



일반논문 (Regular Paper)

방송공학회논문지 제30권 제2호, 2025년 3월 (JBE Vol.30, No.2, March 2025)

<https://doi.org/10.5909/JBE.2025.30.2.238>

ISSN 2287-9137 (Online) ISSN 1226-7953 (Print)

데이터 압축이 시점 생성 딥러닝 모델 성능에 미치는 영향

오 수 민^{a)}, 윤 승 환^{a)}, 김 정 우^{a)}, 이 학 범^{a)}, 서 영 호^{a)}, 김 동 욱^{a)†}

Impact of Data Compression on the Performance of View Synthesis Deep Learning Model

Sumin Oh^{a)}, Seunghwan Yoon^{a)}, Jungwoo Kim^{a)}, Hakbum Lee^{a)},
Young-Ho Seo^{a)}, and Dong-Wook Kim^{a)†}

요 약

본 논문에서는 MV-HEVC(Multiview High Efficiency Video Coding)를 활용한 다중 시점 영상 압축이 딥러닝 기반 시점 생성 모델의 성능에 미치는 영향을 분석하였다. 이를 위해 하나의 씬(Scene)에 대해 촬영된 단일 데이터셋을 활용하여, 21개의 시점을 세 개씩 그룹화한 후, MV-HEVC를 이용하여 다양한 압축 강도(QP 값)로 인코딩 및 디코딩을 수행하였다. 이후, 복호화된 데이터를 학습 데이터로 활용하여 4D 가우시안 스플래팅(4D Gaussian Splatting) 모델을 학습하고, 새로운 시점의 영상을 생성하였다. 실험 결과, QP 값이 증가함에 따라 PSNR이 감소하는 경향을 보였으나, 딥러닝 모델이 생성한 새로운 시점 영상의 품질에는 큰 영향을 미치지 않음을 확인하였다. 또한, 시각적으로 유사한 카메라를 그룹화하는 방식과 COLMAP 기반의 자동 그룹화 방식을 비교한 결과, 두 방식 모두 압축된 데이터셋을 활용하더라도 학습 및 시점 생성 성능이 유지됨을 보였다. 다만, 본 연구는 단일 씬(Scene)을 대상으로 실험을 수행했기 때문에 다양한 환경에서의 일반화 가능성에 대한 추가적인 검증이 필요하다. 그럼에도 불구하고, 본 연구를 통해 다중 시점 영상 데이터셋을 높은 압축률로 저장 및 관리하더라도 딥러닝 기반 시점 생성 모델의 성능이 유지될 가능성을 실험적으로 확인하였다. 이는 향후 VR(Virtual Reality), AR(Augmented Reality), 그리고 3D 영상 합성 기술에서 효율적인 데이터 활용 방안을 제시할 수 있을 것으로 기대된다.

Abstract

This paper analyzes the impact of multiview video compression using MV-HEVC (Multiview High Efficiency Video Coding) on the performance of deep learning-based novel view synthesis models. To achieve this, a single dataset captured for one scene was used, where 21 viewpoints were grouped into sets of three, and encoding and decoding were performed with various compression levels (QP values) using MV-HEVC. The decoded data was then utilized to train a 4D Gaussian Splatting model for generating novel view images. Experimental results showed that as the QP value increased, the PSNR tended to decrease. However, the quality of the synthesized novel view images generated by the deep learning model remained largely unaffected. Additionally, a comparison between manually grouping visually similar cameras and an automatic grouping method based on COLMAP demonstrated that both approaches maintained the performance of the model, even with compressed datasets. Nevertheless, since this study was conducted on a single scene, further verification is required to assess its generalizability to diverse environments. Despite this limitation, the study empirically confirms that multiview video datasets can be stored and managed at high compression rates without significantly degrading the performance of deep learning-based novel view synthesis models. These findings suggest the potential for efficient data utilization in VR (Virtual Reality), AR (Augmented Reality), and 3D video synthesis technologies.

Keyword : MV-HEVC, 4D Gaussian Splatting, Multiview Video Compression, View Synthesis

Copyright © 2025 Korean Institute of Broadcast and Media Engineers. All rights reserved.

“This is an Open-Access article distributed under the terms of the Creative Commons BY-NC-ND (<http://creativecommons.org/licenses/by-nc-nd/3.0>) which permits unrestricted non-commercial use, distribution, and reproduction in any medium, provided the original work is properly cited and not altered.”

1. 서론

최근 열린 2025년 CES에서는 메타버스와 확장현실(XR) 기술이 인공지능(AI)과 결합하여 다시 주목받고 있음을 보여주었다^[1]. 그동안 정체되었던 메타버스 산업은 코로나 상황과 같은 비대면 환경에 대한 수요 증가에 따라 AI 기술의 발전과 더불어 새로운 발전의 계기를 맞고 있다. 이러한 변화는 콘텐츠의 몰입감을 높이는 수으로 이어지고 있으며, 따라서 3D 콘텐츠와 자유 시점 영상 기술에 대한 수요가 확대되고 있다^{[2][3]}.

시점에 대한 이러한 욕구를 충족하기 위해서는 수요자에게 더욱 많은 시점의 영상을 제공하여야 하나, 수요자를 만족시킬 만큼의 시점 데이터를 제공하기 위해서는 엄청난 양의 데이터를 전송하여야 하기 때문에 필요한 데이터를 전송하는 대신 수요자 측에서 요구하는 시점의 데이터를 생성하는 가상시점 합성 기술이 도입되었다^[4]. 이 기술은 다시점 영상(multi-view video)을 사용하여 임의의 가상시점 영상을 생성함으로써 데이터 효율성과 사실감을 동시에 구현할 수 있는 방법으로 주목받고 있다.

가상 시점 합성에는 이미지 공간 시스템(image domain system)^[5], 레이 공간 시스템(ray space system)^[6], 모델-기반 시스템(model-based system)^[7] 등 다양한 시스템이 사용된다. 이 중 모델-기반 시스템의 대표적인 기술로 Neural Radiance Fields(NeRF)^[8]와 3D Gaussian Splatting^[9]이 있다. NeRF는 3차원 공간의 라디언스(radiance) 값을 학습하여 임의의 시점에서의 영상을 생성할 수 있는 기술로, 뛰어난 표현력과 사실감으로 주목받고 있다. 그러나 NeRF는 학습 및 추론 과정에서 대규모 계산 자원을 필요로 하고, 대규모 원본 데이터셋에 의존해야 한다는 한계를 가진다^[10]. 3D

Gaussian Splatting은 이러한 제약을 완화하기 위해 제안된 기술로, 3차원 데이터를 가우시안(Gaussian) 분포로 표현해 복잡한 계산 과정을 단순화하고, 효율적인 학습과 실시간 시점 합성을 가능하게 한다^[11]. 특히, 4D Gaussian Splatting은 3D Gaussian Splatting을 시간적 요소로 확장한 모델로, 동적인 장면을 표현하고 렌더링(rendering) 성능을 향상시키는 데 기여하고 있다^{[12][13][14]}. 이러한 기술들은 시점 생성 네트워크의 장점을 확장하는 데 기여하고 있지만, 여전히 대규모 학습 데이터의 저장 및 관리하여야 하는 실질적인 어려움은 아직 남아 있다.

본 논문은 새로운 시점 생성에 방대한 양의 데이터가 필요하다는 문제를 해결하고자 시점 생성에 원본 데이터가 아닌 압축된 데이터를 사용하면 생성된 영상의 화질에 어떤 영향을 미치는지 분석하여 임의의 시점 영상 생성에 압축된 데이터의 사용 가능성을 확인하고자 한다. 이를 위해 다시점 영상을 MV-HEVC(Multi-View High Efficiency Video Coding)로 압축하여 그 결과의 다시점 영상을 학습 데이터로 활용하는 딥러닝 모델의 성능을 평가한다^[15-22]. 기존 연구에서는 압축 데이터가 딥러닝 모델의 학습과 추론 성능에 미치는 영향을 체계적으로 분석한 연구가 없었으며, 따라서 본 논문에서는 MV-HEVC의 압축 강도(Quantization Parameter, QP)에 따른 학습 데이터의 품질 변화와 딥러닝 모델 성능 간의 관계를 정량적으로 평가한다. 이를 통해 대규모 데이터셋의 효율적인 저장 및 관리 가능성을 제시하고, 딥러닝 모델의 성능을 유지하면서도 데이터 자원 요구를 줄일 수 있는 실질적인 방안을 제시하고자 한다.

