

다시점 동영상의 계층적 깊이영상 기반 표현, 부호화 및 중간 시점 보간 기법

Representation, Encoding and Intermediate View Interpolation Methods for Multi-view Video Using Layered Depth Images

윤 승 욱 · 호 요 성
Seung-Uk Yoon · Yo-Sung Ho

최근 다양한 형태의 멀티미디어를 이용한 새로운 응용분야와 고품질 콘텐츠가 나타나면서 사용자들은 한층 향상된 실감나는 서비스를 즐길 수 있게 되었다. 그 중에서도 다시점 동영상은 여러 시점에서 다수의 카메라로부터 획득된 정보를 이용하여 사용자에게 원하는 시점의 영상을 제공할 수 있는 장점이 있다. 따라서 차세대 서비스인 삼차원 TV, 자유시점 TV, 실감방송 등에 사용될 중요한 실감형 멀티미디어로 각광을 받고 있다. 그러나 다시점 동영상은 카메라 수에 비례하여 데이터양이 늘어나므로 이를 다양한 응용분야에서 효율적으로 사용하기 위해서는 다시점 동영상의 효과적인 표현, 처리, 복원 기법을 개발하는 것이 필요하다. 본 논문에서는 영상기반 렌더링 기법 중에서 계층적 깊이영상의 개념을 이용하여 다시점 동영상을 효과적으로 표현, 처리, 복원하고 복원된 시점 사이를 보간하는 방법을 제안하였다. 제안한 방법은 다시점 동영상을 삼차원 정보가 포함된 계층적 깊이영상의 개념을 이용해 나타냄으로써 다시점 색상과 깊이정보를 구조적으로 표현한다. 이렇게 표현한 데이터를 데이터 모으기와 계층채우기 기법을 이용해 처리하고 이를 구성요소별로 부호화하여 데이터양을 감소시킨다. 기존의 다시점 동영상 부호화 기법은 원본 시점의 복원을 목적으로 하며 중간 시점 보간을 위해서는 영상 사이의 정합정보를 계산하는 독립적인 알고리즘 및 모듈이 필요하다. 하지만 제안한 방식은 다시점 동영상 데이터를 새로운 표현으로 변환함으로써 별도의 모듈 없이 부호화와 중간 시점 보간 기능을 동시에 제공하는 장점이 있다.

주제어: 다시점 동영상, 영상기반 표현, 계층적 깊이영상, 다시점 동영상 부호화, 시점 보간

The multi-view video is a collection of multiple videos, capturing the same scene at different viewpoints. If we acquire multi-view videos from multiple cameras, it is possible to generate scenes at arbitrary view positions. It means that users can change their viewpoints freely and can feel visible depth with view interaction. Therefore, the multi-view video can be used in a variety of applications including three-dimensional TV (3DTV), free viewpoint TV, and immersive broadcasting. However, since the data size of the multi-view video linearly increases as the number of cameras, it is necessary to develop an effective framework to represent, process, and display multi-view video data. In this paper, we propose a system to represent, encode, and reconstruct multi-view video based on an image-based representation, especially, using the concept of layered depth images (LDI). The proposed framework hierarchically represents various information included in multi-view video based on LDI. In addition, we reduce a large amount of multi-view video data to a manageable size by exploiting an encoding technique, reconstruct original multiple viewpoints, and finally generate intermediate images using a view interpolation method of LDI.

Keywords: Multiview video coding, Image-based Representation, Layered Depth Image, View interpolation

I. 서론

우리는 정보혁명과 디지털 시대로 대변되는 21세기에 살고 있다. 이러한 디지털 시대의 핵심은 여러 가지 형태의 멀티미디어의 등장이라고 할 수 있다. 삼차원 컴퓨터 그래픽스 영상을 비롯해서 파노라마 영상, 삼차원 음향과 같은 멀티미디어의 출현이 삶의 질을 향상시키고 있다. 또한 멀티미디어를 이용한 다양한 응용분야와 고품질 콘텐츠가 나타나게 되었다. 이러한 디지털 시대에 사용자의 시청 실감을 한층 높이기 위해 필요한 기술 중 하나가 삼차원 오디오 비주얼(3DAV: Three-Dimensional Audio Visual) 기술이다. 그 중에서도 다시점 동영상(multi-view video)의 표현, 처리, 압축, 재현은 이의 핵심 요소로서 중요한 역할을 할 것이다.

다시점 동영상은 여러 시점에서 다수의 카메라로 한 장면을 촬영한 다중 동영상의 집합이며 획득된 정보를 이용하여 사용자에게 원하는 시점의 영상을 제공하는 것을 주요 목적으로 한다. 응용분야로는 자유시점 동영상(FVV: Free Viewpoint Video), 자유시점 TV(FTV: Free Viewpoint TV), 삼차원 TV, 실감형 방송, 감시 카메라 영상(surveillance), 파노라믹 영상, 교육, 홈 엔터테인먼트 등이 있다. 그러나 다시점 동영상은 카메라 수에 비례하여 데이터양이 늘어나기 때문에 이를 다양한 응용분야에서 효과적으로 사용하기 위해서는 다시점 동영상의 효율적인 표현, 처리, 부호화, 재현 기법을 개발하는 것이 필요하다.

한편, 영상기반 렌더링(IBR: Image-Based Rendering)은 여러 시점의 이차원 영상을 이용하여 삼차원 공간의 임의 시점에서의 영상을 생성하는 기술이다. 이러한 접근 방식은 이차원 영상을 입력으로 사용하므로 생성하려는 화면의 복잡도와 무관하다. 또한 많은 경우에 원하는 장면을 생성하기 위해 복잡한 삼차원 모델을 만드는 것보다 영상이나 사진을 얻는 것이 더 쉽다. 최근에는 이러한 영상기반 렌더링 기법을 이용하여 동영상으로부터 삼차원 정보를 추출하고, 새로운 시점의 영상을 생성하는 연구가 주목을 받고 있다. 그 중에서도 계층적 깊이영상(LDI: Layered Depth Image) 기법[1]은 여러 시점에서 획득한 색상 및 깊이영상을 합성하여 하나의 데이터 구조로 만들고, 다중 깊이정보와 삼차원 워핑(warping) 함수를 사용하여 임의 시점의 영상을 손쉽게 생성할 수 있는 장점이 있다. 이러한 기능은 다시점 동영상을 이용해 원하는 시점을 제공하려는 목적과 유사하므로, 영상기반 렌더링 기법을 이용하면 다시점 동영상을 표현할 수 있다.

본 논문에서는 영상기반 렌더링 기법 중 하나인 계층적 깊이영상의 개념을 이용하여 실사 다시점 동영상을 효과적으로 표현, 처리 및 부호화하고 이를 이용하여 본래의 다시점 영상 복원 및 중간 시점을 보간하는 시스템을 제안한다.

본 논문의 구성은 다음과 같다. 논문의 II장에서는

다시점 동영상에 관련된 기술의 개발 현황을 소개하고, III장에서는 영상기반 및 동영상 기반 렌더링 기술을 살펴본다. IV장에서는 계층적 깊이영상의 개념을 이용한 다시점 동영상 표현, 부호화, 그리고 다시점 복원기법에 대하여 설명한다. 그리고 V장에서는 제안한 방법을 이용하여 실험한 결과를 보여주고, VI장에서 본 논문의 결론을 맺는다.

