



(19) 대한민국특허청(KR)

(12) 등록특허공보(B1)

(45) 공고일자 2022년04월29일

(11) 등록번호 10-2392594

(24) 등록일자 2022년04월26일

(51) 국제특허분류(Int. Cl.)

B25J 9/16 (2006.01) G06N 3/00 (2022.01)

G06N 3/08 (2006.01) G06T 7/11 (2017.01)

(52) CPC특허분류

B25J 9/1602 (2013.01)

B25J 9/1664 (2013.01)

(21) 출원번호 10-2021-0007631

(22) 출원일자 2021년01월19일

심사청구일자 2021년01월19일

(56) 선행기술조사문헌

KR1020200034025 A*

*는 심사관에 의하여 인용된 문헌

(73) 특허권자

고려대학교 산학협력단

서울특별시 성북구 안암로 145, 고려대학교 (안암동5가)

(72) 발명자

송재복

서울특별시 강남구 인주로30길 56, E동2907호(도곡동, 타워팰리스)

김선인

경기도 이천시 장감로77번길 55 (장호원읍)

(뒷면에 계속)

(74) 대리인

특허법인남춘

전체 청구항 수 : 총 16 항

심사관 : 이성수

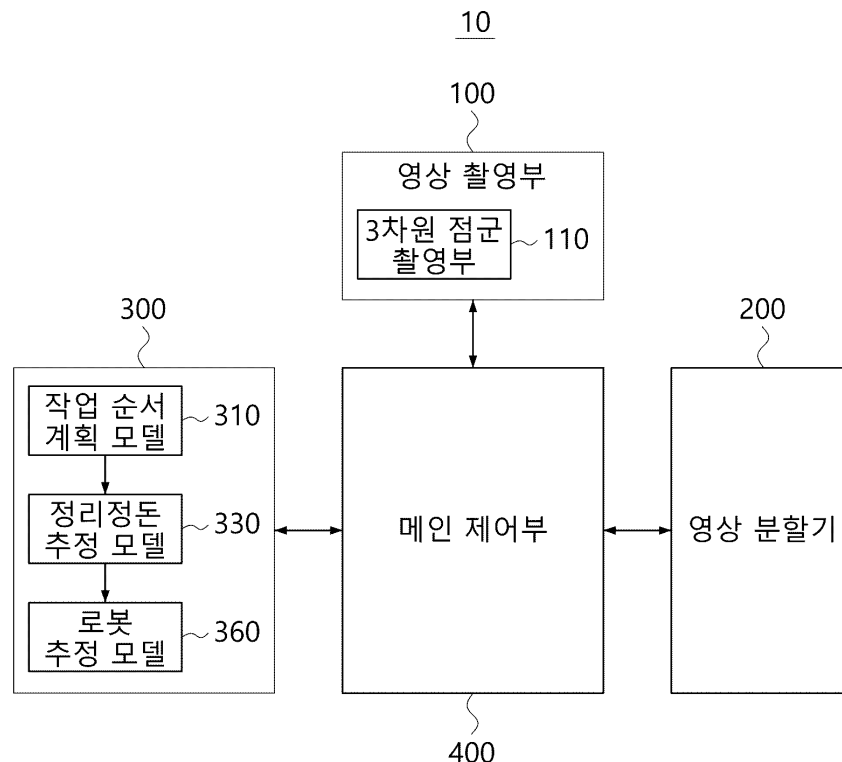
(54) 발명의 명칭 정리정돈 로봇을 이용한 정리정돈 시스템

(57) 요약

본 발명은 정리정돈 로봇을 이용한 정리정돈 시스템에 관한 것으로, 다수의 대상 물체가 분포된 대상 영역을 촬영하는 영상 촬영부와, 목표 영상과 상기 영상 촬영부에 의해 촬영된 현재 영상을 각각 분할하여 목표 분할 영상과 현재 분할 영상을 생성하는 인공지능 기반의 영상 분할기와, 상기 영상 분할기에 의해 생성된 상기 목표 분할

(뒷면에 계속)

대표도 - 도2



영상과 상기 현재 분할 영상을 입력받아, 다수의 상기 대상 물체의 정리정돈 계획을 수립하는 인공지능 기반의 정리정돈 계획기를 포함하며; 상기 정리정돈 계획기는 상기 목표 분할 영상과 상기 현재 분할 영상을 입력받아 다수의 상기 대상 물체 중 정리정돈의 다음 대상을 추정하는 작업 순서 계획 모델과, 상기 작업 순서 계획 모델에 의해 추정된 상기 다음 대상에 대한 정리정돈 동작을 추정하는 정리정돈 추정 모델과, 상기 정리정돈 추정 모델에 의해 추정된 상기 정리정돈 동작에 기초하여, 상기 다음 대상의 정리정돈을 위한 상기 정리정돈 로봇의 로봇 동작을 추정하는 로봇 추정 모델을 포함하는 것을 특징으로 한다. 이에 따라, 정리정돈 시스템이 인공지능을 기반으로 자체적으로 정리정돈 계획, 즉 다음 대상, 다음 대상의 정리정돈 동작, 그리고, 이에 따른 로봇 동작을 추정하여 실행함으로써, 목표 영상과 현재 영상의 입력 만으로 대상 영역의 정리정돈이 가능하게 된다.

(52) CPC특허분류

B25J 9/1679 (2013.01)

B25J 9/1697 (2013.01)

G06N 3/008 (2019.01)

G06N 3/08 (2013.01)

G06T 7/11 (2017.01)

(72) 발명자

양민규

경기도 성남시 분당구 수내로 42, 1109호 (수내동, 두산위브센터움2)

문종술

서울특별시 성북구 종암로 13길 10-22 2-1호

박정관

서울특별시 성북구 안암로 145 안암글로벌하우스

명세서

청구범위

청구항 1

정리정돈 로봇을 이용한 정리정돈 시스템에 있어서,

다수의 대상 물체가 분포된 대상 영역을 촬영하는 영상 촬영부와,

목표 영상과 상기 영상 촬영부에 의해 촬영된 현재 영상을 각각 분할하여 목표 분할 영상과 현재 분할 영상을 생성하는 인공지능 기반의 영상 분할기와,

상기 영상 분할기에 의해 생성된 상기 목표 분할 영상과 상기 현재 분할 영상을 입력받아, 다수의 상기 대상 물체의 정리정돈 계획을 수립하는 인공지능 기반의 정리정돈 계획기를 포함하며;

상기 정리정돈 계획기는

상기 목표 분할 영상과 상기 현재 분할 영상을 입력받아 다수의 상기 대상 물체 중 정리정돈의 다음 대상을 추정하는 작업 순서 계획 모델과,

상기 작업 순서 계획 모델에 의해 추정된 상기 다음 대상에 대한 정리정돈 동작을 추정하는 정리정돈 추정 모델과,

상기 정리정돈 추정 모델에 의해 추정된 상기 정리정돈 동작에 기초하여, 상기 다음 대상의 정리정돈을 위한 상기 정리정돈 로봇의 로봇 동작을 추정하는 로봇 추정 모델을 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 2

제1항에 있어서,

상기 영상 분할기는 기 학습되어 등록된 DeepLap v3+ 알고리즘을 이용하여 상기 목표 분할 영상 및 상기 현재 분할 영상을 생성하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 3

제1항에 있어서,

상기 작업 순서 계획 모델은 심층강화학습모델에 기반하여 학습되어 생성되는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 4

제3항에 있어서,

상기 심층강화학습모델은 가상 영상이 상기 영상 분할기에 의해 분할되어 생성된 가상 분할 영상을 학습 데이터로 하여 학습되어 생성되며;

상기 가상 영상은 각각의 상기 대상 물체에 대응하는 가상 물체가 상기 대상 영역에 대응하는 가상 공간에 분포된 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 5

제3항에 있어서,

상기 심층강화학습모델은 Rainbow DQN(Deep Q-network) 모델을 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 6

제3항에 있어서,

상기 작업 순서 계획 모델은 상기 다음 대상과, 상기 다음 대상의 정리정돈 기본 동작을 추정하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 7

제6항에 있어서,

상기 정리정돈 추정 모델은

상기 다음 대상의 기계적 상태 변화에 따른 상태 조작 동작을 추정하는 정리정돈 상태 추정 모델과;

상기 다음 대상의 자세 변화에 따른 자세 조작 동작을 추정하는 정리정돈 자세 추정 모델을 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 8

제7항에 있어서,

상기 상태 조작 동작은 상기 다음 대상의 상태 조작 위치, 상기 다음 대상의 상태 조작 방법, 및 상기 다음 대상의 상태 영역을 포함하며;

상기 자세 조작 동작은 상기 다음 대상의 이동 동작과 상기 다음 대상의 회전 동작 중 적어도 하나를 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 9

제8항에 있어서,

상기 정리정돈 상태 추정 모델은 OUNET(Object understanding network)을 기반으로 상기 상태 조작 동작을 추정하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 10