본 논문은 다음과 같이 구성된다. 2장에서는 본 논문에서 사용된 주요 기술인 MV-HEVC를 기반으로 한 다중 시점 영상 압축 기술과 4D Gaussian Splatting 모델의 원리 및 주요 구성 요소를 소개한다. 3장에서는 본 논문의 전체적인 실험 설계와 분석 방법을 설명한다. 즉, MV-HEVC를 활용한 데이터 압축 과정과 4D Gaussian Splatting 모델 학습 및 새로운 시점 생성 과정을 구체적으로 설명하며, 압축 강도에 따른 딥러닝 모델 성능 평가 방법을 제시한다. 4장에서는 실험 결과를 통해 압축 데이터셋의 품질 변화와 딥러닝 모델 성능 간의 관계를 정량적으로 분석하고, 이 데이터를 기반으로 5장에서는 본 논문의 결론을 맺는다.

a) 광운대학교 전자재료공학과(Department of Electronic Materials Engineering, Kwangwoon University)

‡ Corresponding Author : 김동욱(Dong-Wook Kim)

E-mail: dwkim@kw.ac.kr

Tel: +82-2-940-8362

ORCID: <https://orcid.org/0000-0001-6106-9894>

※ 본 논문은 2023년도 광운대학교 연구년에 의하여 연구되었습니다.

※ 본 논문은 과학기술정보통신부 및 정보통신기획평가원의 대학ICT연구센터 육성지원사업의 연구결과로 수행되었습니다. (IITP-2024-RS-2022-00156225)

· Manuscript February 27, 2025; Revised February 27, 2025; Accepted February 28, 2025.

II. 관련 기술

1. MV-HEVC

MV-HEVC(H.265)는 표준으로 확장된 기술로, 다중 시점 비디오 데이터를 효율적으로 압축하고 전송하기 위해 설계되었다^[15-22]. 그림 1은 MV-HEVC 부호화기 개념도를 보이고 있다. 그림에서 큰 블록은 한 시점에 대한 압축 부호화를 나타내고, 블록 간에는 시점 간 압축을 보이고 있다. 이 기술은 시점 간 예측(inter-view prediction)을 기반으로, 기존 단일 시점 기반 방법에서 사용하는 화면 내(intra) 예측 및 화면 간(inter) 예측과 더불어, 다른 시점의 데이터를 참조하여 공간적 중복성을 제거한다. 이를 통해 다중 시점 데이터에서 발생하는 중복을 최소화하며, 압축 효율을 높인다.

그림 1을 자세히 살펴보면, 기본 시점(Base view)은 독립적으로 부호화되며, 기존 HEVC와 동일한 방식으로 화면

내(intra) 및 화면 간(inter) 예측을 수행한다. 이후, 예측된 영상은 변환(Transform), 양자화(Quantization), 그리고 엔트로피 부호화(Entropy Encoding) 과정을 거쳐 최종적으로 비트스트림이 생성된다.

반면, 확장 시점(Extended view)은 기본 시점의 복호화 데이터를 참조하여 시점 간 예측(Inter-view Prediction)을 수행하며, 추가적으로 시차 보상 예측(Disparity Compensation Prediction, DCP)을 적용하여 부호화한다. 확장 시점의 부호화 과정에서는 동일한 POC(Picture Order Count)를 갖는 기본 시점의 데이터를 참조하며, 시점 간 중복을 줄이기 위해 시차 벡터(disparity vector)를 활용한다.

또한, MV-HEVC는 부호화된 데이터를 효과적으로 관리하기 위해 디코딩 프레임 버퍼(DPB, Decoded Picture Buffer)를 사용하며, 부호화 과정에서 발생하는 부작용을 줄이기 위해 In-loop 필터(In-loop Filter)를 적용한다. 이 필터는 부호화된 프레임의 품질을 보정하여 다음 프레임의 예측 성능을 향상시킨다.

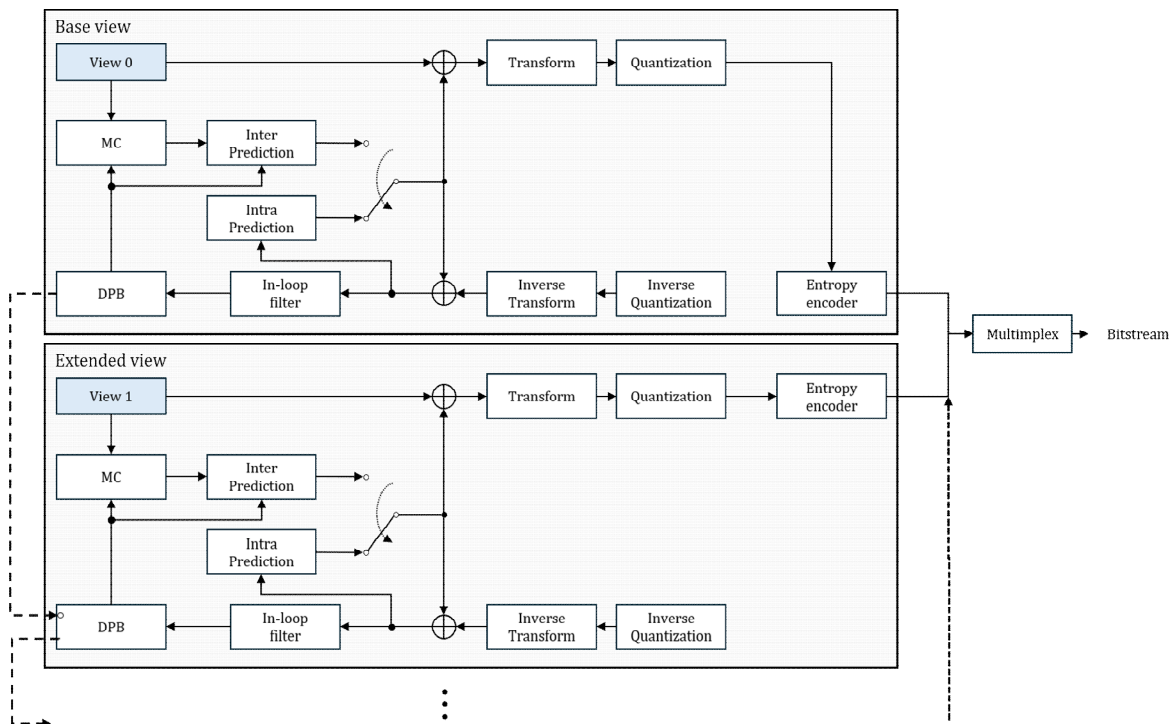


그림 1. MV-HEVC 부호화기 개념도

Fig. 1. MV-HEVC Encoder Architecture Diagram

최종적으로, 부호화된 기본 시점과 확장 시점의 데이터는 멀티플렉서(Multiplexer)를 통해 하나의 비트스트림으로 통합되며, MV-HEVC 기반 다중 시점 비디오 압축 데이터가 생성된다.

본 논문에서는 MV-HEVC의 부호화 기법을 활용하여, 압축 강도에 따른 다중 시점 영상의 품질을 정량적으로 분석하고, 이를 학습 데이터로 활용하여 딥러닝 기반 새로운 시점 생성 모델의 성능에 미치는 영향을 평가한다.