II. 다시점 동영상 기술 개발현황

다시점 동영상의 가장 현실성 있는 응용분야 중 하나가 삼차원 TV이므로 관련 기술의 개발이 삼차원 TV 또는 입체 TV의 연구와 같이 이루어지는 경우가 많다. 유럽에서는 Advanced Three-Dimensional Television System Technologies(ATTEST) 과제를 통해 2002년 3월부터 2년 동안 삼차원 TV에 관한 기초기술을 연구하였으며, 2004년 9월부터는 20여개 기관이 컨소시엄을 구성하여 3DTV 과제를 수행하고 있다. 이 과제는 삼차원 장면의 획득, 표현, 부호화, 전송, 디스플레이까지를 모두 포함하고 있으며 다시점 동영상을 획득하는 부분도 주요 연구분야 중의 하나이다. 미국에서는 NASA에서 주로 삼차원 영상과 관련된 연구를 수행하고 있고 MIT에서는 다시점 동영상 및 홀로그래픽 디스플레이 기술을 연구하고 있다. 일본에서는 초다시점 삼차원 TV 기술개발에 많은 투자를 하고 있다. 일본의 경우, 이미 FTV 시험방송을 할 정도로 다시점 동영상에 대한 관심이 매우 크다.

최근에는 다시점 동영상의 부호화에 대한 관심이 높고 이와 관련하여 일본의 나고야 대학, NTT, 독일의 Fraunhofer Heinrich-Hertz-Institut (HHI) 연구소, 미국에 위치한 Mitsubishi Electric Research Laboratories (MERL), Thomson 등에서 연구 개발이 활발하다. 특히 일본의 나고야 대학에서는 100대의 카메라를 이용하여 다시점 동영상을 획득하는 세 가지 종류의 시스템을 개발했다. HHI, MERL, Microsoft Research 에서도 8대의 카메라로부터 다시점 동영상을 획득하는 시스템을 구축하고 이를 이용하여 다시점 동영상 테스트 데이터를 제작하여 배포하고 있다.

국내의 경우, ETRI에서 2002년 월드컵에서 삼차원 TV 시범방송 서비스를 제공하였다. KIST는 다시점 디스플레이 장치를 연구하고 있고 삼성전자와 LG전자에서도 삼차원 카메라, 삼차원 TV, 안경식 스테레오 LCD 모니터를 연구하고 있다. 그 밖에, 광주과학기술원, 강원대, 광운대, 세종대, 연세대 등 대학에서 다시점 영상 획득, 압축, 전송기술을 연구하고 있다. 이와 같이, 국내외 대학, 연구소, 산업체에서 다시점 동영상에 관한 연구를 진행하고 있으며 이에 대한 관심이 매우 높다.

이와 같은 기술적인 흐름에 발맞춰 ISO/IEC JTC1/SC29/WG11 Moving Picture Experts

Group(MPEG)에서도 다시점 동영상 부호화의 필요성을 인정하여, 2001년 12월부터 새로운 3DAV 부호화 기술의 표준화 활동을 준비해 오고 있다. 최근에 이 분야의 연구가 활발히 진행되어 2004년 8월에는 기초적인 성능 실험을 위한 다시점 동영상 테스트 데이터가 제공되었으며, 2004년 10월에 다시점 동영상 부호화에 대한 Call for Evidence(CfE)가 발행되었다[2]. 2005년 7월에 Call for Proposal(CfP)이 배포되었고[3], 2006년 1월에는 CfP에 대한 응답들이 평가되어[4] 다시점 동영상에 대한 기본 부호화 방식과 참조 소프트웨어가 선정되었다. 현재는 주요 부호화 기법들에 대한 Core Experiment(CE)가 진행 중이며[5] 향후 2년 내에 다시점 동영상 부호화의 표준화 작업이 완료될 예정이다.

III. 영상기반 및 동영상기반 렌더링 기법

1. 영상기반 렌더링 기법

영상기반 렌더링은 여러 시점의 이차원 영상을 입력으로 이용하여 삼차원 공간의 임의 시점에서의 영상을 렌더링하는 방법이다. 이러한 접근 방식은 이차원 영상을 입력으로 사용하므로 생성하려는 장면의 복잡도와 무관하다. 또한 많은 경우에 원하는 장면을 생성하기 위해 복잡한 삼차원 모델을 만드는 것보다 영상이나 사진을 얻는 것이 더 쉽다. 이런 이유로 최근에 영상기반 렌더링은 고전적인 모델기반 렌더링의 대체 기술로서 각광을 받고 있다.

영상기반 렌더링 기술은 기하학 정보의 사용 유무에 따라 크게 장면의 기하정보를 사용하지 않는 렌더링 기법, 화소간 대응관계를 이용하는 방법, 그리고 명확한 기하정보를 사용하는 방식으로 분류할 수 있다[6]. 기하정보를 사용하지 않는 렌더링 방법은 플렌옵틱 함수(plenoptic function)의 특성을 이용하여 새로운 시점의 장면을 재현하는 것으로, 빛 정보와 텍스처 정보만을 이용한다. 대표적으로 루미그래프(lumigraph)[7]와 라이트 필드 렌더링(light field rendering)[8]이 속한다.

라이트 필드 렌더링[8]은 조명을 받은 물체 주위의 투명한 공간을 그 물체 표면에서 반사되는 빛으로 채워 실세계와 같은 장면을 연출한다. 고품질의 영상을 렌더링하기 위해 필요한 라이트 필드의 수는 영상의 해상도에 의해 결정된다. 예를 들면, 256 x 256 크기를 갖는 라이트 필드 영상들을 통해 실사와 같은 장면을 재현하려면 약 200,000장의 영상이 필요하다. 따라서 획득된 라이트 필드는 항상 물리적으로 여러 개의 부분 표본화된 레벨로 나누어서 표현한다. 그러나 부분적으로 표본화된 라이트 필드(sub-sampled light field)라 하더라도 고품질의 영상을 재현하기 위해서는 수천 장이 필요하기 때문에 데이터 압축은 라이트 필드 렌더링 응용에서 필수적이다. 라이트 필드 렌더링 기법과 같이 장면

의 기하정보를 이용하지 않는 방식의 장점은 많은 수의 고화질 영상을 획득하면 자연스러운 출력 영상을 얻을 수 있다는 것이지만 이를 위해 엄청난 양의 데이터를 처리해야 하는 단점이 있다. 또한, 최근 라이트 필드의 개념을 동영상으로 확장하려는 시도가 있으나 아직은 연구가 미흡한 상태이다.

화소간 대응관계를 이용하는 렌더링 방법은 새로운 시점의 장면을 재현하기 위해 입력 영상들의 위치적 관련성을 이용하여 기하정보가 직접적으로 렌더링에 이용되지 않는다. 이 방식은 투영행렬(projection matrix) 계산을 통해 얻은 삼차원 위치정보를 새로운 장면 재현을 위해 사용한다. 시점 모핑(view morphing) 기법[9]이 대표적인 예이다.

시점 모핑 기법은 영상 모핑의 개념을 확장한 개념으로 투영 기하학(projective geometry)의 원리를 기반으로 삼차원 투영 카메라와 장면의 변환을 계산하여 최소한의 왜곡으로 자연스러운 변환 장면을 생성하는 방법이다. 기본적인 개념은 우선 두 장의 영상을 워핑(rewarping)한 후, 시점 보간(interpolation)을 통해 중간 시점 혹은 가상 시점의 워핑된 영상을 계산한다. 그 후, 보간으로 생성된 영상에 카메라의 위치와 시점을 고려하여 역으로 워핑(postwarping)을 적용함으로써 가상의 카메라 위치에서 획득한 영상을 생성해 낸다. 시점 모핑은 주로 영상사이의 자연스러운 변화를 만들어 내는 것을 목적으로 하므로 라이트 필드 렌더링에 비해 매우 적은 영상을 사용하고 카메라의 기하 정보와 워핑 및 시점 보간을 사용해 중간 영상을 생성한다. 하지만 동영상이 아닌 정지된 영상에 적용하므로 다시점 동영상을 표현하는 데는 적합하지 않다.