제9항에 있어서,

상기 정리정돈 상태 추정 모델은

상기 현재 영상으로부터 상기 다음 대상을 추출하는 대상 추출부와,

상기 대상 추출부에 의해 추출된 상기 다음 대상의 상기 상태 조작 위치, 상기 상태 조작 방법 및 상기 상태 영역을 추정하는 물체 특성 추정기를 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 11

제10항에 있어서,

상기 물체 특성 추정기는 U-Net 구조를 기반으로 구성되며;

상기 대상 추출부에 의해 추출된 상기 다음 대상의 특징 벡터를 추출하는 엔코더와,

상기 엔코더와 연결되어 상기 상태 영역 및 상기 상태 조작 위치를 추정하는 디코더와,

상기 엔코더에 의해 추출된 상기 특징 벡터를 입력받아 상기 상태 조작 방법을 분류하는 전연결층을 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 12

제11항에 있어서,

상기 정리정돈 상태 추정 모델은

상기 상태 영역 및 상기 상태 조작 위치 각각의 픽셀의 밀도 및 픽셀의 거리를 기반으로 군집화하여 기 설정된 기준 밀도 이하를 상기 상태 영역 및 상기 상태 조작 위치로부터 제거하여 최종적인 상태 영역 및 상태 조작 위

치를 추정하는 DBSCAM(Density-based spatial clustering of application with noise) 모델을 더 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 13

제8항에 있어서,

상기 영상 촬영부는 상기 대상 영역을 촬영하기 위한 3차원 점군 촬영부를 포함하며;

상기 정리정돈 자세 추정 모델은

상기 3차원 점군 촬영부에 의해 상기 대상 영역에 대해 복수 방향에서 촬영된 3차원 점군 데이터를 상기 대상 물체별로 분할하는 점군 분할 모델과;

상기 점군 분할 모델에 의해 분할된 상기 대상 물체 중 상기 다음 대상에 대한 기준 좌표계를 설정하는 상기 자세 조작 동작을 추정하는 PCA(Principal component analysis) 모델을 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 14

제13항에 있어서,

상기 점군 분할 모델은

상기 3차원 점군 데이터를 상기 정리정돈 로봇에 대한 기준 좌표계로 변환하여 정합하는 데이터 정합부와;

상기 데이터 정합부에 의해 정합된 상기 3차원 점군 데이터로부터 배경 점군을 제거하는 패스 스루 필터(Pass through filter)와;

상기 패스 스루 필터에 의해 상기 배경 점군이 제거된 상기 3차원 점군 데이터의 각 점군 간의 거리 및 색상 차이에 기반하여 복수의 패치로 분할하는 VCCS(Voxel cloud connectivity segmentation) 모델과;

상기 패치의 곡률에 기반하여 상기 대상 물체별로 군집화하는 LCCP(Locally convex connected patches) 모델을 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 15

제13항에 있어서,

상기 로봇 추정 모델은

상기 현재 영상을 이용하여 표준 영상, 분할 영상 및 깊이 영상을 생성하는 적대적 신경망 모델과;

상기 표준 영상, 상기 분할 영상 및 상기 깊이 영상을 기 설정된 사이즈로 압축하는 영상 압축 모델과;

상기 영상 압축 모델에 의해 압축된 상기 표준 영상, 상기 분할 영상 및 상기 깊이 영상과, 상기 정리정돈 로봇의 로봇 상태 정보, 상기 정리정돈 동작을 입력받아 상기 로봇 동작을 추정하는 동작 추정 모델을 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

청구항 16

제15항에 있어서,

상기 적대적 신경망 모델은 RCAN(Randomized-to-canonical adaptation networks) 모델을 포함하고;

상기 영상 압축 모델은 VAE(Variational auto-encoder) 모델을 포함하며;

상기 동작 추정 모델은 SAC(Soft actor critic) 모델을 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템.

발명의 설명

기술 분야

본 발명은 정리정돈 로봇을 이용한 정리정돈 시스템에 관한 것으로서, 보다 상세하게는 불규칙하게 어질러진 대

상 물체를 정리정돈 로봇을 이용하여 기 설정된 목표 영상과 같이 정리정돈할 수 있는 정리정돈 로봇을 이용한 정리정돈 시스템에 관한 것이다.

배경 기술

- [0002] 로봇 기술이 발달하면서 산업용 로봇 뿐만 아니라 점차 가정용 로봇도 많은 연구가 진행되고 있다. 특히, 가정용 로봇은 청소 로봇이 가장 활발하게 상업화가 되어 있는 분야이다.
- [0003] 청소 로봇은 집안 바닥에 떨어져있는 먼지 등과 같은 작은 이물질들을 흡입하는 방식으로 바닥 청소를 수행하는 것으로, 아직 집안에 어질러진 물건을 정리정돈하는 수준에는 이르지 못하고 있다.
- [0004] 로봇 기술에 있어, 현재 기술 수준의 로봇은 다양한 환경에서 영상과 점군(Point cloud)을 통해 상황을 분석하고, 이를 토대로 목표 작업을 수행하지는 못한다. 즉, 일반적으로 로봇은 작업자에 의해 정해진 순서에 따라 정해진 위치에 있는 물체들만 조작할 수 있으므로, 로봇 기술을 정리정돈 로봇에 적용하더라도, 다양한 작업 환경에 대해 작업자가 미리 작업의 순서를 정해두어야 하지만, 이는 사실상 불가능하다.
- [0005] 한국등록특허 제10-2181815호 "개인 맞춤형 정리로봇"에서는 물건의 정리정돈에 로봇을 이용하는 기술을 개시하고 있고, 딥러닝과 같은 인공지능 기술을 적용하고 있으나, 단순히 물건의 위치를 예측하는데 적용하고 있을 뿐, 로봇이 자신의 작업을 계획하고, 이를 적합하게 실행하지는 못하고 있다.

발명의 내용

해결하려는 과제

- [0006] 이에, 본 발명은 상기와 같은 문제점을 해소하기 위해 안출된 것으로서, 정리정돈 로봇이 스스로 수행할 작업을 결정하여 불규칙하게 어질러진 대상 물체를 기 설정된 목표 영상과 같이 정리정돈할 수 있는 정리정돈 로봇을 이용한 정리정돈 시스템을 제공하는데 그 목적이 있다.

과제의 해결 수단

- [0007] 상기 목적은 본 발명에 따라, 정리정돈 로봇을 이용한 정리정돈 시스템에 있어서, 다수의 대상 물체가 분포된 대상 영역을 촬영하는 영상 촬영부와, 목표 영상과 상기 영상 촬영부에 의해 촬영된 현재 영상을 각각 분할하여 목표 분할 영상과 현재 분할 영상을 생성하는 인공지능 기반의 영상 분할기와, 상기 영상 분할기에 의해 생성된 상기 목표 분할 영상과 상기 현재 분할 영상을 입력받아, 다수의 상기 대상 물체의 정리정돈 계획을 수립하는 인공지능 기반의 정리정돈 계획기를 포함하며; 상기 정리정돈 계획기는 상기 목표 분할 영상과 상기 현재 분할 영상을 입력받아 다수의 상기 대상 물체 중 정리정돈의 다음 대상을 추정하는 작업 순서 계획 모델과, 상기 작업 순서 계획 모델에 의해 추정된 상기 다음 대상에 대한 정리정돈 동작을 추정하는 정리정돈 추정 모델과, 상기 정리정돈 추정 모델에 의해 추정된 상기 정리정돈 동작에 기초하여, 상기 다음 대상의 정리정돈을 위한 상기 정리정돈 로봇의 로봇 동작을 추정하는 로봇 추정 모델을 포함하는 것을 특징으로 하는 정리정돈 로봇을 이용한 정리정돈 시스템에 의해서 달성된다.
- [0008] 여기서, 상기 영상 분할기는 기 학습되어 등록된 DeepLap v3+ 알고리즘을 이용하여 상기 목표 분할 영상 및 상기 현재 분할 영상을 생성할 수 있다.
- [0009] 또한, 상기 작업 순서 계획 모델은 심층강화학습모델에 기반하여 학습되어 생성될 수 있다.
- [0010] 그리고, 상기 심층강화학습모델은 가상 영상이 상기 영상 분할기에 의해 분할되어 생성된 가상 분할 영상을 학습 데이터로 하여 학습되어 생성되며; 상기 가상 영상은 각각의 상기 대상 물체에 대응하는 가상 물체가 상기 대상 영역에 대응하는 가상 공간에 분포될 수 있다.
- [0011] 그리고, 상기 심층강화학습모델은 Rainbow DQN(Deep Q-network) 모델을 포함할 수 있다.
- [0012] 그리고, 상기 작업 순서 계획 모델은 상기 다음 대상과, 상기 다음 대상의 정리정돈 기본 동작을 추정할 수 있다.
- [0013] 그리고, 상기 정리정돈 추정 모델은 상기 다음 대상의 기계적 상태 변화에 따른 상태 조작 동작을 추정하는 정리정돈 상태 추정 모델과; 상기 다음 대상의 자세 변화에 따른 자세 조작 동작을 추정하는 정리정돈 자세 추정 모델을 포함할 수 있다.