2. 4D 가우시안 스플래팅 모델

본 논문에서는 3D 생성 딥러닝 모델로 그림 2에 나타난 4D 가우시안 스플래팅^{[4][5][6]} 모델을 채택하였다. 이 모델은 3D 포인트 클라우드 데이터를 기반으로 각 점에 가우시안 분포를 적용하여 3D 장면을 표현하고, 이를 공간적 및 시간적 변화를 반영해 동적으로 렌더링하는 신경 렌더링 기법이다. 이 모델은 3D 장면의 복잡한 구조와 시간적 변화를 학습하며, 이를 통해 새로운 시점에서의 이미지를 생성하는데 활용된다. 이 모델은 세 단계로 구성되며, 각각 3D 포인트 클라우드 및 가우시안 생성, 가우시안 변형 필드 네트워크, 그리고 렌더링 과정으로 이루어진다.

2.1 3D 포인트 클라우드와 가우시안 생성

4D 가우시안 스플래팅에서 사용하는 3D 포인트 클라우드는 SfM(Structure-from-Motion) 기법을 통해 생성된다. 이 과정에서 2D 이미지 간 특징점을 매칭하여 3D 월드 좌표계를 정의하고, 이를 기반으로 포인트 클라우드를 생성한다. 생성된 포인트는 초기 가우시안의 중심으로 설정되며, 각 가우시안은 크기, 방향, 색상 등의 속성을 포함하여 장면을 표현한다. 이후, 가우시안은 변형 과정을 거쳐 동적 장면을 학습하게 된다.

2.2 가우시안 변형 필드 네트워크 및 렌더링

가우시안 변형 필드 네트워크는 공간-시간적 구조 인코더와 가우시안 변형 디코더로 구성된다. 인코더는 3D 가우시안과 시간 정보를 입력받아, 다중 해상도 2D 평면 모듈과 다중 퍼셉트론(MLP)을 활용하여 공간 및 시간 특징을 추출한다. 평면 모듈은 (x,y) , (x,z) , (y,z) , (x,t) , (y,t) , (z,t) 조합으로 구성되며, 3D 정보를 2D로 압축하고 시간 정보는 보간(interpolation)을 통해 통합된다. 이후, 선형 보간법을 사용해 복셀 특징을 계산하고, 이를 MLP를 통해 최종적인 공간-시간적 특징으로 변환한다.

디코더는 인코딩된 특징을 기반으로 3D 가우시안의 위

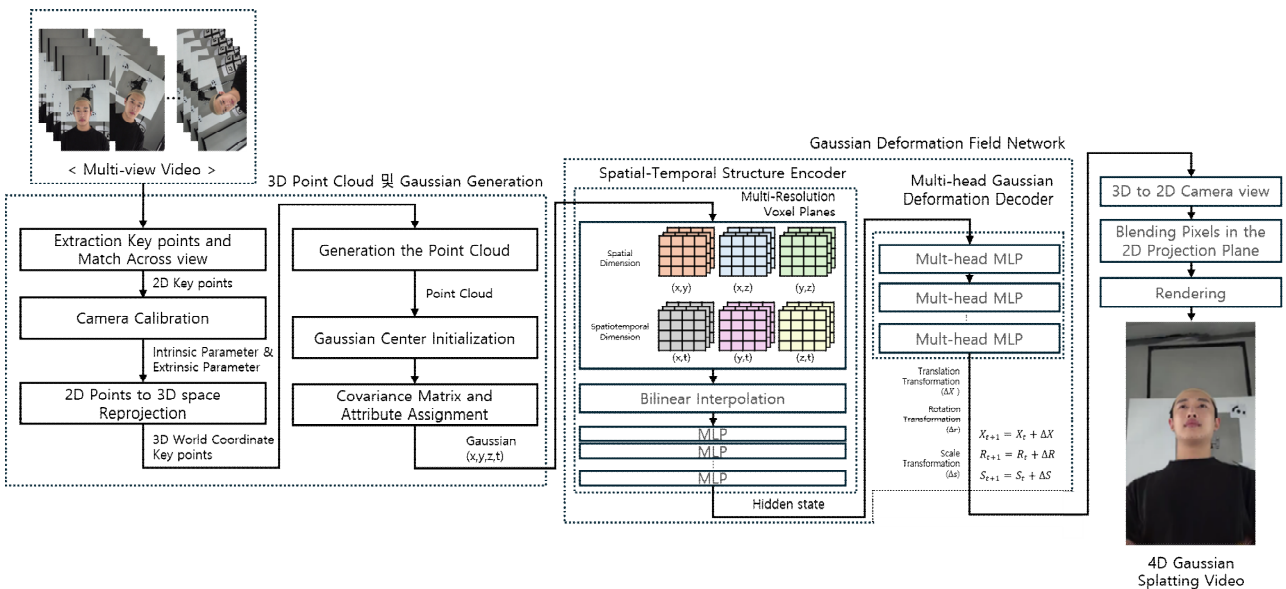


그림 2. 4D Gaussian Splatting 내부 구조도
Fig. 2. Internal Structure of 4D Gaussian Splatting

치, 회전, 스케일 변형 값을 예측한다. 각 변형 값은 독립적인 MLP를 통해 계산되며, 초기 값에 더해져 최종적인 변형 3D 가우시안이 정의된다. 변형된 3D 가우시안은 위치, 회전, 스케일뿐만 아니라 분산 및 색상 정보도 포함하여 장면을 더욱 정교하게 표현한다.

이후, 변형된 3D 가우시안을 기반으로 차분 스플래팅 연산이 수행되어 새로운 시점의 이미지가 생성된다. 이 과정에서 변형된 가우시안의 분포가 2D 이미지 평면으로 투영되며, 각 가우시안의 위치, 크기, 색상, 회전 값을 반영하여 픽셀 값을 계산한 후 새로운 시점의 영상을 렌더링한다.

본 논문에서는 압축된 데이터를 사용하여 이 모델을 학습하고, 그 결과로 임의의 시점을 생성하여, 학습에 사용한 데이터의 압축 정도가 생성된 시점의 화질에 어떤 영향을 미치는지를 분석하고자 한다.

III. 실험 설계

1. 전체적인 실험 네트워크 구조

본 논문의 목적인 압축된 데이터로 학습한 모델로 시점을 생성하여 압축 정도가 생성된 시점의 화질에 미치는 영향을 분석하기 위해 그림 3과 같이 실험을 수행할 네트워크를 설계하였다. 이 설계는 먼저 촬영된 다시점 영상을 세 시점들로 그룹핑하고(그림에서 3-View Grouping), MV-HEVC로 압축 부호화하고 복호화한다(그림의 MV-HEVC).

그 결과 영상 데이터셋을 학습 데이터로 사용하여 4D 가우시안 스플래팅 모델을 학습시키고, 학습된 모델을 통해 새로운 시점 영상을 생성한다(그림의 4D Gaussian Splatting). 다양한 압축 강도로 압축된 데이터셋을 학습 데이터로 사용함으로써 압축 강도가 생성된 시점 영상의 화질에 미치는 영향을 분석한다.

2. MV-HEVC 기반 다시점 영상 압축 부호화 및 복호화

본 논문에서는 총 21개의 시점 데이터를 카메라 위치 관계가 유사한 세 개의 시점으로 그룹화하여 실험을 진행하였다. 총 7개의 그룹 중 관찰 대상 그룹을 제외한 6개의 그룹에서 MV-HEVC 부호화를 수행하였으며, 각 그룹에서 중간 시점을 기본 시점으로 설정하고 이를 기준으로 좌측 및 우측 시점간의 P-I-P 구조를 적용하였다. 그림 4는 P-I-P 참조 구조를 나타내며, 이는 중간 시점이 기본 시점으로 독립적으로 부호화되고, 좌측 시점과 우측 시점이 중간 시점을 참조하여 부호화되는 방식이다. 중간 시점(시점 1)은 가장 먼저 부호화되며, 이후 좌측 시점(시점 0)과 우측 시점(시점 2)은 중간 시점을 참조하여 시점 간 예측을 통해 부호화와 복호화가 이루어진다.