마지막으로 명확한 기하정보가 있는 렌더링 방법은 렌더링시에 직접적으로 깊이정보와 같은 삼차원 정보를 사용하며, 본 논문에서 기술하는 계층적 깊이영상이나 영상기반 모델링을 통한 텍스처 매핑 방법 등이 이 범주에 속한다.

계층적 깊이영상은 복잡한 기하정보를 갖는 삼차원 물체나 장면을 영상기반 렌더링 기법을 이용하여 표현하는 방법 중의 하나이다. 다각형 메쉬를 사용해서 모델을 표현하는 방식과는 달리, 계층적 깊이영상은 여러 시점에서 얻은 다수의 색상 및 깊이영상을 합성하여 하나의 데이터 구조를 생성한다. 따라서 각 계층적 깊이영상 화소는 색상정보 외에 화소와 카메라 사이의 거리를 나타내는 깊이정보와 계층적 깊이영상의 렌더링을 지원하는 추가적인 특성정보를 가지고 있다.

계층적 깊이영상은 다른 시점의 화소정보를 추가적으로 포함하고 있기 때문에, 삼차원 워핑을 수행함으로써 임의의 시점을 생성하는 것이 수월하다. 그림 1은 삼차원 그래픽스 모델로부터 생성된 계층적 깊이영상으로부터 원하는 시점에 따른 영상을 렌더링한 결과를 나타낸다[10]. 그림 1(a)가 삼차원 모델로부터 얻어진 계층적 깊이영상이며 그림 1(b), 1(c), 1(d)는 1(a)를



그림 1. 시점 변화에 따른 계층적 깊이영상 렌더링 결과

이해를 원하는 시점의 영상을 렌더링한 결과이다.

이와 같이 계층적 깊이영상은 다시점 영상들이 가지는 다양한 정보를 한 시점으로 모아 구조적으로 표현할 수 있으며 삼차원 기하정보를 이용하기 때문에 라이트 필드에 비해 상대적으로 적은 수의 영상만으로도 고품질의 최종 영상을 생성해 낼 수 있다.

2. 동영상기반 렌더링 기법

한편, 기존의 영상기반 렌더링 기법은 정지한 장면을 대상으로 적용되었으나 이를 움직이는 동영상으로 확장하려는 시도가 있어왔다. Kanade[11] 등은 원형 돔으로 구성된 공간에 51대의 카메라를 설치하여 움직이는 물체를 촬영하였으며 프레임마다 복셀 컬러링 방식을 이용해 삼차원 물체의 표면을 계산하였다. 그러나 획득한 동영상의 해상도가 낮고 정합 오류와 객체의 윤곽선 부분이 제대로 처리되지 않아 부자연스러운 결과를 나타내었다. 그 이후, Matusik[12] 등은 4대의 카메라로부터 동영상을 획득하여 초당 8 프레임의 속도로 visual hull을 렌더링하였고, Yang[13] 등은 320 x 240의 해상도를 갖는 카메라를 8 x 8 형태의 격자로 배치하여 움직이는 장면을 획득하였다. 이들은 원하는 시점에서 필요한 광선의 정보만을 추출하여 보내는 방식을 사용하여 초당 18 프레임의 속도로 렌더링을 수행하였다. 2004년에 Zitnick[14] 등은 8대의 카메라로부터 획득한 동영상을 이용하여 효과적인 동영상 시점 보간 기법 및 실시간 렌더링 기법을 제안했으며, 이를 이용해 고화질 동영상의 시공간 편집이 가능한 시스템을 개발하였다.

이렇게 동영상기반 렌더링기법을 이용하면 기존에 정지영상에 적용되던 영상기반 렌더링의 장점을 동영상으로 확장할 수 있다. 따라서 본 논문에서는 계층적 깊이영상의 개념을 정지영상이 아닌 동영상으로 확장, 다시점 동영상을 효과적으로 표현 및 부호화하고 원하는 시점을 재현하는 시스템을 제안한다.

IV. 다시점 동영상의 계층적 깊이영상 기반 표현, 부호화 및 중간 시점 보간 기법

1. 계층적 깊이영상의 개념 및 생성

다시점 동영상을 계층적 깊이영상의 형태로 표현하는 구체적 절차를 설명하기 위해 앞서 우선 계층적 깊이영상의 개념 및 생성방법에 대해 살펴보자. 계층적 깊이영상은 다시점에서 획득된 정보를 한 시점(LDI 시점)으로 모은 후, 이를 여러 계층으로 정렬해 표현하는 기법이다. 그림 2는 계층적 깊이영상의 개념을 나타낸다.

광선이 LDI 시점으로부터 물체를 향해 발사된다고 가정하면, 각 화소 위치에 광선과 물체가 교차하는 점의 색상, 기준 시점에서부터 교점까지의 거리, 그리고 렌더링을 위한 부가 정보가 저장된다. 하나의 광선이 물체와 만나는 교점의 수는 물체의 모양에 따라 불규칙적으로 변하므로 각 화소는 서로 다른 수의 교점을 저장한다. 예를 들어 그림 2의 LDI 화소 1은 두 개의 교점 즉, 두 개의 계층을 갖고, 화소 2는 세 개의 계층, 화소 3은 네 개의 계층을 갖는 식이다.

일반적으로 컴퓨터 그래픽스 객체로부터 계층적 깊이영상을 생성하는 경우, 객체를 바라보는 가상의 카메라를 지정하고 광선과 객체의 교점을 계산할 수 있다. 하지만 실세계에서는 광선이 물체를 투과할 수 없으므로, 앞서 설명한 개념적인 방법으로 계층적 깊이영상을 생성하는 것이 불가능하다. 따라서 이 경우에는 다시점 실사영상 및 깊이정보가 필요하며 삼차원 워핑을 통해 다시점 데이터를 한 시점으로 이동시킨 후 이를 계층적으로 정렬하는 방식을 사용한다[1].

2. 삼차원 워핑

실사 다시점 동영상을 계층적 깊이영상으로 표현하기 위한 가장 중요한 구성요소는 삼차원 워핑이다. 일반적으로 이차원 영상 워핑은 매핑을 통해 원본 영상을 원하는 영상으로 왜곡시키는 기법이라고 할 수 있다. 영