- [0014] 그리고, 상기 상태 조작 동작은 상기 다음 대상의 상태 조작 위치, 상기 다음 대상의 상태 조작 방법, 및 상기 다음 대상의 상태 영역을 포함하며; 상기 자세 조작 동작은 상기 다음 대상의 이동 동작과 상기 다음 대상의 회전 동작 중 적어도 하나를 포함할 수 있다.
- [0015] 그리고, 상기 정리정돈 상태 추정 모델은 OUNET(Object understanding network)을 기반으로 상기 상태 조작 동작을 추정할 수 있다.
- [0016] 그리고, 상기 정리정돈 상태 추정 모델은 상기 현재 영상으로부터 상기 다음 대상을 추출하는 대상 추출부와, 상기 대상 추출부에 의해 추출된 상기 다음 대상의 상기 상태 조작 위치, 상기 상태 조작 방법 및 상기 상태 영역을 추정하는 물체 특성 추정기를 포함할 수 있다.
- [0017] 그리고, 상기 물체 특성 추정기는 U-Net 구조를 기반으로 구성되며; 상기 대상 추출부에 의해 추출된 상기 다음 대상의 특징 벡터를 추출하는 엔코더와, 상기 엔코더와 연결되어 상기 상태 영역 및 상기 상태 조작 위치를 추정하는 디코더와, 상기 엔코더에 의해 추출된 상기 특징 벡터를 입력받아 상기 상태 조작 방법을 분류하는 전연결층을 포함할 수 있다.
- [0018] 그리고, 상기 정리정돈 상태 추정 모델은 상기 상태 영역 및 상기 상태 조작 위치 각각의 픽셀의 밀도 및 픽셀의 거리를 기반으로 군집화하여 기 설정된 기준 밀도 이하를 상기 상태 영역 및 상기 상태 조작 위치로부터 제거하여 최종적인 상태 영역 및 상태 조작 위치를 추정하는 DBSCAM(Density-based spatial clustering of application with noise) 모델을 더 포함할 수 있다.
- [0019] 그리고, 상기 영상 촬영부는 상기 대상 영역을 촬영하기 위한 3차원 점군 촬영부를 포함하며; 상기 정리정돈 자세 추정 모델은 상기 3차원 점군 촬영부에 의해 상기 대상 영역에 대해 복수 방향에서 촬영된 3차원 점군 데이터를 상기 대상 물체별로 분할하는 점군 분할 모델과; 상기 점군 분할 모델에 의해 분할된 상기 대상 물체 중 상기 다음 대상에 대한 기준 좌표계를 설정하는 상기 자세 조작 동작을 추정하는 PCA(Principal component analysis) 모델을 포함할 수 있다.
- [0020] 그리고, 상기 점군 분할 모델은 상기 3차원 점군 데이터를 상기 정리정돈 로봇에 대한 기준 좌표계로 변환하여 정합하는 데이터 정합부와; 상기 데이터 정합부에 의해 정합된 상기 3차원 점군 데이터로부터 배경 점군을 제거하는 패스 스루 필터(Pass through filter)와; 상기 패스 스루 필터에 의해 상기 배경 점군이 제거된 상기 3차원 점군 데이터의 각 점군 간의 거리 및 색상 차이에 기반하여 복수의 패치로 분할하는 VCCS(Voxel cloud connectivity segmentation) 모델과; 상기 패치의 곡률에 기반하여 상기 대상 물체별로 군집화하는 LCCP(Locally convex connected patches) 모델을 포함할 수 있다.
- [0021] 그리고, 상기 로봇 추정 모델은 상기 현재 영상을 이용하여 표준 영상, 분할 영상 및 깊이 영상을 생성하는 적대적 신경망 모델과; 상기 표준 영상, 상기 분할 영상 및 상기 깊이 영상을 기 설정된 사이즈로 압축하는 영상 압축 모델과; 상기 영상 압축 모델에 의해 압축된 상기 표준 영상, 상기 분할 영상 및 상기 깊이 영상과, 상기 정리정돈 로봇의 로봇 상태 정보, 상기 정리정돈 동작을 입력받아 상기 로봇 동작을 추정하는 동작 추정 모델을 포함할 수 있다.
- [0022] 그리고, 상기 적대적 신경망 모델은 RCAN(Randomized-to-canonical adaptation networks) 모델을 포함하고; 상기 영상 압축 모델은 VAE(Variational auto-encoder) 모델을 포함하며; 상기 동작 추정 모델은 SAC(Soft actor critic) 모델을 포함할 수 있다.

발명의 효과

- [0023] 상기와 같은 구성에 따라, 본 발명에 따르면 정리정돈 시스템이 인공지능을 기반으로 자체적으로 정리정돈 계획, 즉 다음 대상, 다음 대상의 정리정돈 동작, 그리고, 이에 따른 로봇 동작을 추정하여 실행함으로써, 목표 영상과 현재 영상의 입력 만으로 대상 영역의 정리정돈이 가능하게 된다.

도면의 간단한 설명

- [0024] 도 1은 본 발명의 실시예에 따른 정리정돈 로봇을 이용한 정리정돈 시스템의 구성을 나타낸 도면이고,
 도 2 내지 도 5는 본 발명의 실시예에 따른 정리정돈 로봇을 이용한 정리정돈 시스템의 제어블록도이고,
 도 6은 본 발명의 실시예에 따른 정리정돈 시스템에서 현재 영상, 목표 영상, 현재 분할 영상, 및 목표 분할 영상의 예를 나타낸 도면이고,

도 7은 본 발명의 실시예에 따른 정리정돈 시스템의 영상 분할기가 가상 영상과 실제 영상을 가상 분할 영상과 실제 분할 영상으로 분할하여 생성한 예를 나타낸 도면이고,

도 8은 본 발명의 실시예에 따른 정리정돈 시스템에서 U-Net 구조의 물체 특성 추정기의 예를 나타낸 도면이고,

도 9는 본 발명의 실시예에 따른 정리정돈 시스템에서 적대적 신경망 모델에 의해 생성되는 표준 영상, 분할 영상 및 깊이 영상의 예를 나타낸 도면이다.

발명을 실시하기 위한 구체적인 내용

- [0025] 본 발명의 이점 및 특징, 그리고 그것들을 달성하는 방법은 첨부되는 도면과 함께 상세하게 후술되어 있는 실시예들을 참조하면 명확해질 것이다. 그러나 본 발명은 이하에서 개시되는 실시예들에 한정되는 것이 아니라 서로 다른 다양한 형태로 구현될 수 있으며, 단지 본 실시예들은 본 발명의 개시가 완전하도록 하고, 본 발명이 속하는 기술 분야에서 통상의 지식을 가진 자에게 발명의 범주를 완전하게 알려주기 위해 제공되는 것이며, 본 발명은 청구항의 범주에 의해 정의될 뿐이다. 명세서 전체에 걸쳐 동일 참조 부호는 동일 구성 요소를 지칭한다.
- [0026] 이하에서는 첨부된 도면을 참조하여 본 발명에 따른 실시예들을 상세히 설명한다.
- [0027] 도 1은 본 발명의 실시예에 따른 정리정돈 로봇(20)을 이용한 정리정돈 시스템(100)의 구성을 나타낸 도면이고, 도 2 내지 도 5는 본 발명의 실시예에 따른 정리정돈 로봇(20)을 이용한 정리정돈 시스템(100)의 제어블록도이다.
- [0028] 도 1에 도시된 바와 같이, 본 발명의 실시예에 따른 정리정돈 시스템(100)은 정리정돈 로봇(20)을 이용하여, 대상 물체(0)가 다수 분포되어 있는 대상 영역(S)에서 대상 물체(0)를 정리정돈한다.
- [0029] 본 발명에서는 정리정돈 로봇(20)이 다관절 로봇 형태로 마련되고, 말단에 대상 물체(0)를 파지할 수 있는 그립퍼가 마련되어, 대상 물체(0)를 파지하여 이동 또는 회전시키거나, 대상 물체(0)의 상태, 예를 들어, 서랍을 닫거나 여는 등의 동작이 가능한 형태로 마련되는 것을 예로 한다.
- [0030] 도 2 내지 도 5를 참조하여 설명하면, 본 발명의 실시예에 따른 정리정돈 시스템(100)은 영상 촬영부(100), 영상 분할기(200) 및 정리정돈 계획기(300)를 포함할 수 있다.
- [0031] 영상 촬영부(100)는 다수의 대상 물체(0)가 분포된 대상 영역(S)을 촬영한다. 본 발명에서는 영상 촬영부(100)가 대상 영역(S)의 컬러 영상, 예컨대 RGB 영상을 촬영하는 것을 예로 한다. 이하에서는 영상 촬영부(100)에 의해 촬영된 영상을 "현재 영상"이라 정의하여 설명한다.
- [0032] 영상 분할기(200)는 인공 지능을 기반으로 하여 동작하며, 입력된 목표 영상과 영상 촬영부(100)에 의해 촬영된 현재 영상을 각각 분할하여, 목표 분할 영상과 현재 분할 영상을 생성할 수 있다.
- [0033] 여기서, 목표 영상은 대상 영역(S)에 대상 물체(0)들이 정리정돈된 상태에 대해 미리 촬영되어 등록된 것으로, 본 발명의 실시예에 따른 정리정돈 로봇(20)이 최종적으로 대상 물체(0)들을 정리정돈하기 위한 상태가 된다.
- [0034] 본 발명에서는 영상 분할기(200)가 기 학습되어 등록된 DeepLap v3+ 알고리즘을 이용하여 목표 분할 영상 및 현재 분할 영상을 생성하는 것을 예로 한다. 도 6은 본 발명의 실시예에 따른 정리정돈 시스템(100)에서 현재 영상, 목표 영상, 현재 분할 영상, 및 목표 분할 영상의 예를 나타낸 도면으로, 도 6의 (a)가 현재 영상이고, 도 6의 (b)가 목표 영상이고, 도 6의 (c)가 현재 분할 영상이고, 도 6의 (d)가 목표 분할 영상이다.
- [0035] 본 발명의 실시예에 따른 정리정돈 계획기(300)의 후술할 작업 순서 계획 모델(310)은 가상 영상이 영상 분할기(200)에 의해 분할되어 생성된 가상 분할 영상을 학습 데이터로 하여 학습되어 생성된다.
- [0036] 즉, 작업 순서 계획 모델(310)의 학습 과정에서 실제 영상을 사용하여 학습할 경우, 학습에 많은 학습 데이터가 필요하게 되는데, 이를 실제 영상으로부터 얻는데 소요되는 시간 및 비용적인 측면을 고려하여 가상 영상을 이용하여 학습이 수행된다.
- [0037] 따라서, 작업 순서 계획 모델(310)이 학습에 사용한 가상 영상과 실제 영상인 현재 영상과의 차이가 최소화되어, 작업 순서 계획 모델(310)의 추정 정확도를 높일 수 있도록 영상 분할기(200)에 의해 목표 영상과 현재 영상이 분할되어 목표 분할 영상 및 현재 분할 영상이 생성된다.
- [0038] 도 7은 본 발명의 실시예에 따른 정리정돈 시스템(100)의 영상 분할기(200)가 가상 영상과 실제 영상을 가상 분할 영상과 실제 분할 영상으로 분할하여 생성한 예를 나타낸 도면이다.