본 논문에서는 그림 4에서 제시된 P-I-P 구조를 기반으로 기본 시점을 GOP(Group of Pictures) 크기 8의 계층적 참조 구조로 부호화하였다. 화면 내 예측을 통해 생성된 I-픽처는 기준 영상으로 활용되며, 비 기준 영상(Non-key

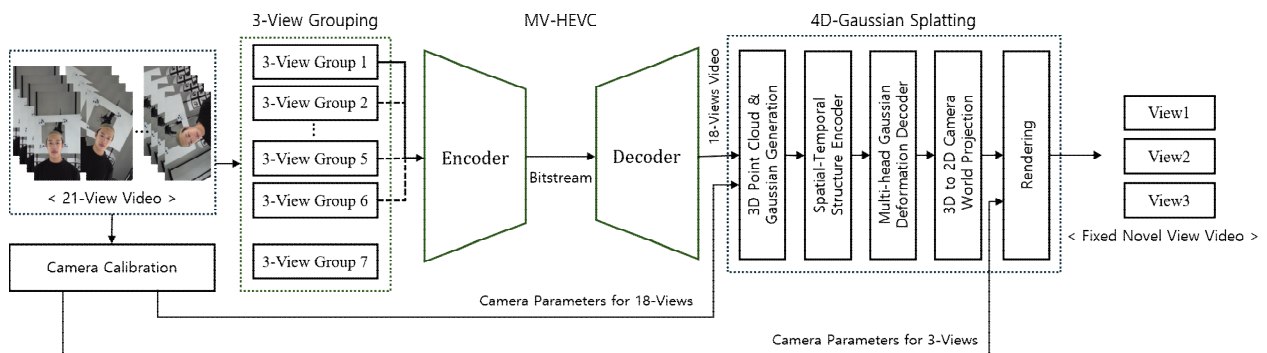


그림 3. 실험 설계 구조 개요

Fig. 3. Overview of the Experimental Design Structure

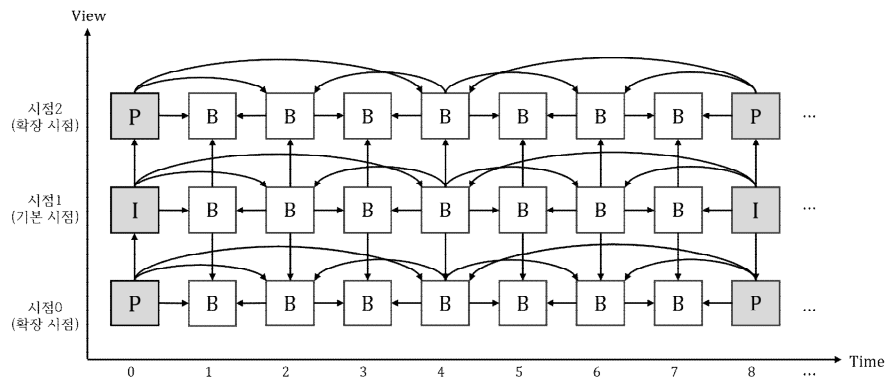


그림 4. P-I-P 참조 구조
Fig. 4. P-I-P Reference Structure

picture)은 기존의 계층적 참조 구조를 따르거나 동일한 POC(Picture Order Count)를 가지는 영상을 참조하여 부호화된다. 이러한 방식으로 부호화된 다시점 영상은 3D 생성 딥러닝 모델의 학습 데이터로 활용되며, 압축된 학습 데이터가 딥러닝 모델 성능에 미치는 영향을 평가하는 데 사용되었다.

3. 4D 가우시안 스피래팅 기반 학습 및 시점 생성

압축된 다시점 영상 데이터를 활용하여 원하는 새로운 시점에서의 영상을 생성하기 위해 본 논문은 4D 가우시안 스피래팅 구조를 채택하였다. 그림 2와 같이 4D 가우시안 스피래팅은 3D 포인트 클라우드 생성, 가우시안 생성 및 변형, 그리고 렌더링의 세 단계로 구성되며, 본 논문에서는 이 구조를 차용하면서 일부 모델 구조를 변경하여 사용하였다.

기존 4D 가우시안 스피래팅에서 사용되는 3D 포인트 클라우드는 SfM 기반의 SIFT(Scale Invariant Feature Transform) 알고리즘을 통해 생성된다. SIFT 알고리즘은 다시점 영상에서 각 2D 이미지의 특징점을 추출하고, 서로 다른 시점에서 촬영된 이미지 간의 동일한 특징점을 매칭하여 2D 키포인트를 생성한다. 이후 카메라 보정을 통해 카메라의 내부 파라미터와 외부 파라미터를 추정한다. 2D 키포인트는 추정된 카메라 매트릭스를 통해 3D 공간으로 재투영되어 2D 이미지의 픽셀 좌표가 3D 월드 좌표계로 변환된다. 결과적으로 생성된 3D 월드 키포인트는 포인트

클라우드를 형성하게 된다. 이 과정은 그림 2의 3D 포인트 클라우드 및 가우시안 생성 단계에 해당한다.

본 논문에서는 21개의 시점 중 3개를 관찰 시점으로 선택하였으며, 이는 그림 3의 3-View Group 7에 해당한다. SIFT 알고리즘을 활용하여 21개 시점의 포인트 클라우드와 카메라 월드 좌표계를 생성한 후, 이 중 관찰 시점에 해당하는 3개의 카메라 정보를 추출하였다.

4D 가우시안 스피래팅의 가우시안 변형 필드 네트워크에서는 관찰 시점을 제외한 18개 시점의 포인트 클라우드 정보를 입력으로 사용하여 학습을 진행하였다. 추론 과정에서는 학습에 제외되었던 관찰 시점의 카메라 정보를 이용하여 해당 3개 시점에서 고정된 영상을 생성하였다.

생성된 영상은 4D 가우시안 스피래팅 모델의 성능 평가 기준이 되며, 원본 영상과 비교하여 PSNR을 계산함으로써 딥러닝 모델의 성능을 측정하였다. 본 연구에서는 이러한 실험 구조를 통해 압축된 데이터셋이 모델 성능에 미치는 영향을 효과적으로 분석할 수 있는 환경을 구축하였다.

IV. 실험 및 결과

1. 실험

1.1 데이터셋

본 논문의 목적을 위해 원본 고화질 다시점 비디오 영상이 필요하였다. 그러나 기존에 공개된 다시점 영상 데이터

는 대부분 MP4 형식으로 이미 압축된 상태로 제공되므로, 원본 영상과 압축 영상 간의 딥러닝 성능 차이를 분석하려는 본 논문의 목적에 부적합하였다. 따라서 고화질 원본 다 시점 영상 촬영을 통해 실험에 필요한 데이터셋을 직접 확보하였다.

다시점 영상은 한 명의 피사체를 대상으로 하나의 씬(Scene)에 대해 촬영되었으며, 총 63대의 3,840×2,160(4K UHD) 해상도 카메라를 활용하여 8초 길이의 영상을 30fps로 촬영하였다. 피사체를 다양한 각도에서 기록할 수 있도록 카메라 장비를 원형으로 배치하였으며, 촬영 과정에서 각 카메라는 동일한 시간대의 장면을 기록하도록 동기화되었다. 이를 통해 다중 시점 간 정확한 상관성을 유지하였다. 이 중 상대적으로 피사체의 정면에 해당하는 21대의 카메라를 선별하여 실험에 사용하였다. 선별된 카메라는 피사체를 중심으로 좌우 대칭적으로 배치되어, 정면을 포함한 다양한 각도의 시점을 포함하고 있다.

촬영된 영상은 한국어 발화를 수행하는 피사체가 다양한 얼굴 표정과 움직임을 포함하여 자연스럽게 말하는 장면을 담고 있다. 이를 통해 딥러닝 모델 학습과 압축 강도에 따른 성능 평가에 적합한 데이터를 확보하였다. 또한, 딥러닝 학습에 활용한 기준 데이터를 구축하기 위해 촬영 과정에서 정밀한 스틸-컷(still-cut) 촬영을 병행하였다. 스틸-컷은 각 시점에서 완벽히 정지된 이미지를 제공하며, 정밀한 품질 비교와 학습 데이터 생성을 위한 기준 데이터로 사용하였다.