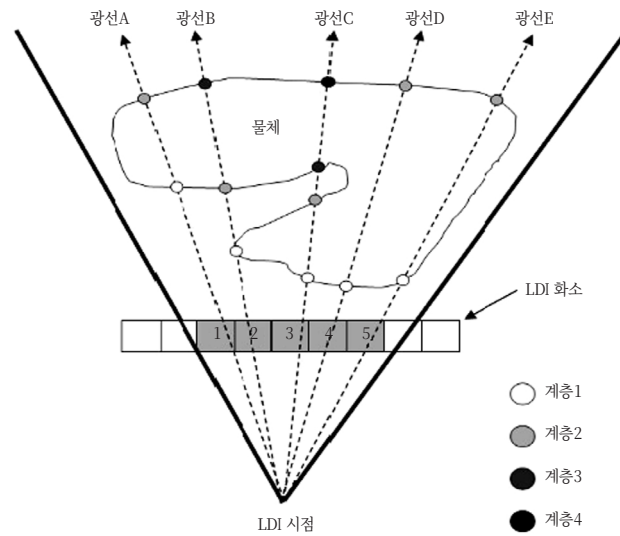


그림 2. 계층적 깊이영상의 개념

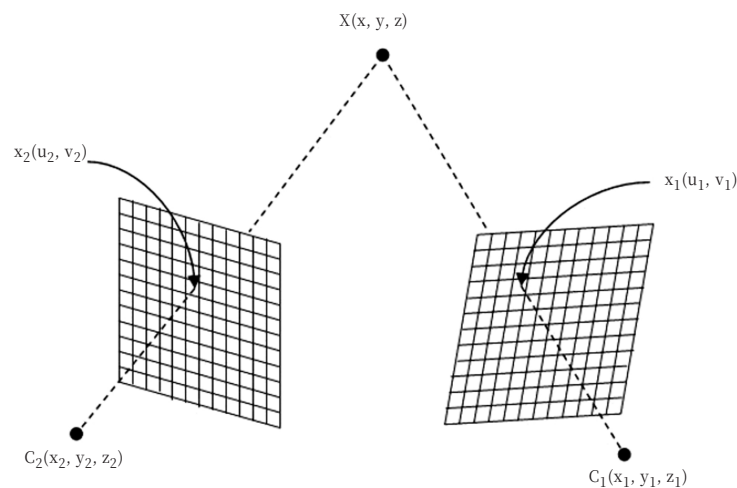


그림 3. 삼차원 워핑의 예

상 워핑은 주로 영상 획득 시스템에서 발생하는 영상의 기하학적인 왜곡을 보정하기 위한 목적으로 많이 사용되어왔다. 컴퓨터 그래픽스에서는 이를 이용해 의도적으로 영상에 왜곡을 줌으로써 예술적 효과나 특수효과를 생성하는데 사용하기도 한다. 워핑을 위해서는 원본과 목적으로 하는 영상사이에 왜곡을 생성할 수 있는 매핑을 정의해야 하며, 화소간 대응관계를 계산하여 이를 정의할 수 있다.

삼차원 워핑은 한 시점에서 획득한 영상을 다른 임의의 시점으로 이동할 때 사용하며 두 영상의 화소간 대응관계 대신 카메라 매개변수와 화소당 깊이정보를 이

용한다. 그림 3은 삼차원 워핑의 예를 나타낸다. 그림에서 삼차원 공간의 기저 벡터는 $\vec{x}$, $\vec{y}$, $\vec{z}$ 로, 영상평면의 기저 벡터는 $\vec{u}$, $\vec{v}$ 로 가정하였다. 삼차원 공간상의 한 점 C 가 카메라의 투영중심(center of projection)일 때, $C_1(x_1, y_1, z_1)$ 시점에서 본 영상평면의 한 화소 $x_1(u_1, v_1)$ 은 삼차원 공간상의 점 $X(x, y, z)$ 에 대응된다. 이 점을 $C_2(x_2, y_2, z_2)$ 시점으로 투영하게 되면 $C_2(x_2, y_2, z_2)$ 위치에서 본 영상평면의 한 점 $x_2(u_2, v_2)$ 를 얻는다. 이 과정을 모든 화소에 대해 반복하면 새로운 시점의 영상을 획득할 수 있다[15].

이를 수식으로 표현하기 위해 우선 삼차원 월드 좌

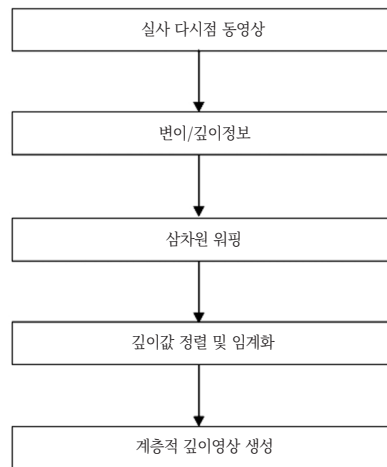


그림 4. 다시점 동영상의 계층적 깊이영상 표현 과정

표계의 한 점을 카메라에 투영된 이차원 평면의 한 점으로 변환하는 행렬 C_1, C_2 를 정의한다.

$$C_1 = V_1 \cdot P_1 \cdot A_1, C_2 = V_2 \cdot P_2 \cdot A_2 \quad (1)$$

여기서 V 는 뷰포트 (viewport) 행렬, P 는 투영 행렬, A 는 어파인(affine) 행렬을 나타낸다. $x_1(u_1, v_1)$ 을 $x_2(u_2, v_2)$ 로 이동하는 변환행렬을 $T_{1,2} = C_2 \cdot C_1^{-1}$ 이라 하면, 다음 식에 의해 카메라 시점 $C_1(x_1, y_1, z_1)$ 에서 본 영상을 카메라 시점 $C_2(x_2, y_2, z_2)$ 로 워핑할 수 있다.

$$T_{1,2} \cdot \begin{bmatrix} x_1 \\ y_1 \\ z_1 \\ 1 \end{bmatrix} = \begin{bmatrix} x_2 \cdot w_2 \\ y_2 \cdot w_2 \\ z_2 \cdot w_2 \\ w_2 \end{bmatrix} \quad (2)$$

z_1 은 화소의 깊이값이며 새로운 시점의 영상 좌표 (x_2, y_2) 는 결과를 w_2 로 나눈 후 얻게 된다[1].

3. 실사 다시점 동영상의 계층적 깊이영상 표현

컴퓨터 그래픽스 객체가 아닌 실사 다시점 영상을 계층적 깊이영상으로 표현하기 위해서는 삼차원 워핑을 통해 다시점 영상을 하나의 기준시점(LDI 시점)으로 이동시켜야 한다. 이 과정은 영상의 시점이 바뀌는 과정이므로, 폐색(occlusion)과 비폐색(disocclusion)이 발생하게 되며 이를 효과적으로 처리하는 것이 중요하다.

폐색으로 인해 여러 화소가 한 화소로 겹치는 문제

는 이동된 화소들을 깊이값에 따라 정렬하여 보는 시점에 가장 가까운 화소를 표현함으로써 해결할 수 있다. 컴퓨터 그래픽스 분야에서는 화소의 깊이값 또는 변이값(disparity)을 이용한 Z-buffering 기법 등을 사용하여 이 문제의 해결책을 모색해 왔다. 비폐색으로 인해 이전 영상에서는 보이지 않던 새로운 부분이 생겨나는 경우에는 공백(hole)이 발생한다. 공백을 채우기 위해 주변 화소값을 이용한 보간 방법이 많이 사용되어 왔다.

이들 방식과는 달리 계층적 깊이영상 표현에서는 기준 시점의 각 화소 위치에 다른 시점에서 가져온 화소들을 여러 계층으로 저장한다. 만약 두 시점에서 기준 시점으로 이동된 화소의 깊이값의 차이가 미리 정해진 임계치보다 작다면 평균을 취해 이를 하나의 화소로 저장한다. 반대로 깊이값의 차이가 임계치보다 크면, 새로운 계층을 생성해 화소를 저장한다. 따라서 8시점에서 획득된 다시점 영상이 있을 경우, 최대 8개의 계층이 생성될 수 있다. 이 방식의 장점은 공백이 발생했을 때 후위 계층의 화소값을 이용해 이를 채울 수 있다는 점이다. 후위 계층의 화소는 다른 시점에서 이동되어 왔으므로 비폐색으로 발생한 영역에 대한 정보를 가지고 있다. 따라서 단순히 주변값을 이용하는 것보다 효과적으로 이 문제를 해결할 수 있다.

이렇게 워핑과 깊이값에 따른 화소의 정렬 그리고 계층생성 과정을 거쳐, 각 프레임마다 획득된 실사 다시점 영상들로부터 계층적 깊이영상을 생성한다. 그러면 최종적으로 다시점 동영상에 대한 계층적 깊이영상 시퀀스를 얻을 수 있다. 다시점 동영상의 계층적 깊이영상 표현 과정을 그림 4에 나타내었다.