- [0039] 도 7의 (a)는 가상 영상이고, 도 7의 (b)는 가상 분할 영상이고, 도 7의 (c)는 실제 영상이고, 도 7의 (d)는 실제 분할 영상이다. 도 7에 도시된 바와 같이, 실제 영상과 가상 영상은 큰 차이를 보이지만, 영상 분할기(200)에 의해 분할되어 생성된 가상 분할 영상과 실제 분할 영상은 그 차이가 미비함을 확인할 수 있다.
- [0040] 본 발명의 실시예에 따른 정리정돈 계획기(300)는 영상 분할기(200)에 의해 생성된 목표 분할 영상과 현재 분할 영상을 입력받고, 인공지능 기술을 기반으로 하여, 다수의 대상 물체(0)의 정리정돈 계획을 수립한다.
- [0041] 본 발명의 실시예에 따른 정리정돈 계획기(300)는 전술한 작업 순서 계획 모델(310), 정리정돈 추정 모델(330) 및 로봇 추정 모델(360)을 포함하는 것을 예로 한다.
- [0042] 작업 순서 계획 모델(310)은 목표 분할 영상과 현재 분할 영상을 입력받아 다수의 대상 물체(0) 중 정리정돈의 다음 대상을 추정할 수 있다.
- [0043] 본 발명의 실시예에서는 작업 순서 계획 모델(310)이 심층강화학습모델을 기반으로 하여 학습되는 것을 예로 한다. 앞서 설명한 바와 같이, 작업 순서 계획 모델(310)의 학습에는 가상 영상이 영상 분할기(200)에 의해 분할되어 생성된 가상 분할 영상이 학습 데이터로 하여 학습되어 생성된다. 여기서, 가상 영상은, 도 7에 도시된 바와 같이, 각각의 대상 물체(0)에 대응하는 가상 물체가 대상 영역(S)에 대응하는 가상 공간에 분포되어 있다.
- [0044] 본 발명의 실시예에서는 작업 순서 계획 모델(310)에 적용되는 심층강화학습모델로 Rainbow DQN(Deep Q-network) 모델이 적용되는 것을 예로 한다. Rainbow DQN(Deep Q-network) 모델은 목표 분할 영상과 현재 분할 영상을 입력받고, 다음 대상과, 다음 대상의 정리정돈 기본 동작을 추정하여 출력하도록 학습되는 것을 예로 한다.
- [0045] Rainbow DQN(Deep Q-network) 모델은 기본적인 DQN에 Double Q-learning, Prioritized replay buffer, Dueling networks, Noisy Nets, Multi-step learning, 그리고 Distributional Learning의 여섯 가지 기법들이 더해져 총 7가지 요소를 가진 DQN이다.
- [0046] Rainbow DQN(Deep Q-network) 모델의 학습 과정에 대해 설명하면, 목표 분할 영상과 현재 분할 영상을 입력으로 받아, 정리정돈 기본 동작에 대한 가치를 추정한다. 이 때, 추정된 정리정돈 기본 동작에 대한 가치 중 가장 높은 가치를 가지는 작업을 다음 대상 및 정리정돈 기본 동작으로 추정하게 된다.
- [0047] 이후, 가상 환경에서의 가상 로봇이 추정된 작업을 수행하고, 그 결과에 따라 가상 환경으로부터 두 가지 경우의 보상 중 하나를 받게 된다. 첫 번째 보상은 모든 정리를 완료한 경우로 '1'을 보상받고, 두 번째 보상은 모든 정리를 완료하지 못한 경우로 '0'을 보상받게 된다. 즉, 상황에 맞는 적절한 작업을 통해 모든 대상 물체(0)를 정리한 경우에만 보상을 받고, 그 외의 경우에는 보상을 받지 못한다.
- [0048] 이와 같은 보상을 기반으로 평가가 진행되고, 이러한 평가가 학습에 반영됨으로써, 최적의 다음 대상 및 정리정돈 기본 동작을 추정하는 방향으로 학습이 진행된다.
- [0049] 여기서, 대상 물체(0)가 노트북, 텀블러, 서랍, 핸드크림, 상자인 경우를 예로 하여 설명하면, 정리정돈 기본 동작은 "노트북을 닫아라", "텀블러를 정리해라", "서랍을 닫아라", "서랍을 열어라", "상자를 닫아라", "상자를 열어라", "핸드크림을 정리해라" 등을 포함할 수 있다.
- [0050] 또한, "핸드크림을 정리해라"라는 정리정돈 기본 동작은 "핸드크림을 책상 위로 정리해라", "핸드크림을 서랍에 넣어라", "핸드크림을 상자에 넣어라" 등으로 세분화될 수 있다.
- [0051] 본 발명의 실시예에 따른 작업 순서 계획 모델(310)이 정리정돈 기본 동작을 추정하는 경우, 후술할 정리정돈 추정 모델(330)의 정리정돈 상태 추정 모델(331)과 정리정돈 자세 추정 모델(333)이 이에 따른 상태 조작 동작 및 자세 조작 동작을 추정하게 된다.
- [0052] 예를 들어, "핸드크림을 서랍에 넣어라"는 동작을 수행하기 위해, 작업 순서 계획 모델(310)은 "서랍을 열어라" -> "핸드크림을 서랍에 넣어라" -> "서랍을 닫아라"라는 순서로 작업 계획을 수립하게 된다.
- [0053] 한편, 정리정돈 추정 모델(330)은 작업 순서 계획 모델(310)에 의해 추정된 다음 대상에 대한 정리정돈 동작을 추정하게 된다. 즉, 본 발명의 실시예에 따른 정리정돈 추정 모델(330)은 작업 순서 계획 모델(310)에 의해 추정된 다음 대상과, 다음 대상에 대한 정리정돈 기본 동작에 따라 정리정돈 동작을 추정하게 된다.
- [0054] 도 3은 본 발명의 실시예에 따른 정리정돈 시스템(100)의 정리정돈 추정 모델(330)의 제어블록도의 예를 나타낸 도면이다.