1.2 실험환경

본 논문에서는 실험 데이터의 특성을 고려하여 MV-HEVC 압축 파라미터 값을 표 1과 같이 설정하였다. 실험에 사용된 영상은 3,840×2,160 해상도로 구성되며, 30fps로 총 240프레임(FramesToBeEncoded=240)을 인코딩 대상으로 설정하였다. 영상의 휘도와 색도 정보를 처리하기 위해, 입력·출력·내부 비트 심도(InputBitDepth, OutputBitDepth, InternalBitDepth)와 색도 비트 심도(InputBitDepthC, OutputBitDepthC, InternalBitDepthC)는 모두 8비트로 동일하게 설정하였다. 또한, 색 정보는 4:2:0 크로마 포맷(InputChromaFormat=420)을 적용하여, 고해상도에서도 데이터 양을 효율적으로 관리하면서 안정적인 화질을 유지할 수 있도록 하였다. 압축 설정은 Random Access 모드를 기반

으로 하였으며, GOP 크기를 8로 설정하여 장면 전환이 없고 시간 경과에 따라 변화가 적은 정적인 발화 영상의 특성에 맞게 효율적으로 부호화를 수행하였다. IntraPeriod는 -1로 지정하여 첫 번째 프레임만 I-Frame으로 부호화함으로써, 고정된 장면에서 불필요한 부호화를 최소화하였다. 특히, 발화 영상의 정적인 특성을 고려하여 Search Range를 64로 설정하여 시점 간 참조 과정에서의 탐색 범위를 제한하여 계산량을 줄이고 인코딩 시간을 단축하였다. 또한, 압축 강도를 조절하기 위해 양자화 매개변수(QP:Quantization Parameter)를 18, 24, 30, 36, 42로 설정하였으며, 이를 통해 압축 비율에 따른 품질 변화를 분석하였다. 이외에도 루프 필터(LoopFilter)와 샘플 적응 오프셋(SAO), 변환 스킵(Transform Skip)을 활성화하여 실험을 진행하였다.

표 1. MV-HEVC 파라미터들
Table 1. MV-HEVC parameters

Parameter	Value
Mode	Random access
SourceHeight	2,160
SourceWidth	3,840
InputBitDepth	8
OutputBitDepth	8
InternalBitDepth	8
InputBitDepthC	8
OutputBitDepthC	8
InternalBitDepthC	8
InputChromaFormat	420
FrameRate	30
FramesToBeEncoded	240
Profile	main main multiview-main
GOP	8
IntraPeriod	-1
Search Range	64
QP	18, 24, 30, 36, 42
LoopFilterDisable	0
SAO	1
TransformSkip	1

1.3 실험순서

그림 5는 본 논문의 실험 순서도를 나타낸다. 우선, MV-HEVC 압축을 위해 21개 시점의 원본 데이터셋을 3개 씩 묶어 총 7개의 그룹으로 분할한다. 그룹화 방식은 두 가

지로 진행된다. 첫 번째 방식은 21개 시점 중 시각적으로 유사한 3개 시점을 묶는 것이며, 두 번째 방식은 카메라 캘리브레이션(Camera Calibration) 기법 중 COLMAP을 활용하는 것이다. 두 번째 방식을 자세히 설명하자면, 먼저 COLMAP을 활용하여 각 카메라의 외부 파라미터(회전 벡터 및 변환 벡터)를 추출한다. 이후, 이 외부 파라미터를 이용해 월드 좌표계에서 각 카메라의 센터를 계산하고, 카메라 센터들 간의 유클리드 거리를 측정한다. 마지막으로, 이러한 거리 정보를 바탕으로 공간적으로 가장 가까운 3개의 카메라를 자동으로 클러스터링하는 방식이다.

그룹화를 진행한 후, 7개의 그룹 중 새로운 시점을 생성하기 위한 한 그룹을 제외하고, 나머지 6개 그룹의 원본 데이터를 사용하여 4D 가우시안 스플래팅 모델을 학습하였다. 이후, 제외된 그룹의 3개 카메라 위치를 기준으로 새로운 시점의 영상을 추론하였으며, 생성된 추론 영상과 원본 영상을 비교하여 PSNR(Peak Signal-to-Noise Ratio) 및 SSIM(Structural Similarity Index Measure)을 측정하였다. 이 측정값은 본 논문의 데이터셋에 대한 4D 가우시안 스플래팅 모델 성능 기준으로 활용하였다.

이후, 본 논문에서는 데이터 압축이 시점 생성 딥러닝 모델 성능에 미치는 영향을 분석하기 위해, 앞서 그룹화한 7

개 그룹 중 1개를 제외한 나머지 6개 그룹에 대해 압축 실험을 수행하였다. 구체적으로, 3개의 다중 시점 영상을 대상으로 MV-HEVC 압축을 진행하였으며, 총 6개 그룹에 대해 압축을 적용하였다. 압축 강도는 QP 값을 18, 24, 30, 36, 42의 5가지 조건으로 설정하여 실험을 진행하였다. 압축된 데이터셋은 복호화된 후 4D 가우시안 스플래팅 모델 학습에 활용되었으며, 이후 제외된 그룹에 포함된 3개 카메라 위치를 기준으로 새로운 시점의 영상을 추론하였다. 최종적으로, 생성된 추론 영상과 원본 영상을 비교하여 PSNR과 SSIM을 측정하였고, 이 값을 본 연구에서 산출한 4D 가우시안 스플래팅 모델 성능 기준과 비교하여 데이터셋 압축이 시점 생성 딥러닝 모델 성능에 미치는 영향을 분석하였다.

2. 실험 결과

그림 6은 21개의 시점을 3개의 시점으로 그룹화한 7개의 그룹 중 6개 그룹에 대해 5가지 QP 값(18, 24, 30, 36, 42)으로 MV-HEVC 압축을 수행한 결과 영상을 보이고 있으며, QP 값에 따라 화질이 달라지는 것을 시각적으로 확인할 수 있다. 표 2는 QP 값의 변화에 따라 압축된 다시점 영상과

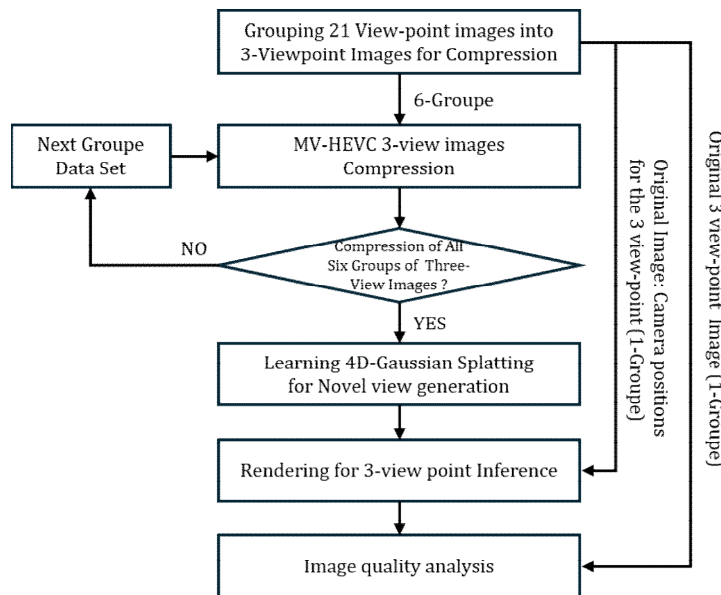


그림 5. 실험 순서도
Fig. 5. Diagram of experiment procedure



그림 6. MV-HEVC의 QP=18, 24, 30, 36, 42로 압축된 영상과 원본 영상의 비교 (각 그림 하단에 압축 강도 표시)

Fig. 6. Comparison of compressed images with QP=18, 24, 30, 36, and 40 of MV-HEVC and the original image (compression strength is shown below each image)

표 2. 압축강도에 따른 생성 영상의 PSNR 비교

Table 2. PSNR comparison of the generated images for the compression strength

QP	PSNR
18	42.00
24	41.89
30	41.01
36	39.55
42	36.31

원본 영상 간의 평균 PSNR 값을 보이고 있다. 예측한 바와 같이, QP 값이 증가할수록 PSNR 값이 점차 감소하는 경향을 보였다. 이는 QP 값이 증가함에 따라 압축 강도가 높아지고, 그에 따라 복원 영상의 품질이 점차 저하됨을 의미한다.

