Ding, X., Huai, T., Liu, S., Chen, Q., Xie, Y., He, L., 2025g. Recent advances of foundation language models-based continual learning: A survey. *ACM Computing Surveys* 57, 1–38.
- [264] Yang, Z., Chen, Y., Wang, J., Manivasagam, S., Ma, W.C., Yang, A.J., Urtasun, R., 2023b. Unisim: A neural closed-loop sensor simulator, in: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 1389–1399.
- [265] Yang, Z., Garrett, C., Fox, D., Lozano-Pérez, T., Kaelbling, L.P., 2025h. Guiding long-horizon task and motion planning with vision language models, in: *2025 IEEE International Conference on Robotics and Automation (ICRA)*, IEEE. pp. 16847–16853.
- [266] Ye, S., Jang, J., Jeon, B., Joo, S., Yang, J., Peng, B., Mandelkar, A., Tan, R., Chao, Y.W., Lin, B.Y., et al., 2024. Latent action pretraining from videos. *arXiv preprint arXiv:2410.11758*.
- [267] Ye, Y., Ma, J., Cen, J., Lu, Z., 2025. Token expand-merge: Training-free token compression for vision-language-action models. *arXiv preprint arXiv:2512.09927*.
- [268] Yu, Z., Wang, B., Zeng, P., Zhang, H., Zhang, J., Gao, L., Song, J., Sebe, N., Shen, H.T., 2025. A survey on efficient vision-language-action models. *arXiv preprint arXiv:2510.24795*.
- [269] Yue, Y., Wang, Y., Kang, B., Han, Y., Wang, S., Song, S., Feng, J., Huang, G., 2024. Deer-vla: Dynamic inference of multimodal large language models for efficient robot execution. *Advances in Neural Information Processing Systems* 37, 56619–56643.
- [270] Zawalski, M., Chen, W., Pertsch, K., Mees, O., Finn, C., Levine, S., 2024. Robotic control via embodied chain-of-thought reasoning. *arXiv preprint arXiv:2407.08693*.
- [271] Zhai, X., Mustafa, B., Kolesnikov, A., Beyer, L., 2023. Sigmoid loss for language image pre-training, in: *Proceedings of the IEEE/CVF international conference on computer vision*, pp. 11975–11986.
- [272] Zhan, Z., Chen, Y., Zhou, J., Lv, Q., Liu, H., Wang, K., Lin, L., Wang, G., 2026. Stable language guidance for vision-language-action models. *arXiv preprint arXiv:2601.04052*.
- [273] Zhang, B., Li, J., Shen, J., Cai, Y., Zhang, Y., Chen, Y., Dai, J., Ji, J., Yang, Y., 2025a. Vla-arena: An open-source framework for benchmarking vision-language-action models. *arXiv preprint arXiv:2512.22539*.
- [274] Zhang, B., Zhang, Y., Ji, J., Lei, Y., Dai, J., Chen, Y., Yang, Y., 2025b. Safevla: Towards safety alignment of vision-language-action model via safe reinforcement learning. *arXiv preprint arXiv:2503.03480*.
- [275] Zhang, D., Sun, J., Hu, C., Wu, X., Yuan, Z., Zhou, R., Shen, F., Zhou, Q., 2025c. Pure vision language action (vla) models: A comprehensive survey. *arXiv preprint arXiv:2509.19012*.
- [276] Zhang, H., Ding, P., Lyu, S., Peng, Y., Wang, D., 2025d. Gevrn: Goal-expressive video generation model for robust visual manipulation. *arXiv preprint arXiv:2502.09268*.
- [277] Zhang, H., Yu, H., Zhao, L., Choi, A., Bai, Q., Yang, B., Xu, W., 2025e. Slim: Sim-to-real legged instructive manipulation via long-horizon visuomotor learning. *arXiv preprint arXiv:2501.09905*.

- [278] Zhang, H., Zantout, N., Kachana, P., Wu, Z., Zhang, J., Wang, W., 2024a. Vla-3d: A dataset for 3d semantic scene understanding and navigation. *arXiv preprint arXiv:2411.03540* .
- [279] Zhang, J., Guo, Y., Hu, Y., Chen, X., Zhu, X., Chen, J., 2025f. Up-vla: A unified understanding and prediction model for embodied agent. *arXiv preprint arXiv:2501.18867* .
- [280] Zhang, J., Wang, K., Wang, S., Li, M., Liu, H., Wei, S., Wang, Z., Zhang, Z., Wang, H., 2024b. Uni-navid: A video-based vision-language-action model for unifying embodied navigation tasks. *arXiv preprint arXiv:2412.06224* .
- [281] Zhang, K., Yin, Z.H., Ye, W., Gao, Y., 2024c. Learning manipulation skills through robot chain-of-thought with sparse failure guidance. *arXiv preprint arXiv:2405.13573* .