- [0055] 도 3을 참조하여 설명하면, 본 발명의 실시예에 따른 정리정돈 추정 모델(330)은, 앞서 설명한 바와 같이, 정리정돈 상태 추정 모델(331)과, 정리정돈 자세 추정 모델(333)을 포함할 수 있다.
- [0056] 본 발명의 실시예에 따른 정리정돈 추정 모델(330)은 다음 대상의 기계적 상태 변화에 따른 상태 조작 동작을 추정한다.
- [0057] 여기서, 다음 대상, 즉 대상 물체(0)의 기계적 상태 변화는 앞서 설명한 예시에서, 서랍을 열거나 닫는 행위, 노트북을 닫는 행위를 포함할 수 있다. 기계적 상태 변화는 비탄성 물체가 기계적인 장치, 즉 본 발명의 실시예에 따른 정리정돈 로봇(20)으로 인해 변환할 수 있는 형태를 의미한다. 다른 예로, 책을 덮는 행위도 이에 포함될 수 있다.
- [0058] 본 발명의 실시예에 따른 정리정돈 상태 추정 모델(331)은 상태 조작 동작으로 다음 대상의 상태 조작 위치, 상태 조작 방법 및 상태 영역을 추정하는 것을 예로 한다.
- [0059] 본 발명에서는 정리정돈 상태 추정 모델(331)이 OUNET(Object understanding network)을 기반으로 상태 조작 동작, 즉 상태 조작 위치, 상태 조작 방법 및 상태 영역을 추정하는 것을 예로 한다.
- [0060] 본 발명의 실시예에서는 정리정돈 상태 추정 모델(331)이, 도 3에 도시된 바와 같이, 대상 추출부(3311), 물체 특성 추정기(3312)를 포함하는 것을 예로 한다. 또한, 본 발명의 실시예에 따른 정리정돈 상태 추정 모델(331)은 DBSCAM(Density-based spatial clustering of application with noise) 모델(3313)을 더 포함할 수 있다.
- [0061] 대상 추출부(3311)는 현재 영상, 즉, 전술한 바와 같이, RGB 영상으로부터 정리정돈 계획 모델에 의해 추정된 다음 대상을 추출한다. 본 발명에서는 정리정돈 계획 모델에 의해 추출된 다음 대상에 대한 좌표 정보를 받아 현재 영상으로부터 다음 대상에 대한 물체 영역을 추출하는 것을 예로 한다.
- [0062] 다른 예로, 대상 추출부(3311)는 인공지능 기반의 물체 인식기(Object detector), 예를 들어, 다양한 물체들이 존재하는 RGB 영상에서 물체들을 검출하는 EfficientDet D7와 같은 물체 인식기로 구현될 수 있다. 그리고, 물체 인식기는 RGB 영상으로부터 추출된 물체 중, 정리정돈 계획 모델로부터 제공되는 다음 대상에 대한 정보에 기초하여, 대상 물체(0)를 추출하도록 마련될 수 있다.
- [0063] 물체 특성 추정기(3312)는 앞서 설명한 바와 같이, OUNET(Object understanding network)을 기반으로 동작하며, 대상 추출부(3311)에 의해 추출된 다음 대상의 상태 조작 위치, 상태 조작 방법 및 상태 영역을 추정한다.
- [0064] 도 8은 본 발명의 실시예에 따른 정리정돈 시스템(100)에서 U-Net 구조의 물체 특성 추정기(3312)의 예를 나타낸 도면이다.
- [0065] 도 8을 참조하여 설명하면, 본 발명의 실시예에 따른 물체 특성 추정기(3312)는 엔코더, 디코더 및 전연결층(Fully connected layer)을 포함할 수 있다.
- [0066] 엔코더는 대상 추출부(3311)에 의해 추출된 다음 대상에 대한 물체 영역을 입력받아 다음 대상의 특징 벡터를 추출한다. 그리고, 디코더는 엔코더와 연결되어 상태 영역 및 상태 조작 위치를 추정하게 된다. 엔코더에 의해 추출된 특징 벡터는 전연결층을 통과하면서, 상태 조작 방법이 분류될 수 있다. 여기서, 본 발명의 실시예에서는 상태 영역 및 상태 조작 위치가 점수 지도(Score map) 형태로 추정되는 것을 예로 한다.
- [0067] 여기서, 상태 영역은 대상 물체(0)의 상태, 예를 들어, 열려있는 상태나 닫혀있는 상태를 특정하기 위한 영역을 의미하며, 도 8에 도시된 바와 같이, 분할 영상 형태로 출력되는 것을 예로 한다.
- [0068] 그리고, 상태 조작 위치는 정리정돈 로봇(20)의 다음 대상의 기계적 상태를 변화시키기 위해 접촉하는 영역을 나타내며, 예를 들어, 서랍의 경우 중간 영역으로 학습되어 추출될 수 있다.
- [0069] 여기서, DBSCAM(Density-based spatial clustering of application with noise) 모델(3313)은 상태 영역 및 상태 조작 위치 각각의 픽셀의 밀도 및 픽셀의 거리를 기반으로 군집화하여 기 설정된 기준 밀도 이하를 상태 영역 및 상태 조작 위치로부터 제거하여 최종적인 상태 영역 및 상태 조작 위치를 추정한다.
- [0070] 이를 통해, 상태 영역이나 상태 조작 위치와 같이 영역의 추출에서 발생하는 오검출을 최소화할 수 있게 된다.
- [0071] 한편, 본 발명의 실시예에 따른 정리정돈 자세 추정 모델(333)은 다음 대상의 자세 변화에 따른 자세 조작 동작을 추정할 수 있다. 여기서, 자세 조작 동작은 다음 대상의 이동 동작과 다음 대상의 회전 동작 중 적어도 하나를 포함할 수 있다.

- [0072] 본 발명의 실시예에서는 정리정돈 자세 추정 모델(333)이 3차원 점군 데이터를 이용하여 자세 조작 동작을 추정한다. 이를 위해, 본 발명의 실시예에 따른 영상 촬영부(100)는 다음 대상을 촬영하기 위한 3차원 점군 촬영부(110)를 포함할 수 있다. 본 발명에서는 영상 촬영부(100)와 3차원 점군 촬영부(110)가 각각 개별적으로 구성되는 것을 예로 하고 있으나, 3차원 점군 촬영부(110)에 의해 촬영된 영상을 현재 영상으로 사용할 수 있음은 물론이다.
- [0073] 자세 조작 동작은 다음 대상의 자세 변화, 즉 이동 동작과 회전 동작을 포함할 수 있으며, 이를 위해서는 대상 물체(0)의 현재 자세가 중요하다.
- [0074] 이를 위해, 정리정돈 자세 추정 모델(333)은, 도 3에 도시된 바와 같이, 점군 분할 모델(3331)과, PCA(Principal component analysis) 모델(3332)을 포함하는 것을 예로 한다.
- [0075] 점군 분할 모델(3331)은 3차원 점군 촬영부(110)에 의해 대상 영역(S)에 대해 복수 방향에서 촬영된 3차원 점군 데이터를 대상 물체(0)별로 분할할 수 있다. 그리고, PCA(Principal component analysis) 모델(3332)은 점군 분할 모델(3331)에 의해 분할된 대상 물체(0) 중 다음 대상에 대한 기준 좌표계를 설정하여 자세 조작 동작을 추정하게 된다.
- [0076] 여기서, 본 발명의 실시예에 따른 점군 분할 모델(3331)은, 도 4에 도시된 바와 같이, 데이터 정합부(3331a), 패스 스루 필터(Pass through filter)(3331b), VCCS(Voxel cloud connectivity segmentation) 모델(3331c), LCCP(Locally convex connected patches) 모델을 포함할 수 있다.
- [0077] 데이터 정합부(3331a)는 다수의 방향에서 촬영된 3차원 점군 데이터를 정리정돈 로봇(20)에 대한 기준 좌표계로 변환하여 정합한다. 어느 일 방향에서만 촬영된 3차원 점군 데이터는 폐색(Occlusion)으로 인해 일부 영역에서 누락이 발생할 수 있는 바, 본 발명에서는 복수의 시점에서 3차원 점군 데이터를 획득하고, 이를 정리정돈 로봇(20)에 대한 기준 좌표계로 변환하여 이를 정합한다.
- [0078] 데이터 정합부(3331a)에 의해 정합된 3차원 점군 데이터는 패스 스루 필터(Pass through filter)(3331b)를 통과하면서 바닥과 같은 배경 점군들이 제거될 수 있다.
- [0079] 그런 다음, VCCS(Voxel cloud connectivity segmentation) 모델(3331c)이 패스 스루 필터를 통과한 3차원 점군 데이터의 각 점군 간의 거리와 색상 차이에 기반하여 3차원 점군 데이터를 복수의 패치(Patch)로 분할하게 된다.
- [0080] 그리고, LCCP(Locally convex connected patches) 모델이 패치들의 곡률, 즉 오목한지 볼록한지 여부를 추정하고, 이를 기반으로 패치들을 합하여 3차원 점군 데이터를 대상 물체(0)별로 군집화하게 된다.
- [0081] 여기서, PCA(Principal component analysis) 모델(3332)은 점군 분할 모델(3331)에 의해 분할된 대상 물체(0) 중 다음 대상에 대한 기준 좌표계를 설정하게 되는데, PCA(Principal component analysis) 모델(3332)은 일정한 규칙을 통해 다음 대상의 기준 좌표계를 설정하게 된다.
- [0082] 예컨대, PCA(Principal component analysis) 모델(3332)에 다음 대상의 3차원 점군 데이터가 입력되면, 점군의 분포에 따라 3개의 주성분을 얻을 수 있다. 이와 같이 얻어진 3개의 주성분 중 큰 순서대로 기준 좌표계의 z축, y축, x축으로 정의함으로써, 다음 대상에 대한 기준 좌표계가 설정될 수 있다.
- [0083] 상기와 같이 다음 대상에 대한 기준 좌표계가 설정되면, 목표 영상과 현재 영상에서 다음 대상의 자세를 얻을 수 있으며, 현재 영상에서의 다음 대상의 자세, 즉 좌표로 부터 목표 영상에서의 다음 대상의 자세, 즉 좌표로의 이동 또는 회전을 위한 자세 조작 동작의 추정이 가능하게 된다.
- [0084] 한편, 본 발명의 실시예에 따른 로봇 추정 모델(360)은, 정리정돈 추정 모델(330)에 의해 추정된 정리정돈 동작에 따라, 정리정돈 로봇(20)이 다음 대상의 정리정돈을 위해 다음 대상에 대해 취해야하는 실제 로봇 동작을 추정한다.
- [0085] 정리정돈 추정 모델(330)에 의해 정리정돈 동작이 추정되더라도, 대상 물체(0)에 따라 실제 조작하는 방법은 다르기 때문에, 이를 수식화하여 학습하기는 현실적으로 불가능하다. 예를 들어, 서랍을 "닫는" 동작과, 노트북을 "닫는" 동작, 상자를 "닫는" 동작이 모두, 닫는 동작이더라도 실제 정리정돈 로봇(20)의 동작은 물체마다 상이하다.
- [0086] 이에, 본 발명에서는 영상 기반의 심층강화학습을 이용하여 로봇 추정 모델(360)을 학습시켜 로봇 동작을 추정하도록 한다. 여기서, 로봇 추정 모델(360)의 학습에 실제 정리정돈 로봇(20)을 이용하여 학습시키는 것은 비효