2.1 시각적 카메라 위치 관계 기반 그룹화 방식 실험 결과
그림 7과 표 3은 21개 시점의 원본 데이터셋에서, 카메

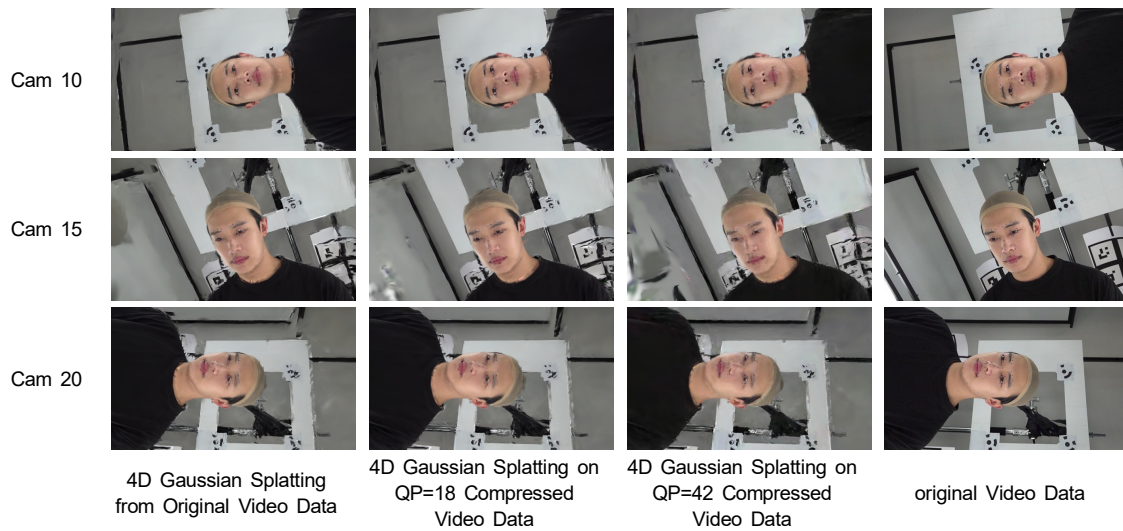


그림 7. 가장 왼쪽부터 원본 영상과 MV-HEVC로 압축(QP=18, 42)된 다시점 영상을 각각 4D 가우시안 스플래팅 학습 후 새로운 세 시점 (cam10, cam15, cam20)에 대해 추론한 영상과 원본 영상의 예

Fig. 7. From the left, we present examples of the original video and the multi-view videos compressed with MV-HEVC (QP=18, QP=42). Each was used to train a 4D Gaussian Splatting model, and the inferred results for three novel viewpoints (cam10, cam15, cam20) are shown alongside the original video

표 3. 압축 강도에 따른 딥러닝 모델의 시점 성능 비교

Table 3. Comparison of deep learning performance at different compression levels

	QP	Original	18	24	30	36	42
Cam 10	PSNR	26.1	26.58	26.96	26.85	26.3	26.4
	SSIM	0.904	0.905	0.909	0.906	0.903	0.898
Cam 15	PSNR	18.23	18.78	18.61	18.58	18.64	18.14
	SSIM	0.754	0.905	0.909	0.906	0.903	0.898
Cam 20	PSNR	23.79	23.77	23.92	23.5	23.84	24.1
	SSIM	0.858	0.862	0.862	0.856	0.856	0.852

라 위치가 시각적으로 유사한 3개 시점을 그룹화하는 방식을 적용하여, 그림 5의 실험 순서도에 따른 결과를 보여준다. Cam 10, 15, 20번은 21개 시점 중 각각 정면, 우측, 좌측 시점으로, 시각적 유사성을 바탕으로 하나의 그룹으로 묶은 것이며, 이 그룹은 새로운 시점을 생성하기 위한 카메라 위치로 활용하였다.

실험에서는 6가지 QP 값(18, 24, 30, 36, 42)으로 압축된 데이터셋을 사용하여 시점 생성 딥러닝 모델(4D 가우시안 스피래딩)을 학습한 후, 제외된 그룹의 3개 시점(Cam10, Cam15, Cam20)에 대해 새로운 영상을 추론하였다. 그림 7은 이 중 원본과 두 가지 QP 값(18, 42)을 적용하여 생성된 시점 영상의 예시를, 표 3은 각 QP 값별로 추론된 3개 시점 영상의 평균 화질(PSNR 및 SSIM)을 정량적으로 보여준다. 여기서 ‘Original’은 원본 영상을 학습 데이터셋으로 사용한 결과를 의미하며, 이는 본 논문에서 제시하는 4D 가우시안 스피래딩 모델의 기준 성능을 나타낸다.

표 3의 값들을 보면, 세 시점에 대해 화질이 꽤 많은 차이를 보이고 있는데, 이것은 학습 데이터셋이 특정 카메라 위치에 대한 충분한 정보를 포함하지 못한 시점의 경우 상대적으로 낮은 화질을, 반대로 정보가 풍부한 경우 높은 화질을 보이는 것을 의미한다. 5가지 다른 QP 값에 따른 PSNR 및 SSIM 값을 기준 성능과 비교해보면 QP 값의 증가에 대해 생성된 시점 영상의 화질은 크게 영향을 받지 않는 것으로 나타났으며, QP에 대한 어떤 경향성도 찾아볼 수 없었다.

2.2 Camera Calibration을 활용한 그룹화 방식 실험 결과

21개 시점에 대해 칼리브레이션 및 그룹화가 결과에 영향을 미치는지를 확인하기 위해서 COLMAP^{[23][24][25]}을 사용하여 실험을 진행하였다. 즉, COLMAP을 이용하여 21개

시점의 카메라 외부 파라미터를 추출하고, 이를 통해 각 카메라 센터를 계산한 후, 유클리드 거리를 기반으로 공간적으로 인접한 카메라들을 자동 그룹핑하였다. 이 과정을 통해 Cam 4, Cam 15, Cam 20이 하나의 그룹으로 묶였으며, 이 그룹은 새로운 시점 생성을 위한 기준 카메라 위치로 활용하였다. 실험방법은 그림 5와 동일하였으며, 각 데이터 압축에 대해 Cam 4, Cam 15, Cam 20 위치에서 새로운 시점 영상을 추론하였다.

이 실험의 결과는 그림 8에서 원본 영상과 QP=18과 QP=42로 압축된 영상에서 학습된 모델을 통해 추론된 결과를 순서대로 보이고 있다. 표 4는 각 QP 조건별로 Cam 4, Cam 15, Cam 20에 대해 산출된 PSNR 및 SSIM 평균값을 정량적으로 보이고 있다. 표 4의 값을 보면 Cam 15와 Cam 20은 그림 7과 표 3의 실험에도 포함된 시점이며, 그 값도 약간의 차이를 보일 뿐임을 알 수 있다. Cam 4는 앞의 실험에서는 포함되지 않은 시점이다. 표 4의 세 시점 모두 앞의 실험과 유사하게 학습에 사용한 데이터의 압축 강도에 아무런 경향성을 보이고 있지 않음을 알 수 있다.

두 실험 결과, 압축된 데이터셋을 활용하여도 시점 생성 딥러닝 모델의 성능이 크게 저하되지 않음을 시사한다. 따라서 방대한 다시점 영상 데이터셋을 저장 및 관리할 때 높은 압축률을 적용하더라도 새로운 시점 생성 딥러닝 모델의 성능 유지에는 문제가 없는 것을 뜻한다.

다만, 본 연구는 하나의 씬(Scene)에 대해 촬영된 단일 데이터셋을 기반으로 실험이 진행되었기 때문에 결과의 일반화에는 한계가 있을 수 있다. 그럼에도 불구하고, 본 연구를 통해 압축된 다시점 데이터셋이 시점 생성 모델 성능에 미치는 영향을 분석할 수 있는 가능성을 확인하였으며, 향후 다양한 데이터셋과 추가 실험을 통해 이를 더욱 검증할 필요가 있다.