- [282] Zhang, K., Yun, P., Cen, J., Cai, J., Zhu, D., Yuan, H., Zhao, C., Feng, T., Wang, M.Y., Chen, Q., et al., 2025g. Generative artificial intelligence in robotic manipulation: A survey. *arXiv preprint arXiv:2503.03464* .
- [283] Zhang, R., Dong, M., Zhang, Y., Heng, L., Chi, X., Dai, G., Du, L., Wang, D., Du, Y., Zhang, S., 2025h. Mole-vla: Dynamic layer-skipping vision language action model via mixture-of-layers for efficient robot manipulation. *arXiv preprint arXiv:2503.20384* .
- [284] Zhang, Z., Bao, C., Pan, X., Chang, C.M., Igarashi, T., Zhang, G., 2025i. Through the lens of privacy: Exploring privacy protection in vision-language model interactions on smart glasses, in: *Proceedings of the Extended Abstracts of the CHI Conference on Human Factors in Computing Systems*, pp. 1–8.
- [285] Zhao, H., Song, W., Wang, D., Tong, X., Ding, P., Cheng, X., Ge, Z., 2025a. More: Unlocking scalability in reinforcement learning for quadruped vision-language-action models. *arXiv preprint arXiv:2503.08007* .
- [286] Zhao, Q., Lu, Y., Kim, M.J., Fu, Z., Zhang, Z., Wu, Y., Li, Z., Ma, Q., Han, S., Finn, C., et al., 2025b. Cot-vla: Visual chain-of-thought reasoning for vision-language-action models. *arXiv preprint arXiv:2503.22020* .
- [287] Zhao, T.Z., Kumar, V., Levine, S., Finn, C., 2023. Learning fine-grained bimanual manipulation with low-cost hardware. *arXiv preprint arXiv:2304.13705* .
- [288] Zhen, H., Qiu, X., Chen, P., Yang, J., Yan, X., Du, Y., Hong, Y., Gan, C., 2024. 3d-vla: A 3d vision-language-action generative world model. *arXiv preprint arXiv:2403.09631* .
- [289] Zheng, J., Li, J., Liu, D., Zheng, Y., Wang, Z., Ou, Z., Liu, Y., Liu, J., Zhang, Y.Q., Zhan, X., 2025. Universal actions for enhanced embodied foundation models. *arXiv preprint arXiv:2501.10105* .
- [290] Zheng, J., Shi, C., Cai, X., Li, Q., Zhang, D., Li, C., Yu, D., Ma, Q., 2026. Lifelong learning of large language model based agents: A roadmap. *IEEE Transactions on Pattern Analysis and Machine Intelligence* .
- [291] Zhong, Y., Huang, X., Li, R., Zhang, C., Liang, Y., Yang, Y., Chen, Y., 2025. Dexgraspvla: A vision-language-action framework towards general dexterous grasping. *arXiv preprint arXiv:2502.20900* .
- [292] Zhou, D.W., Zhang, Y., Wang, Y., Ning, J., Ye, H.J., Zhan, D.C., Liu, Z., 2025a. Learning without forgetting for vision-language models. *IEEE Transactions on Pattern Analysis and Machine Intelligence* .
- [293] Zhou, X., Han, X., Yang, F., Ma, Y., Knoll, A.C., 2025b. Opendrivevla: Towards end-to-end autonomous driving with large vision language action model. *arXiv preprint arXiv:2503.23463* .
- [294] Zhou, Z., Cai, T., Zhao, S.Z., Zhang, Y., Huang, Z., Zhou, B., Ma, J., 2025c. Autovla: A vision-language-action model for end-to-end autonomous driving with adaptive reasoning and reinforcement fine-tuning. *arXiv preprint arXiv:2506.13757* .
- [295] Zhou, Z., Zhu, Y., Zhu, M., Wen, J., Liu, N., Xu, Z., Meng, W., Cheng, R., Peng, Y., Shen, C., et al., 2025d. Chatvla: Unified multimodal understanding and robot control with vision-language-action model. *arXiv preprint arXiv:2502.14420* .
- [296] Zhu, D.H., Chang, Y.P., 2020. Robot with humanoid hands cooks food better? effect of robotic chef anthropomorphism on food quality prediction. *International Journal of Contemporary Hospitality Management* 32, 1367–1383.
- [297] Zhu, M., Zhu, Y., Li, J., Zhou, Z., Wen, J., Liu, X., Shen, C., Peng, Y., Feng, F., 2025. Objectvla: End-to-end open-world object manipulation without demonstration. *arXiv preprint arXiv:2502.19250* .
- [298] Zhu, Y., Zhou, Y., Wang, C., Cao, Y., Han, J., Hou, L., Xu, H., 2024. Unit: Unifying image and text recognition in one vision encoder, in: *NeurIPS*.
- [299] Zitkovich, B., Yu, T., Xu, S., Xu, P., Xiao, T., Xia, F., Wu, J., Wohlhart, P., Welker, S., Wahid, A., et al., 2023. Rt-2: Vision-language-action models transfer web knowledge to robotic control, in: *Conference on Robot Learning*, PMLR. pp. 2165–2183.