본 논문에서는 Microsoft Research(MSR)에서 배포한 다시점 동영상을 사용하여 실험을 수행하였다. 현재 다양한 기관에서 다시점 동영상 테스트 데이터를 제작하여 배포하고 있지만 MSR에서 제공하는 데이터만

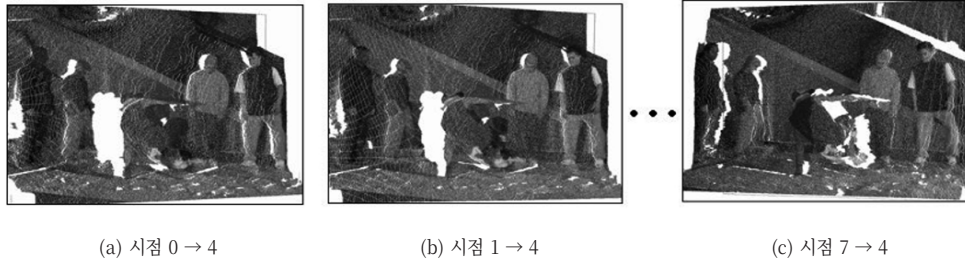


그림 5. 기준 카메라 시점으로 삼차원 워핑 수행 결과

이 다시점 깊이정보를 포함하고 있다. 깊이정보가 없는 다시점 동영상에 대해서는 스테레오 알고리즘을 사용하여 변이정보를 예측할 수 있다. 그러나 알고리즘의 종류와 성능이 매우 다양하여 적절한 기법의 선택에 어려움이 있다고 판단하여 여기서는 제공되는 데이터를 이용하였다. MSR 데이터는 일차원 원호형(arc)으로 배치된 8대의 카메라로부터 획득한 다시점 동영상이며, 각 카메라마다 색상 및 깊이동영상과 카메라 매개변수가 제공된다[14],[16].

기존의 계층적 깊이영상은 주로 삼차원 컴퓨터 그래픽스 모델로부터 생성되었기 때문에 실사 영상에 대해 식 (1)의 행렬을 계산하기가 어렵다. 따라서 본 논문에서는 카메라 행렬을 주어진 카메라 매개변수로부터 새롭게 계산하였다. 3×4 투영행렬이 제공되므로 워핑에 사용될 카메라 행렬 C 는 다음 식에 의해 계산된다[17].

$$C' = A \cdot E = \begin{bmatrix} c'_{11} & c'_{12} & c'_{13} & c'_{14} \\ c'_{21} & c'_{22} & c'_{23} & c'_{24} \\ c'_{31} & c'_{32} & c'_{33} & c'_{34} \end{bmatrix} \quad (3)$$

$$A = \begin{bmatrix} \alpha_x & s & x_0 \\ 0 & \alpha_y & y_0 \\ 0 & 0 & 1 \end{bmatrix}, E = \begin{bmatrix} R_{11} & R_{12} & R_{13} & T_1 \\ R_{21} & R_{22} & R_{23} & T_2 \\ R_{31} & R_{32} & R_{33} & T_3 \end{bmatrix} \quad (4)$$

$$C = \begin{bmatrix} c_{11} & c_{12} & c_{13} & c_{14} \\ c_{21} & c_{22} & c_{23} & c_{24} \\ c_{31} & c_{32} & c_{33} & c_{34} \\ 0 & 0 & 0 & 1 \end{bmatrix} \quad (5)$$

여기서 A 는 카메라의 내부 매개변수(intrinsic parameter)를 나타내는 행렬이며, α_x , α_y 는 각각 x , y 축방향의 화소단위 초점거리, s 는 카메라 촬상면(CCD)의 두 축 사이의 기울어짐(skew) 정도, x_0 , y_0 는 화소

단위 주점(principle point)의 위치이다. 일반적으로 (x_0, y_0) 는 획득되는 영상의 중심점에 근접한다. E 는 외부 매개변수(extrinsic parameter)를 나타내는 행렬로 구성요소인 R 은 회전, T 는 이동행렬을 나타낸다. C 는 4×4 행렬이므로 C' 에 마지막 행 $[0 \ 0 \ 0 \ 1]$ 을 추가하여 생성한다[17]. 테스트 데이터는 모든 카메라에 대한 매개변수와 깊이값을 제공하므로 이를 이용해 C_0 , C_1 , ..., C_7 를 계산하고 식 (2)를 통해 최종 워핑을 수행한다.

그림 5는 MSR의 다시점 동영상 테스트 데이터 중 “Breakdancers” 시퀀스를 이용하여 다시점 영상을 기준 시점(카메라 4)으로 워핑한 결과를 보여준다. 그림에서 하얀 화소로 처리된 부분은 시점이 이동하여 새롭게 드러난 부분이며 삼차원 워핑의 결과를 명확히 나타내기 위해 보간을 하지 않고 하얀 화소로 표현하였다[17].

이렇게 기준 시점 위치로 모든 시점의 영상을 워핑한 후에는 이동된 화소들을 정렬하고 깊이값을 비교하여 계층적 깊이영상을 생성하게 된다. 앞서 언급한 바와 같이, 최대 생성될 수 있는 계층의 수는 다시점 동영상의 최대 시점의 수와 동일하다. 깊이값 비교시 임계값(threshold) 5.0을 사용하였을 때 생성된 계층적 깊이영상의 계층별 색상영상을 그림 5에 나타내었다.

그림 6에서 관찰할 수 있는 계층적 깊이영상의 중요한 특징은 화소별로 다른 개수의 계층이 존재하며 후위 계층으로 갈수록 화소의 밀도가 줄어드는 것이다. 후자의 주요 원인은 다시점 데이터를 획득할 때 주로 카메라 시스템이 물체의 정면에 위치하기 때문이다. 현재 배포된 다시점 동영상 테스트 데이터의 경우 대부분 작은 크기의 물체가 아닌 복잡한 실사 장면을 담고 있다. 따라서 카메라의 공간적인 배치와 제한된 카메라 개수 등을 고려할 때 밀집된 카메라들을 장면의 전면에 배치하는 것이 일반적인 추세이다.

4. 다시점 동영상의 계층적 깊이영상 기반 부호화

다시점 동영상을 계층적 깊이영상의 형태로 표현하면 임계화 과정에서 깊이값의 차이가 적은 화소들이 평균값을 가진 화소로 합쳐지므로 상당한 양의 데이터를 줄일