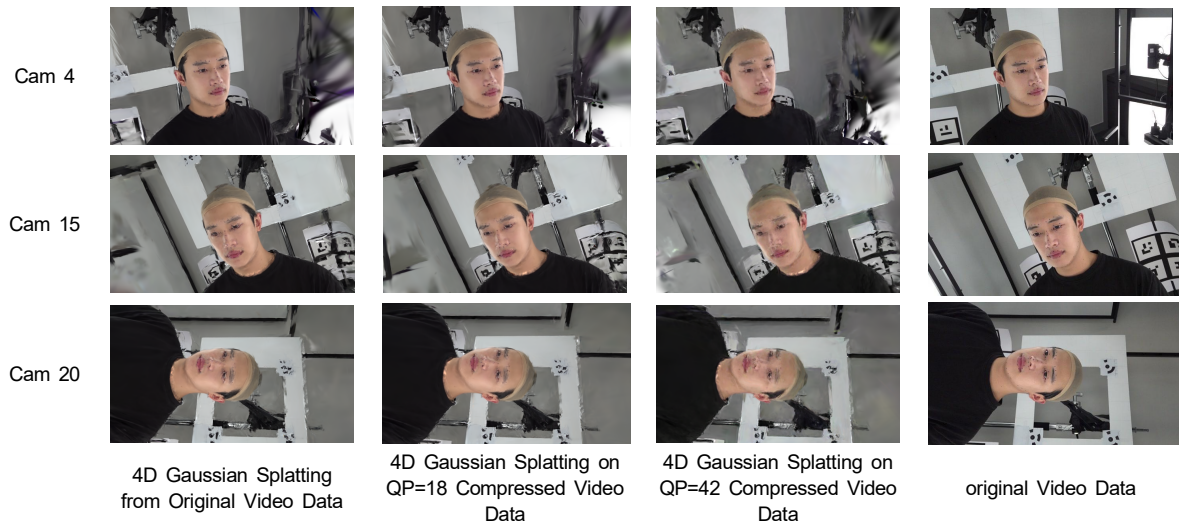


그림 8. 가장 왼쪽부터 원본 영상과 MV-HEVC로 압축(QP=18, 42)된 다시점 영상을 각각 4D 가우시안 스플래팅 학습 후 새로운 세 시점(cam4, cam15, cam20)에 대해 추론한 영상과 원본 영상의 예

Fig. 8. From the left, we present examples of the original video and the multi-view videos compressed with MV-HEVC (QP=18, QP=42). Each was used to train a 4D Gaussian Splatting model, and the inferred results for three novel viewpoints (cam4, cam15, cam20) are shown alongside the original video

표 4. 압축 강도에 따른 딥러닝 모델의 시점 성능 비교

Table 4. Comparison of deep learning performance at different compression levels

	QP	Original	18	24	30	36	42
Cam 4	PSNR	12.25	12.24	12.21	12.72	12.52	11.95
	SSIM	0.712	0.707	0.7	0.716	0.722	0.702
Cam 15	PSNR	18.15	18.07	18.89	18.53	18.8	18.9
	SSIM	0.747	0.746	0.758	0.755	0.756	0.742
Cam 20	PSNR	23.97	24.2	23.9	24.22	23.92	23.8
	SSIM	0.863	0.867	0.857	0.864	0.857	0.848

V. 결 론

본 논문에서는 MV-HEVC를 활용하여 다중 시점 영상을 압축하고, 이를 학습 데이터로 사용한 딥러닝 기반 시점 생성 모델의 성능에 미치는 영향을 분석하였다. 실험을 위해 21개의 시점을 세 개씩 그룹화하여 압축 및 복호화를 수행하고, 이를 4D 가우시안 스플래팅 모델의 학습 데이터로 활용하였다. 다양한 압축 강도(QP 값)를 적용하여 실험을 진행한 결과, 압축 강도가 증가할수록 PSNR이 감소하는 경향을 보였지만, 딥러닝 모델의 최종 시점 생성 성능에는 큰 영향을 미치지 않음을 확인하였다.

또한, 시점 그룹화 방식에 따라 실험을 진행한 결과, 카페

라 위치가 시각적으로 유사한 시점을 기준으로 그룹화한 경우와 COLMAP을 활용하여 자동으로 그룹화한 경우 모두 유사한 결과를 보였다. 두 실험에서 압축된 데이터셋을 활용하더라도 시점 생성 모델의 성능이 유지되었으며, 높은 압축률을 적용한 데이터셋도 학습 데이터로 활용 가능함을 확인하였다. 이는 다중 시점 영상 데이터를 저장 및 관리할 때 높은 압축을 적용하더라도, 새로운 시점 생성 모델의 성능 유지에는 문제가 없음을 의미한다.

다만, 본 연구는 하나의 씬(Scene)에 대해 촬영된 단일 데이터셋을 기반으로 실험이 진행되었기 때문에, 다양한 환경에서의 일반화 가능성에 대한 추가적인 검증이 필요하다. 이러한 한계에도 불구하고, 본 연구를 통해 압축된 다시

점 데이터셋이 시점 생성 모델의 성능에 미치는 영향을 정량적으로 분석할 가능성을 확인하였으며, 높은 압축률을 적용하더라도 학습 및 추론 성능이 유지될 수 있음을 실험적으로 입증하였다. 이를 통해 향후 보다 다양한 데이터셋을 활용한 추가 연구를 진행함으로써, 대규모 다시점 영상 데이터셋을 보다 효율적으로 저장 및 활용할 수 있는 방향을 제시할 수 있을 것으로 기대된다.

참 고 문 헌 (References)