울적이므로, 본 발명에서는 가상 환경에서 학습을 수행하고, 실제 환경에 적용하는 실제 환경 전이(sim-to-real) 기술을 사용하는 것을 예로 한다.

- [0087] 도 5는 본 발명의 실시예에 따른 로봇 추정 모델(360)의 구성의 예를 나타낸 도면이다. 도 5를 참조하여 설명하면, 본 발명의 실시예에 따른 로봇 추정 모델(360)은 적대적 신경망 모델(361), 영상 압축 모델(363) 및 동작 추정 모델(365)을 포함하는 것을 예로 한다.
- [0088] 적대적 신경망 모델(361)은 현재 영상을 이용하여 표준 영상(Canonical image), 분할 영상(Segmented image), 그리고 깊이 영상(Depth image)을 생성한다. 이는 앞서 설명한 바와 같이, 가상 환경에서의 학습과 실제 환경에서의 적용에 따른 영상 간의 격차를 줄이기 위해서이다.
- [0089] 본 발명에서는 적대적 신경망 모델(361)로 RCAN(Randomized-to-canonical adaptation networks) 모델이 적용되는 것을 예로 한다.
- [0090] RCAN(Randomized-to-canonical adaptation networks) 모델은 물체들의 색이 무작위로 변하는 가상 환경에서 얻은 영상으로부터 물체들의 색이 고정된 또 다른 가상 환경에서의 표준 영상(Canonical image), 분할 영상(Segmented image), 그리고 깊이 영상(Depth image)을 생성한다. 도 9의 (a)는 실제 영상이고, 도 9의 (b)는 표준 영상이고, 도 9의 (c)는 깊이 영상이고, 도 9의 (d)는 분할 영상의 예를 나타내고 있다.
- [0091] 본 발명의 실시예에서는 RCAN(Randomized-to-canonical adaptation networks) 모델이, 도 5에 도시된 바와 같이, 생성기(3611)와 판별기(3612)를 포함하는 것을 예로 한다.
- [0092] 생성기(3611)는 입력되는 학습 데이터, 즉 가상 환경에 대한 가상 영상이 입력되면, 이를 표준 영상, 분할 영상 그리고 깊이 영상으로 변환한다. 그리고, 판별기(3612)는 생성기(3611)에 의해 변환된 표준 영상, 분할 영상 및 깊이 영상을 원본 영상과 구분하는 방법으로 학습이 수행된다.
- [0093] 그리고, 실제 정리정돈 시스템(100)에 적용될 경우, 현재 영상이 입력되면 생성기(3611)가 이를 표준 영상, 분할 영상 및 깊이 영상으로 변환하여 출력하게 된다.
- [0094] 영상 압축 모델(363)은 표준 영상, 분할 영상 및 깊이 영상을 기 설정된 사이즈로 압축한다. 본 발명에서는 영상 압축 모델(363)로 VAE(Variational auto-encoder) 모델이 적용되는 것을 예로 한다.
- [0095] 현재 영상에서 표준 영상, 분할 영상 및 깊이 영상으로 변환되더라도 해당 영상들은 많은 정보를 포함하고 있어, 본 발명에서는 작업의 효율을 높이기 위해 이를 압축하는 것을 예로 한다.
- [0096] 예를 들어, CNN 기반의 인공 신경망인 VAE(Variational auto-encoder) 모델은 $256 \times 256 \times (3(\text{표준 영상}) + 6(\text{분할 영상}) + 1(\text{깊이 영상}))$ 의 영상 정보를 256개의 숫자로 구성된 벡터로 압축할 수 있다. 즉, 세가지 영상을 256개의 숫자로 구성된 벡터로 줄인 후 다시 원래 영상을 복원하는 방법으로 학습되어 생성될 수 있다. 이를 통해, 압축된 256개의 숫자를 통해 현재 상태를 파악함으로써, 용량이 큰 영상 정보보다 효율적으로 학습이 가능하게 된다.
- [0097] 한편, 동작 추정 모델(365)은 영상 압축 모델(363)에 의해 압축된 표준 영상, 분할 영상 및 깊이 영상과, 정리정돈 로봇(20)의 로봇 상태 정보, 정리정돈 동작을 입력받아 로봇 동작을 추정하게 된다. 여기서, 로봇 상태 정보는 정리정돈 로봇(20)의 말단 좌표와 속도를 포함할 수 있다.
- [0098] 본 발명에서는 동작 추정 모델(365)로 SAC(Soft actor critic) 모델이 적용되는 것을 예로 한다.
- [0099] SAC(Soft actor critic) 모델은, 도 5에 도시된 바와 같이, 행동기(3651)와 평가기(3652)를 포함할 수 있다.
- [0100] 행동기(3651)는 영상 압축 모델(363)에 의해 압축된 정보, 그리고 로봇 상태 정보 및 정리정돈 동작을 입력받아 정리정돈 로봇(20)의 다음 동작을 결정하게 된다. 그리고, 평가기(3652)는 가상 환경에서의 로봇 말단의 위치, 속도, 다음 대상의 상태, 정리정돈 동작 그리고 행동기(3651)가 추정한 다음 동작을 입력으로 받아 행동기(3651)가 결정한 행동의 가치를 평가한다.
- [0101] 그리고 행동기(3651)는 평가기(3652)가 평가한 행동의 가치를 이용하여 더 높은 평가를 받도록 학습하고, 평가기(3652)는 주어진 상황과 이때의 로봇의 행동에 대한 정확한 가치를 평가하도록 학습한다. 이 둘을 번갈아 학습시키며 점차적으로 성능을 높일 수 있다.
- [0102] 그리고, 실제 환경에서의 적용에 있어서는 행동기(3651)가 앞서 설명한 정보를 입력받고, 정리정돈 로봇(20)의 다음 동작을 추정하게 된다.