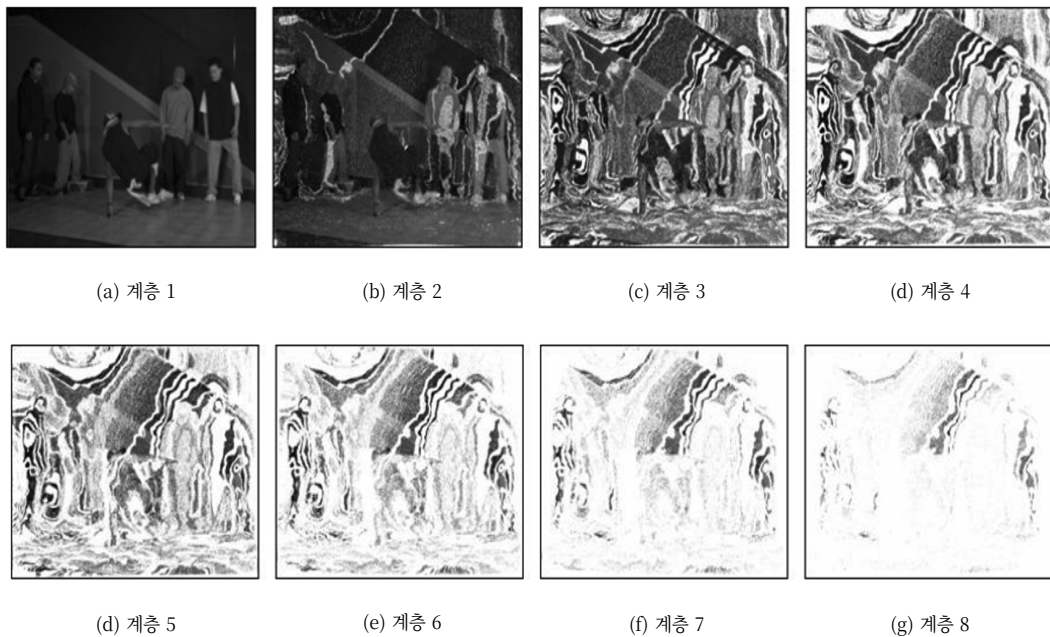


그림 6. 생성된 계층적 깊이영상의 계층별 색상영상

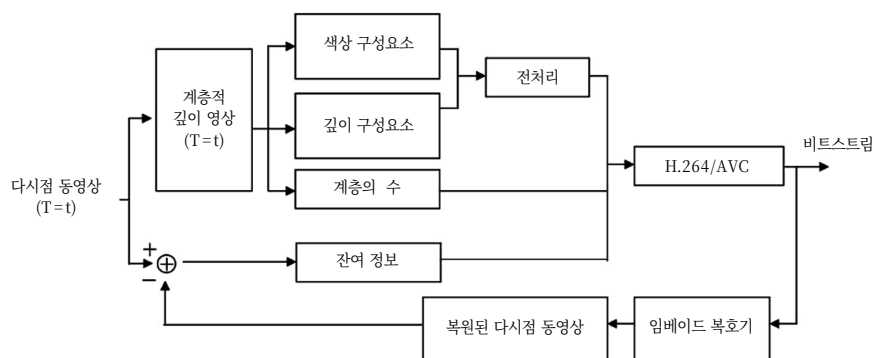


그림 7. 계층적 깊이영상 기반 다시점 동영상 부호화기

수 있다. 즉, 각 카메라에서 획득된 영상사이에 존재하는 공간적인 상관관계가 계층적 깊이영상의 계층별로 재배치되므로 이를 이용해 공간적 중복 데이터를 줄일 수 있다. 기존의 논문에서는 계층적 깊이영상의 개념을 이용해 다시점 동영상을 표현 및 처리하는 프레임워크를 제안하였으며[18], 본 논문에서는 구체적으로 다시점 동영상으로부터 생성된 계층적 깊이영상의 전처리 및 부호화 방법을 제안한다. 그림 7은 제안한 계층적 깊이영상 기반 다시점 동영상 부호화기의 구조이다[19].

계층적 깊이영상의 각 계층영상은 화소의 분포가 일정하지 않고, 후위 계층으로 갈수록 밀도가 줄어들게 되

는 특징이 있다. 따라서 기존의 계층적 깊이영상 압축방식은 데이터 모으기(data aggregation)를 통해 분산된 화소를 한 방향으로 모으고 여기에 임의의 형상 기반 웨이블릿 부호화 기법을 적용하였다[20]. 그러나 부호화 효율 측면에서 웨이블릿 기반 코덱보다 H.264/AVC의 성능이 우수한 것이 입증되고 있기 때문에 본 논문에서는 H.264/AVC 기반으로 부호화를 수행하였다. 하지만 H.264/AVC는 임의의 형상 부호화를 지원하지 않으므로 계층별 영상을 부호화가 가능한 형태로 변환해 주는 전처리 과정이 필요하다.

본 논문에서는 전처리 기법으로 데이터 모으기와 계

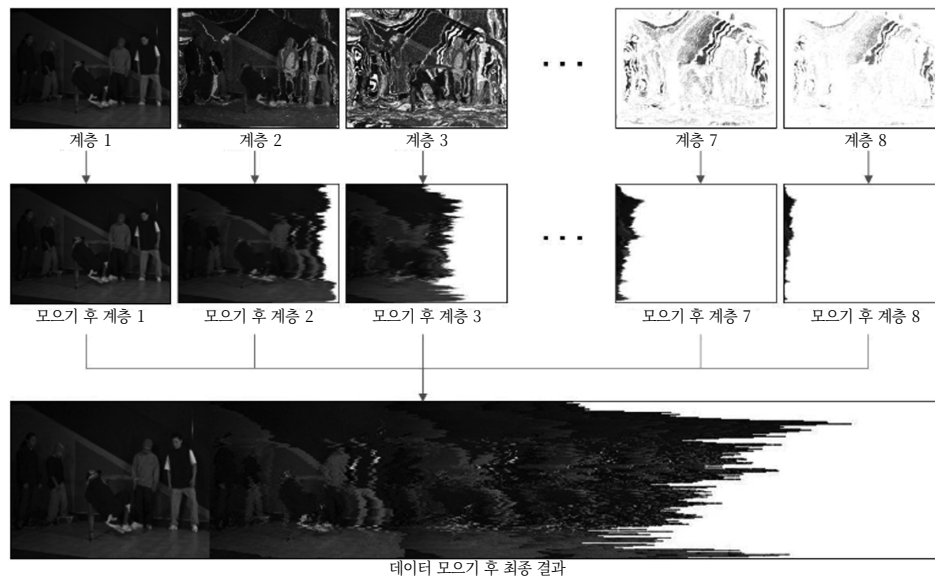


그림 8. 수평방향 데이터 모으기 결과

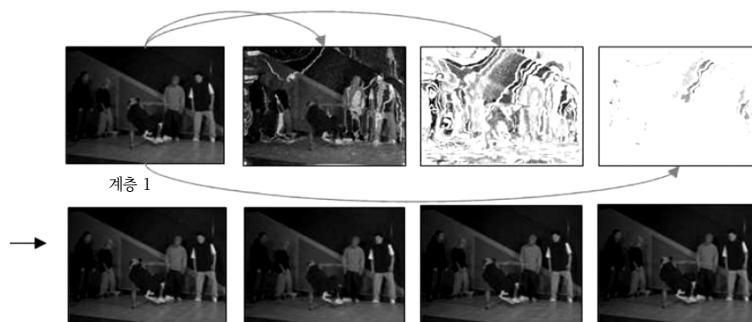


그림 9. 계층 채우기 기법을 사용한 결과

층 채우기 (layer filling) 방법을 사용한다[19]. 기존의 방식에서는 각 계층별로 데이터 모으기를 수행한 후, 각 계층에 독립적으로 임의형상 부호화를 적용한다. 이와는 달리 제안한 방식은 계층별로 모으기를 적용하고, 계층영상 전체를 하나의 영상으로 결합한다. 여기에 한 번 더 데이터 모으기를 적용한 후, 남은 부분은 각 스캔라인의 마지막 화소값으로 패딩한다. 마지막으로 패딩된 전체 영상을 원본 영상 크기에 맞게 분할한 후, 이를 최종 시퀀스 형태로 만들어 H.264/AVC를 적용한다. 데이터 모으기를 수행하는 과정을 그림 8에 나타내었다.

한편, 계층 채우기 기법은 계층적 깊이영상의 첫 번째 계층에 있는 화소값으로 후위 계층의 빈 화소 위치를 채우는 방식이다. 이는 계층적 깊이영상에서 앞선 계층의 화소가 존재할 때만 후위 계층의 화소가 존재하는 특

성을 이용하여 계층별 영상 사이에 상관도를 높이는 방법이다. 그림 9는 계층 채우기 과정을 나타낸다.