- [1] B. Hyewon and P. Jimin, "The resurrection of the metaverse with AI technology," Chosun Ilbo, Jan. 10, 2025. https://www.chosun.com/economy/tech_it/2025/01/10/2J2I5ICHJPJDUTD72I3W7HUULEA [Accessed: Jan. 10, 2025].
- [2] Wang, Y., Huang, Z., Zhu, H., Li, W., Cao, X., & Yang, R. (2020). Interactive Free-Viewpoint Video Generation. *Virtual Reality & Intelligent Hardware*, 2(3), 247 - 260.
doi: <https://doi.org/10.1016/j.vrih.2020.04.004>
- [3] Carballeira, P., Carmona, C., Díaz, C., Berjón, D., Corregidor, D., Cabrera, J., Morán, F., Doblado, C., Arnaldo, S., Martín, M. del M., & García, N. (2020). FVV Live: A Real-Time Free-Viewpoint Video System with Consumer Electronics Hardware. *arXiv:2007.00558*.
doi: <https://doi.org/10.1109/TMM.2021.3079711>
- [4] C. Fehn, "Depth-Image-Based Rendering (DIBR), Compression, and Transmission for a New Approach on 3D-TV," in *Proceedings of SPIE - Stereoscopic Displays and Virtual Reality Systems XI*, vol. 5291, 2004.
doi: <https://doi.org/10.1117/12.524762>
- [5] H. Shum and S. B. Kang, "Review of image-based rendering techniques," in *Proceedings of SPIE - Visual Communications and Image Processing 2000*, vol. 4067, May 2000.
doi: <https://doi.org/10.1117/12.386541>
- [6] M. Levoy and P. Hanrahan, "Light field rendering," in *Proceedings of the 23rd Annual Conference on Computer Graphics and Interactive Techniques (SIGGRAPH '96)*, Aug. 1996, pp. 31-42.
doi: <https://doi.org/10.1145/237170.237199>
- [7] T. Łuczyński, P. Łuczyński, L. Pehle, M. Wirsum, and A. Birk, "Model based design of a stereo vision system for intelligent deep-sea operations," *Measurement*, vol. 144, pp. 298-310, 2019.
doi: <https://doi.org/10.1016/j.measurement.2019.05.004>
- [8] Ben Mildenhall, Pratul P Srinivasan, Matthew Tancik, Jonathan T Barron, Ravi Ramamoorthi, and Ren Ng. "Nerf: Representing scenes as neural radiance fields for view synthesis." *Communications of the ACM*, 65(1):99 - 106, 2021.
doi: <https://doi.org/10.48550/arXiv.2003.08934>
- [9] Bernhard Kerbl, Georgios Kopanas, Thomas Leimkuhler, and George Drettakis. "3D Gaussian Splatting for Real-Time Radiance Field Rendering." *ACM Transactions on Graphics (ToG)*, 42(4):1 - 14, 2023.
doi: <https://doi.org/10.48550/arXiv.2308.04079>
- [10] A. Pumarola, E. Corona, G. Pons-Moll, and F. Moreno-Noguer, "D-NeRF: Neural Radiance Fields for Dynamic Scenes," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021.
doi: <https://doi.org/10.48550/arXiv.2011.13961>
- [11] Z. Guo, W. Zhou, L. Li, M. Wang, and H. Li, "Motion-aware 3D Gaussian Splatting for Efficient Dynamic Scene Reconstruction," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 34, no. 3, pp. 789-803, Mar. 2024.
doi: <https://doi.org/10.48550/arXiv.2403.11447>
- [12] Guanjun Wu, Taoran Yi, Jiemin Fang, Lingxi Xie, Xiaopeng Zhang, Wei Wei, Wenyu Liu, Qi Tian, and Xinggang Wang, "4D Gaussian Splatting for Real-Time Dynamic Scene Rendering," School of CS, Huazhong University of Science and Technology; School of EIC, Huazhong University of Science and Technology; Huawei Inc., 2023.
doi: <https://doi.org/10.48550/arXiv.2310.08528>
- [13] Ang Cao and Justin Johnson. Hexplane: A fast representation for dynamic scenes. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 130 - 141, 2023.
doi: <https://doi.org/10.48550/arXiv.2301.09632>
- [14] Robert A Drebin, Loren Carpenter, and Pat Hanrahan. Volume rendering. *ACM Siggraph Computer Graphics*, 22(4): 65 - 74, 1988.
doi: <https://doi.org/10.1145/378456.378484>
- [15] M. Flierl and B. Girod, "Multiview video compression," *IEEE Signal Processing Magazine*, vol. 24, no. 6, pp. 66-76, Nov. 2007.
doi: <https://doi.org/10.1109/MSP.2007.905699>
- [16] P. Merkle, K. Muller, A. Smolic, and T. Wiegand, "Efficient Compression of Multi-View Video Exploiting Inter-View Dependencies Based on H.264/MPEG4-AVC," in *Proceedings of the IEEE International Conference on Image Processing (ICIP)*, Jul. 2006, pp. 201 - 204.
doi: <https://doi.org/10.1109/ICME.2006.262881>
- [17] G. J. Sullivan, J.-R. Ohm, W.-J. Han, and T. Wiegand, "Overview of the High Efficiency Video Coding (HEVC) standard," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 22, no. 12, pp. 1649-1668, Dec. 2012.
doi: <https://doi.org/10.1109/TCSVT.2012.2221191>
- [18] D. Flynn, D. Marpe, M. Naccari, T. Nguyen, C. Rosewarne, and K. Sharman, "Overview of the Range Extensions for the HEVC Standard: Tools, Profiles, and Performance," *IEEE Transactions on Circuits and Systems for Video Technology*, vol. 26, no. 1, pp. 4-19, Jan. 2016.
doi: <https://doi.org/10.1109/TCSVT.2015.2478707>
- [19] M. M. Hannuksela, Y. Yan, X. Huang, and H. Li, "Overview of the multiview high efficiency video coding (MV-HEVC) standard," in *Proceedings of the 2015 IEEE International Conference on Image Processing (ICIP)*, Quebec City, QC, Canada, Sep. 2015, pp. 1389-1393.
doi: <https://doi.org/10.1109/ICIP.2015.7351182>
- [20] J.-Y. Lee and J.-K. Han, "Fast disparity motion vector searching method for the MV-HEVC," *Journal of Broadcast Engineering*, vol.

22, no. 2, pp. 240-252, Mar. 2017.

doi: <https://doi.org/10.5909/JBE.2017.22.2.240>

- [21] C. Yang, P. An, L. Shen, R. Ma, and N. Feng, "Fast depth map coding based on virtual view quality," *Signal Processing: Image Communication*, vol. 47, pp. 183-192, Sep. 2016.
doi: <https://doi.org/10.1016/j.image.2016.05.015>
- [22] ISO/IEC JTC1/SC29/WG11, "Call for proposals on 3D video coding technology," document N12036, Geneva, CH, Mar. 2011
- [23] J. L. Schönberger and J.-M. Frahm, "Structure-from-Motion Revisited," in *Proceedings of the IEEE Conference on Computer Vision*

and Pattern Recognition (CVPR), 2016.

doi: <https://doi.org/10.1109/CVPR.2016.445>

- [24] M. She, F. Seeger, D. Nakath, and K. Köser, "Refractive COLMAP: Refractive Structure-from-Motion Revisited," in *Proceedings of the IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2024.
doi: <https://doi.org/10.1109/IROS58592.2024.10802043>
- [25] Y. Fu, S. Liu, A. Kulkarni, J. Kautz, A. A. Efros, and X. Wang, "COLMAP-Free 3D Gaussian Splatting,"
doi: <https://doi.org/10.48550/arXiv.2312.07504>

— 저 자 소 개 —

오 수 민



- 2025년 2월 : 광운대학교 전자재료공학과 졸업(공학사)
- 2025년 3월 ~ 현재 : 광운대학교 전자재료공학과 일반대학원(석사과정)
- ORCID : <https://orcid.org/0009-0008-9583-5399>
- 주관심분야 : 하드웨어 설계, 딥러닝, 영상 코덱

윤 승 환



- 2025년 2월 : 광운대학교 전자재료공학과 졸업(공학사)
- 2025년 3월 ~ 현재 : 광운대학교 전자재료공학과 일반대학원(석사과정)
- ORCID : <https://orcid.org/0009-0004-2066-6376>
- 주관심분야 : 하드웨어 설계, 딥러닝, 영상 코덱

김 정 우



- 2024년 2월 : 광운대학교 전자재료공학과 졸업(공학사)
- 2024년 3월 ~ 현재 : 광운대학교 전자재료공학과 일반대학원(석사과정)
- ORCID : <https://orcid.org/0009-0003-0913-0709>
- 주관심분야 : 하드웨어 설계, 딥러닝, 3D 영상 처리

이 학 범



- 2024년 2월 : 광운대학교 전자재료공학과학과 졸업(공학사)
- 2024년 3월 ~ 현재 : 광운대학교 전자재료공학과(석사과정)
- ORCID : <https://orcid.org/0000-0003-0721-4944>
- 주관심분야 : 다중 카메라 Calibration, 3D 휴먼 모션 재구성, NPU설계

저 자 소 개



서 영 호

- 1999년 2월 : 광운대학교 전자재료공학과 졸업(공학사)
- 2001년 2월 : 광운대학교 일반대학원 졸업(공학석사)
- 2004년 8월 : 광운대학교 일반대학원 졸업(공학박사)
- 2005년 9월 ~ 2008년 2월 : 한성대학교 조교수
- 2008년 3월 ~ 현재 : 광운대학교 전자재료공학과 교수
- ORCID : <https://orcid.org/0000-0003-1046-395X>
- 주관심분야 : 실감미디어, 2D/3D 영상 신호처리, SoC 설계, 디지털 홀로그램



김 동 욱

- 1983년 2월 : 한양대학교 전자공학과 졸업(공학사)
- 1985년 2월 : 한양대학교 공학석사
- 1991년 9월 : Georgia공과대학 전기공학과(공학박사)
- 1992년 3월 ~ 현재 : 광운대학교 전자재료공학과 정교수
- ORCID : <https://orcid.org/0000-0001-6106-9894>
- 주관심분야 : 3D 영상처리, 디지털 홀로그램, 디지털 VLSI Testability, VLSI CAD, DSP 설계, Wireless Communication