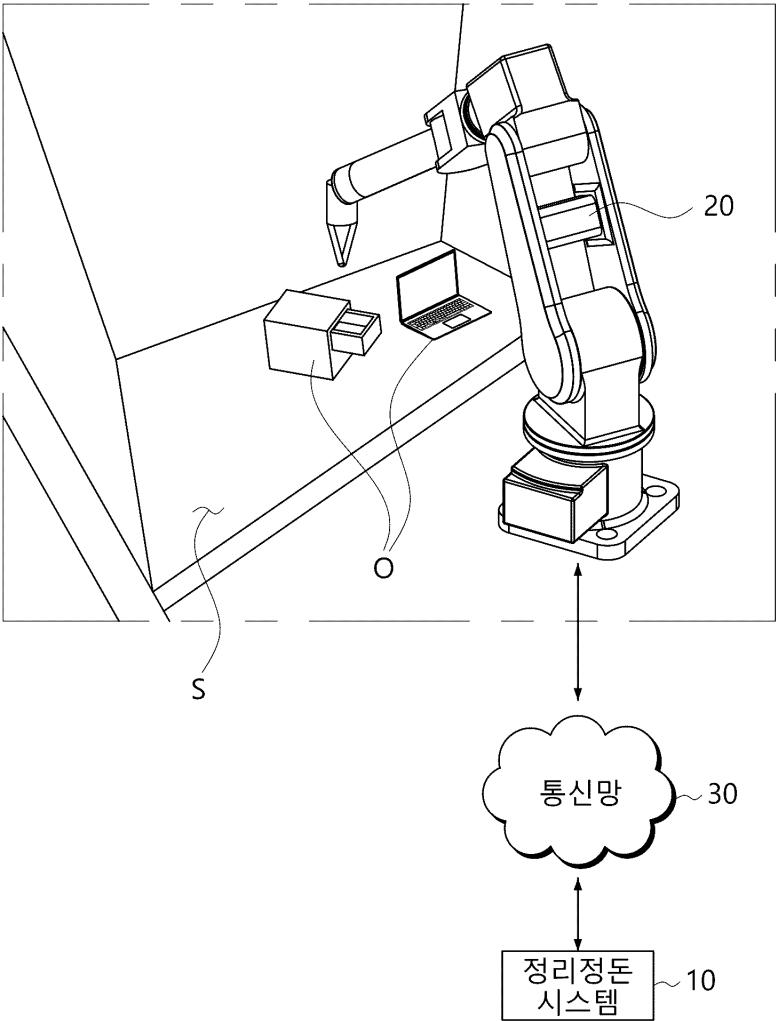
- [0103] 상기와 같은 구성에 따라, 정리정돈 시스템(100)이 인공지능을 기반으로 자체적으로 정리정돈 계획, 즉 다음 대상, 다음 대상의 정리정돈 동작, 그리고, 이에 따른 로봇 동작을 추정하여 실행함으로써, 목표 영상과 현재 영상의 입력 만으로 대상 영역(S)의 정리정돈이 가능하게 된다.
- [0104] 상기와 같이 정리정돈 로봇(20)의 로봇 동작이 추정되면, 메인 제어부(400)는 정리정돈 로봇(20)의 실제 동작을 위한 제어 명령을 정리정돈 로봇(20)으로 전송할 수 있다. 본 발명에서는 도 1에 도시된 바와 같이, 정리정돈 시스템(100)이 통신망(30)을 통해 정리정돈 로봇(20)과 통신하는 것을 예로 한다.
- [0105] 비록 본 발명의 몇몇 실시예들이 도시되고 설명되었지만, 본 발명이 속하는 기술분야의 통상의 지식을 가진 당업자라면 본 발명의 원칙이나 정신에서 벗어나지 않으면서 본 실시예를 변형할 수 있음을 알 수 있을 것이다. 발명의 범위는 첨부된 청구항과 그 균등물에 의해 정해질 것이다.

부호의 설명

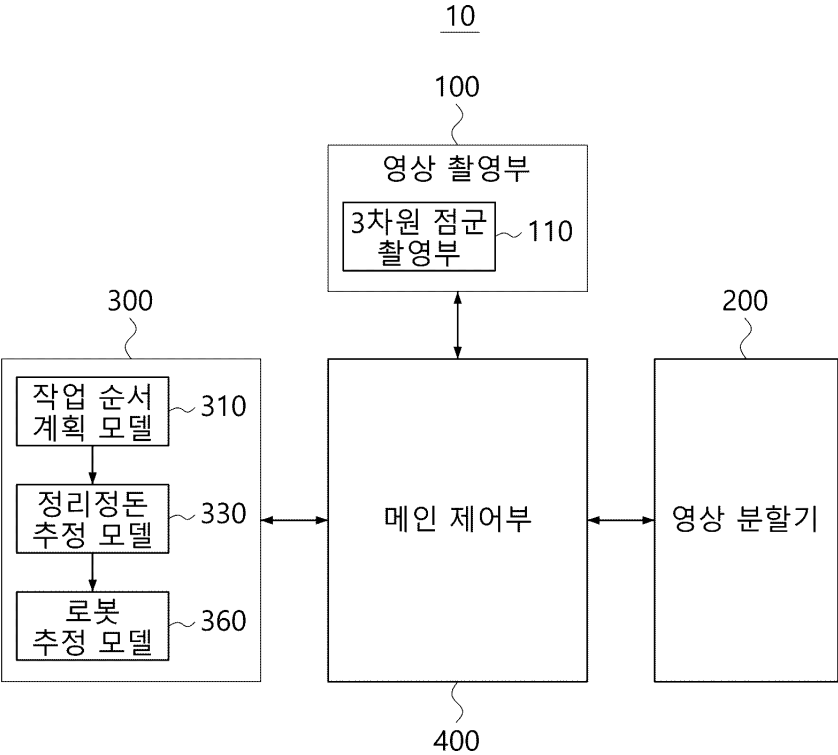
- [0106]
- | | | | |
|-----|---------------|-----|--------------|
| 10 | : 정리정돈 시스템 | 20 | : 정리정돈 로봇 |
| 30 | : 통신망 | | |
| 100 | : 영상 촬영부 | 110 | : 3차원 점군 촬영부 |
| 200 | : 영상 분할기 | 300 | : 정리정돈 계획기 |
| 310 | : 작업 순서 계획 모델 | 330 | : 정리정돈 추정 모델 |
| 360 | : 로봇 추정 모델 | 400 | : 메인 제어부 |