계층 채우기가 가능한 이유는 계층적 깊이영상이 각 화소위치에 존재하는 계층의 수(Number of Layer, NOL)를 저장하기 때문이다. NOL은 화소당 1에서 최대 계층의 수 (테스트 데이터에서는 8) 만큼의 숫자를 저장하는 한 장의 영상으로 생각할 수 있다. NOL 정보를 이용해 계층별 화소의 위치를 찾을 수 있기 때문에, NOL은 잔여 정보와 함께 무손실 부호화된다.

5. 계층적 깊이영상의 개념을 이용한 다시점 동영상의 복원 및 중간 시점 보간

계층적 깊이영상 표현으로부터 원래의 다시점 동영상

표 1. "Breakdancers" 데이터 크기 비교

[단위: kB]

"Breakdancers"	첫 번째 8 프레임	두 번째 8 프레임
8시점 영상의 합계	25,166.0	25,166.0
독립적 부호화(색상+깊이)	137.7	132.5
독립적 부호화(색상)	97.4	96.3
LDI 프레임(임계값: 3.0)	13,924.0	13,803.0
부호화된 LDI(데이터 모으기)	131.7	133.8
부호화된 LDI(계층 채우기)	48.4	48.2
LDI 프레임(임계값: 5.0)	13,808.0	13,723.0
부호화된 LDI(데이터 모으기)	91.7	93.0
부호화된 LDI(계층 채우기)	46.3	47.0

상을 재현하고 새로운 시점을 생성하는 과정은 앞서 기술한 계층적 깊이영상 생성 과정의 역으로 생각할 수 있다. 우선 복호된 계층적 깊이영상 프레임을 각 카메라 시점으로 역워핑한다. 이 때, 이미 생성단계에서 존재 하던 폐색/비폐색 영역과 데이터 손실로 인해 역워핑된 결과만으로는 복원이 불가능하다. 따라서 앞서 설명한 바와 같이 후위 계층에 포함된 다시점 색상 및 깊이정보를 사용하여 현재 프레임의 빈 화소 위치를 최대한 보상한다. 하지만 카메라 0번과 7번의 좌우측 끝부분 데이터와 같이 다시점 데이터 획득시 카메라의 제한된 시야 각과 위치에서 야기되는 비폐색영역은 여전히 복원이 어렵다. 따라서 이 부분에 대해서는 원본 다시점 동영상으로부터 보상에 필요한 정보를 추가로 보내주어야 한다. 최종적으로 역워핑된 결과에 후위 계층에 포함된 다시점 색상 및 깊이정보와 원본으로부터 추출된 잔여 정보를 추가하면 원래의 다시점을 복원할 수 있다[19].

한편, 원래의 시점이 아닌 중간 시점 혹은 사용자가 원하는 가상의 시점을 복원하기 위해 본 논문에서는 카메라 매개변수를 보간하는 방법을 사용하였다. 카메라 행렬인 C 에 대해서는 시점에 따라 카메라 외부 매개변수 행렬에 가중치를 두고 선형적으로 보간하여 새로운 카메라 행렬을 계산하였다. 예를 들어, i 와 j 번째 ($0 \leq i, j < \text{카메라 개수}$) 카메라의 행렬을 C_i, C_j 라고 할 때, 다음 식을 사용하여 카메라 매개변수 행렬을 선형 보간한다. 식에서 α ($0 \leq \alpha \leq 1$)는 선택한 시점이 어느 카메라에 가까운지를 나타내는 가중치이다.

$$C_{ij} = \alpha \cdot C_i + (1 - \alpha) \cdot C_j \quad (6)$$

회전성분을 제외한 매개변수들은 식(6)에 의해 보간 하며, 회전성분에 대한 보간은 다음과 같다. 두 평면의 단위 법선 벡터를 계산한 후, 이 벡터들에 수직인 축으로 두 벡터가 이루는 각도만큼 회전이동을 수행한다.

$$N_{ij} = N_i \times N_j, R(\theta) = \cos^{-1}(N_i \cdot N_j) \quad (7)$$

카메라 i 와 j ($0 \leq i, j < \text{카메라 개수}$)에서 바라본 두 영상평면의 단위 법선 벡터를 N_i, N_j 라고 할 때, N_{ij} 은 두 벡터의 외적으로 회전할 축을 나타낸다. 회전각 $R(\theta)$ 는 두 벡터의 사이 각으로 계산할 수 있다.

V. 실험 결과 및 분석

1. 다시점 동영상의 계층적 깊이영상 기반 부호화

본 논문에서는 생성되는 계층적 깊이영상 시퀀스의 데이터양을 줄임으로써 복호화 성능을 향상시키고, 렌더링 속도를 높이기 위해 계층적 깊이영상을 부호화하였다. 데이터 모으기와 계층 채우기 기법을 이용한 전처리를 수행한 후, H.264/AVC를 이용해 부호화를 수행하였다.

표 1은 원본 다시점 동영상과 계층적 깊이영상 표현을 부호화 후 데이터의 크기를 비교한 결과를 나타낸다. 부호화에는 H.264/AVC JM 9.5가 사용되었으며, 데이터 크기는 kB로 나타내었다. 표 1에서 8시점 영상의 합계는 BMP 형식의 원본 데이터양을 나타낸다. 각 시점마다 색상과 깊이영상이 존재하므로 총 16장의 영상 크기의 합계가 된다. 독립적 부호화는 각 동영상을 H.264/AVC 기반으로 각각 부호화하여 합계를 계산한 것을 말한다. 색상과 깊이영상 프레임 각각에 대해 인트라 프레임 부호화를 수행했으며 이 때 사용된 탐색범위는 32, QP는 30이다. 계층적 깊이영상 프레임은 이 다시점 영상을 계층적 깊이영상으로 표현하였을 경우 데이터양을 나타낸다[19].

실험 결과 두 전처리 방식 중 계층 채우기 방식의 성능이 우수한 것을 확인할 수 있다. 이는 데이터 모으기 방식이 계층별로 화소를 모으면서 공간적인 유사성을 감소시킨 반면, 계층 채우기는 최대한 계층간 유사성을



그림 10. 각 시점으로 역 삼차원 위핑을 수행한 결과



그림 11. 후위 계층 정보를 이용한 다시점 복원 결과

유지하였기 때문이다.

2. 다시점 동영상 복원 및 중간 시점 보간

복호된 계층적 깊이영상으로부터 각 시점으로 역으로 삼차원 위핑을 수행한 결과를 그림 10에 나타내었다 [19]. 가운데 위치한 영상이 기준 카메라로 삼차원 위핑을 수행한 결과이다. 위쪽의 영상들은 각각 카메라 0과 1번 위치로, 아래쪽 영상들은 카메라 6과 7번 위치로 역위핑한 결과이다.

그림 10을 살펴보면 비페색과 깊이값 임계화를 통한 정보손실로 많은 부분이 빈화소로 표시됨을 관찰할

수 있다. 따라서 그림 10의 결과를 후위 계층 화소에 포함된 색상 및 깊이정보를 이용하여 보상해 주어야 하며, 보상된 다시점 복원 영상을 그림 11에 나타내었다 [19].

그림 11에서 보는 것처럼 대부분의 영역이 보상되었지만 계층적 깊이영상 생성시부터 존재했던 페색/비페색 영역과 물체의 움직임이 심한 부분에 결함이 존재한다. 여기에 원본에서 추출된 부가정보를 이용하여 최종 복원 영상을 생성하였다. 그림 12는 최종 다시점 재현 결과를 나타내며 부자연스러운 부분이 있기는 하나 대부분의 영역이 제대로 복원된 것을 알 수 있다[19].

마지막으로 그림 13은 중간 시점 보간 기법을 이용



그림 12. 최종 복원된 다시점 영상



(a) 제안한 방법을 이용한 중간 시점 보간



(b) 매칭을 이용한 고화질 중간 시점 복원 알고리즘

그림 13. 중간 시점 보간 결과 비교

해 생성한 0번과 1번 카메라 사이의 다시점 영상을 나타낸다. 그림 13 (a)는 제안한 방법을 사용해 보간한 결과이며 13 (b)는 Zitnick[14] 등의 시점 보간 기법으로 보간한 결과이다. 이 방식은 MSR에서 테스트 데이터를 제공하면서 함께 제안한 것으로 물체의 경계부분을 디지털 매칭기법을 사용해 별도로 처리하고 다시점 영상과 이 경계면 데이터를 함께 블렌딩하여 보간하는

방식이다. 별도의 알고리즘을 통해 고화질의 중간 시점 복원을 목적으로 하는 만큼 제안한 방식보다 더욱 자연스러운 시점을 재생함을 관찰할 수 있다. 그림 13 (a)의 결과가 13 (b)에 비해 흐리고 명확하지 않은 것을 확인할 수 있다.

기존의 다시점 동영상 부호화 기법은 원본 시점의 복원을 목적으로 하며 중간 시점 보간을 위해서는 영상

사이의 정합정보를 계산하는 독립적인 알고리즘 및 모듈이 필요하다. 하지만 제안한 방식은 다시점 동영상 데이터를 새로운 표현으로 변환함으로써 별도의 모듈 없이 효율적인 부호화와 중간 시점 보간 기능을 동시에 제공한다는 점에 의미가 있다.

VI. 결 론

다시점 동영상은 삼차원 TV를 비롯한 다양한 응용 분야에 사용될 중요한 실감형 멀티미디어의 하나이다. 본 논문에서는 영상기반 렌더링 기법 중 계층적 깊이영상의 개념을 이용하여 다시점 동영상을 효과적으로 표현 및 부호화하고 원본 시점과 중간 시점을 복원하는 시스템을 제안하였다. 제안한 방법은 다시점 동영상을 삼차원 정보가 포함된 계층적 깊이영상 표현으로 나타냄으로써 다시점 영상들이 갖는 다양한 정보를 구조적으로 표현하였다. 또한 이를 압축함으로써 데이터양을 줄이고, 본래의 다시점 영상 및 중간 시점을 보간하였다. 향후, 다시점 복원 및 중간 시점 보간 결과의 화질 향상을 위한 기법에 대해 추가로 연구할 예정이다.

[참고문헌]

- [1] J. Shade, S. Gortler, and R. Szeliski, "Layered Depth Image," Proc. of ACM SIGGRAPH, 1998, pp. 291-298.
- [2] "Call for Evidence on Multi-view Video Coding," ISO/IEC JTC1/SC29/WG11 N6720, Oct. 2004.
- [3] "Call for Proposals on Multi-view Video Coding," ISO/IEC JTC1/SC29/WG11 N7327, Jul. 2005.
- [4] "Subjective Test Results for the CIP on Multi-view Video Coding," ISO/IEC JTC1/SC29/WG11 N7779, Jan. 2006.
- [5] "Description of Core Experiments in MVC," ISO/IEC JTC1/SC29/WG11 N8019, Apr. 2006.
- [6] H. Shum and S. Kang, "Survey of Image-based Representations and Compression Techniques," IEEE Trans. on Circuits and Systems for Video Technology, Vol. 13, No. 11, 2003, pp. 1020-1037.
- [7] S. Gortler, R. Grzeszczuk, R. Szeliski, and M. Cohen, "The Lumigraph," Proc. of ACM SIGGRAPH, 1996, pp. 43-54.
- [8] M. Levoy and P. Hanrahan, "Light Field Rendering," Proc. of ACM SIGGRAPH, 1996, pp. 31-42.
- [9] S. Seitz and C. Dyer, "View Morphing," Proc. of ACM SIGGRAPH, 1996, pp. 21-30.
- [10] <http://grail.cs.washington.edu/projects/ldi/>
- [11] T. Kanade, P. Rander, and P. Narayanan, "Virtualized Reality: Constructing Virtual Worlds from Real Scenes," IEEE Multimedia Magazine, Vol.1, No. 1, 1997, pp. 34-47.
- [12] W. Matusik, C. Buehler, R. Raskar, L. McMillan, and S. Gortler, "Image-based Visual Hulls," Proc. of ACM SIGGRAPH, 2000., pp. 369-374
- [13] J. Yang, M. Everett, C. Buehler, and L. McMillan, "A Real-time Distributed Light Field Camera," Eurographics Workshop on Rendering, 2002, pp. 77-85.
- [14] C. Zitnick, S. Kang, M. Uyttendaele, S. Winder, and R. Szeliski, "High-quality Video View Interpolation using a Layered Representation," Proc. of ACM SIGGRAPH, 2004, pp. 600-608.
- [15] L. McMillan, "An Image-Based Approach to Three-dimensional Computer Graphics," Ph.D. dissertation, Univ. North Carolina at Chapel Hill, 1997.
- [16] Interactive visual media group at Microsoft Research, <http://research.microsoft.com/vision/InteractiveVisualMediaGroup/3DVideoDownload/>
- [17] "Generation and Coding of Layered Depth Images for Multi-view Video," ISO/IEC JTC1/SC29/WG11 m12485, Oct. 2005.
- [18] S. Yoon, S. Kim, E. Lee, and Y. Ho, "A Framework for Multi-view Video Coding using Layered Depth Images," Lecture Notes in Computer Science (LNCS), Vol. 3767, 2005, pp. 431-442.
- [19] S. Yoon, E. Lee, S. Kim, Y. Ho, K. Yun, S. Cho, and N. Hur, "Coding of Layered Depth Images Representing Multiple Viewpoint Video," Proc. of Picture Coding Symposium, SS3-2, 2006, pp. 1-6.
- [20] J. Duan and J. Li, "Compression of the LDI," IEEE Trans. on Image Processing, Vol.12, No. 3, 2003, pp.365-372.



윤 승 옥
(Seung-Uk Yoon)

2000. 2: 서강대학교 전자공학 학사

2002. 8: 광주과학기술원 정보통신공학 석사

2002. 9~현재: 광주과학기술원 정보통신공학과
박사과정

관심분야: 다시점 색상 및 깊이정보 표현 및 부호화,
계층적 깊이영상 부호화 및 영상기반 렌더링,
실감미디어 처리 및 실감방송

E-mail: suyoon@gist.ac.kr

Tel: +82-62-970-2263

Fax: +82-62-970-3164



호 요 성
(Yo-Sung Ho)

1981. 2: 서울대학교 전자공학 학사

1983. 2: 서울대학교 전자공학 석사

1989. 12: Univ. of California at Santa Barbara,
전기컴퓨터 공학 박사

1983. 3~1995. 9: 한국전자통신연구소 선임연구원

1990. 1~1993. 5: 미국 Philips 연구소 선임연구원

1995. 9~현재: 광주과학기술원 정보통신공학과 교수

관심분야: 디지털 신호처리, 영상 신호처리 및 압축,
멀티미디어 시스템, 디지털 TV와
고선명 TV, MPEG 표준, 3차원 TV,
실감방송

E-mail: hoyo@gist.ac.kr

Tel: +82-62-970-2211

Fax: +82-62-970-3164