도면

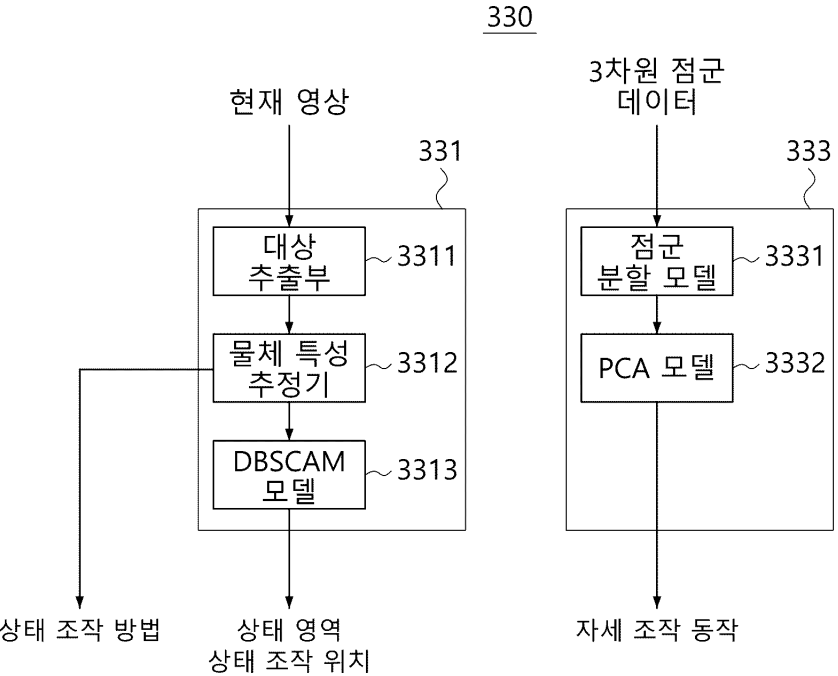
도면1



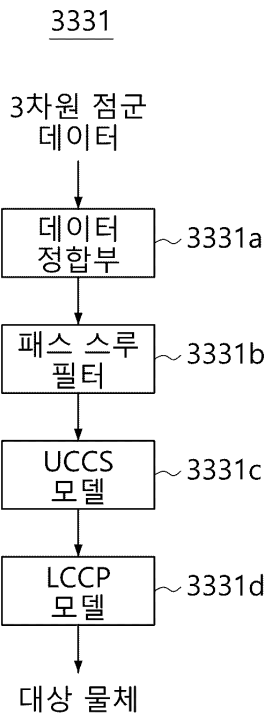
도면2



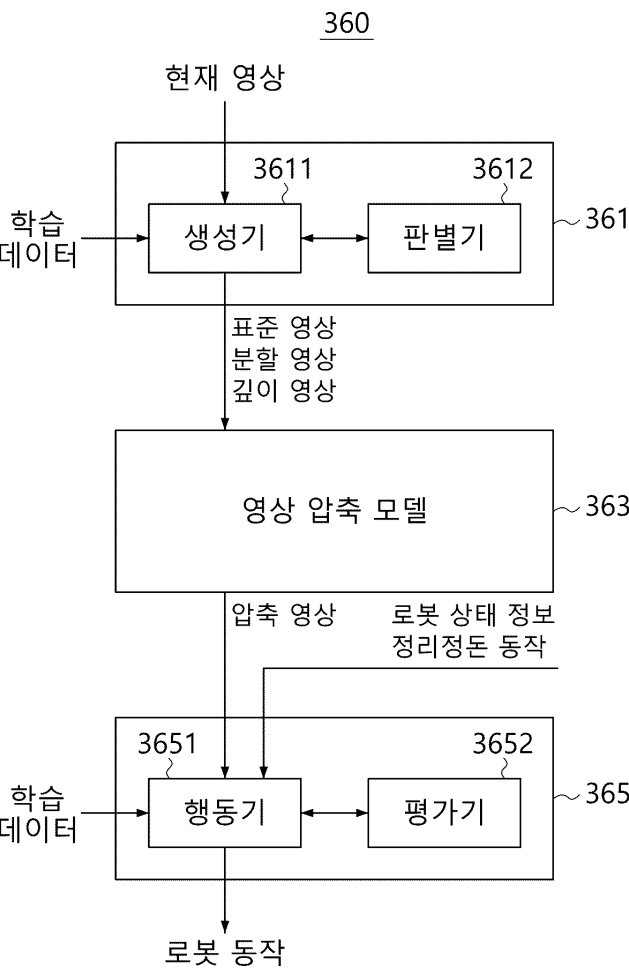
도면3



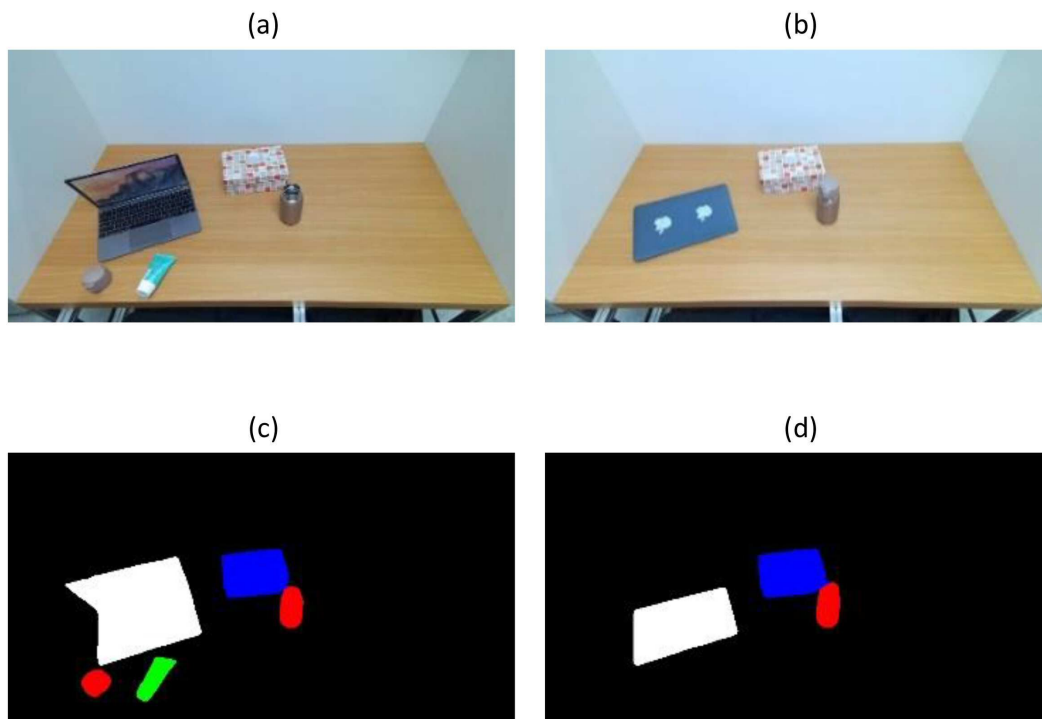
도면4



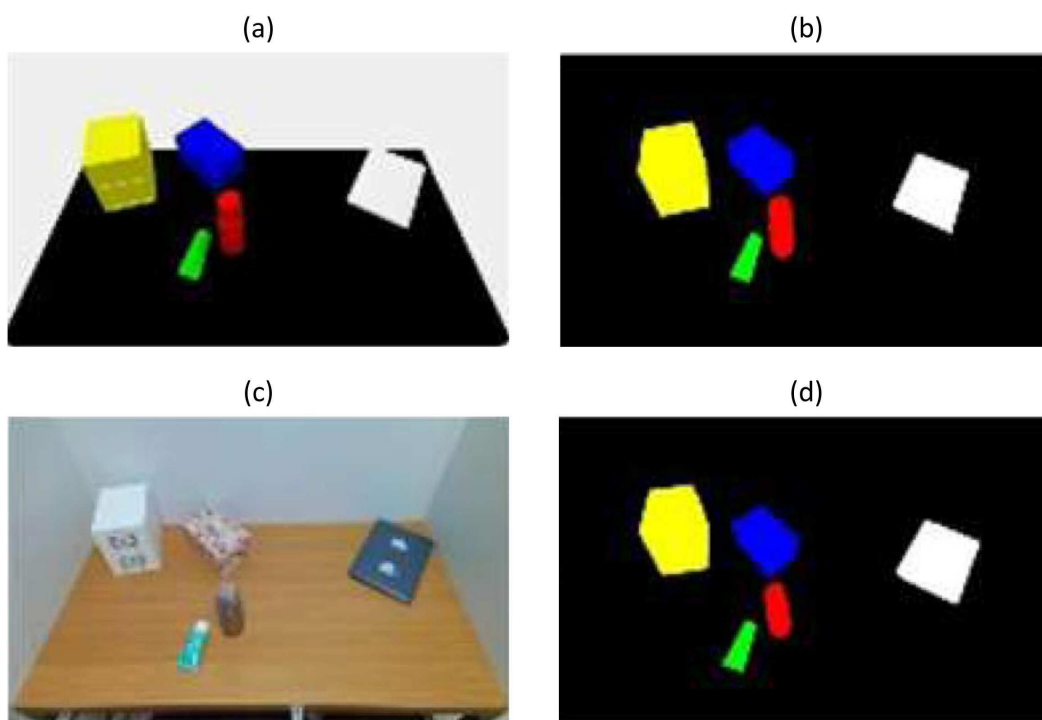
도면5



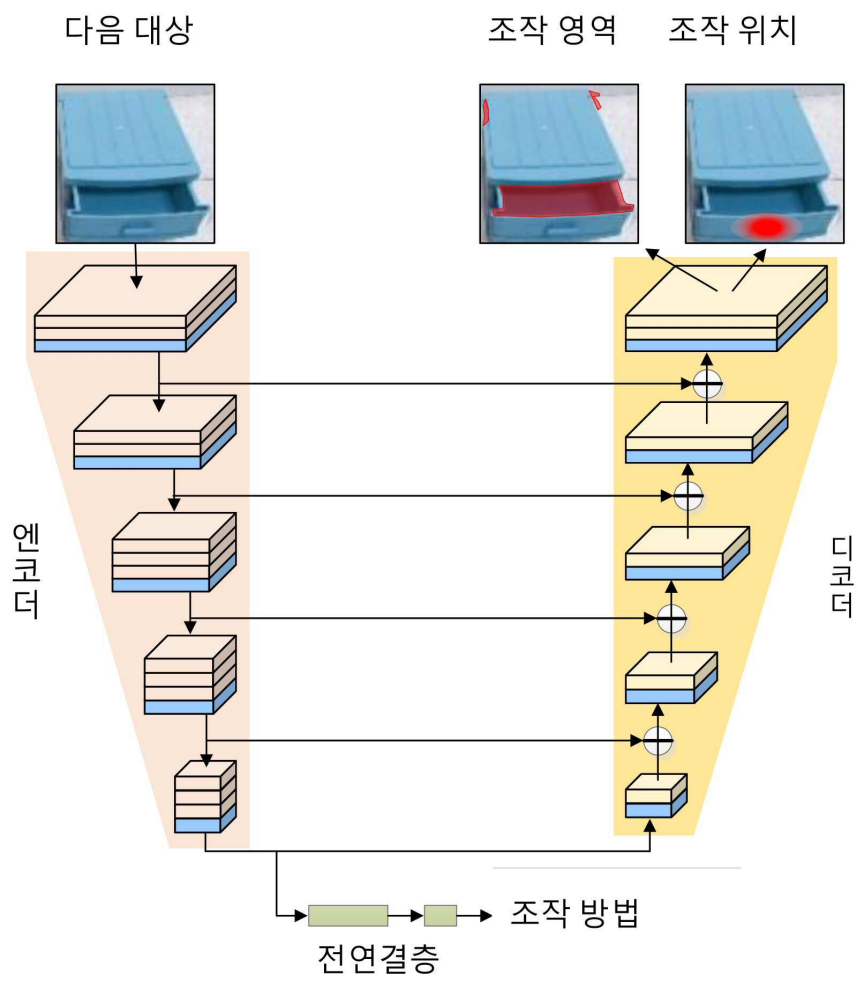
도면6



도면7

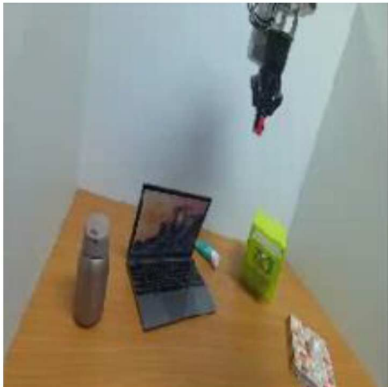


도면8



도면9

(a)



(b)



(c)



(d)

