



(19) 대한민국특허청(KR)
(12) 공개특허공보(A)

(11) 공개번호 10-2020-0044435
(43) 공개일자 2020년04월29일

(51) 국제특허분류(Int. Cl.)
H04N 21/466 (2011.01) G06K 9/00 (2006.01)
G06N 99/00 (2019.01)
(52) CPC특허분류
H04N 21/4668 (2013.01)
G06K 9/00744 (2013.01)
(21) 출원번호 10-2018-0125197
(22) 출원일자 2018년10월19일
심사청구일자 2018년10월19일

(71) 출원인
인하대학교 산학협력단
인천광역시 미추홀구 인하로 100(용현동, 인하대 학교)
(72) 발명자
권장우
부산광역시 남구 오륙도로 85, 109동 2503호(용호 동, 오륙도 에스케이뷰 아파트)
설현우
경기도 고양시 일산동구 위시티4로 45, 412동 2501호(식사동, 위시티일산자이4단지아파트)
(74) 대리인
양성보

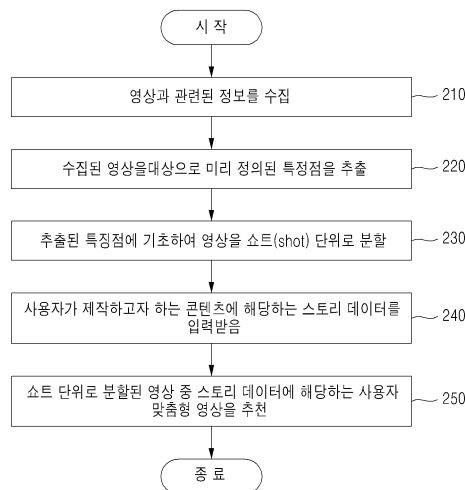
전체 청구항 수 : 총 8 항

(54) 발명의 명칭 **영상의 쇼트 분류를 이용한 사용자 맞춤형 영상 추천 시스템**

(57) 요약

영상의 쇼트 분류를 이용한 사용자 맞춤형 영상 추천 시스템이 개시된다. 사용자 맞춤형 영상 추천 방법에 있어서, 영상과 관련된 정보를 수집하는 단계, 수집된 영상을 대상으로 미리 정의된 특징점을 추출하는 단계, 추출된 상기 특징점에 기초하여 상기 영상을 쇼트(shot) 단위로 분할하는 단계, 사용자가 제작하고자 하는 콘텐츠에 해당하는 스토리 데이터를 입력받는 단계, 및 상기 쇼트 단위로 분할된 영상 중 상기 스토리 데이터에 해당하는 사용자 맞춤형 영상을 추천하는 단계를 포함할 수 있다.

대표도 - 도2



(52) CPC특허분류

G06N 20/00 (2019.01)

H04N 21/4665 (2013.01)

이 발명을 지원한 국가연구개발사업

과제고유번호 1711070451(세부과제번호: 2017-0-01642-002)

부처명 과학기술정보통신부

연구관리전문기관 정보통신기술진흥센터

연구사업명 정보통신기술인력양성(R&D, 정보화)

연구과제명 인공지능을 활용한 콘텐츠 창작 기술

기 여 율 1/1

주관기관 인하대학교 산학협력단

연구기간 2018.01.01 ~ 2018.12.31

명세서

청구범위

청구항 1

사용자 맞춤형 영상 추천 방법에 있어서,
영상과 관련된 정보를 수집하는 단계;
수집된 영상을 대상으로 미리 정의된 특징점을 추출하는 단계;
추출된 상기 특징점에 기초하여 상기 영상을 쇼트(shot) 단위로 분할하는 단계;
사용자가 제작하고자 하는 콘텐츠에 해당하는 스토리 데이터를 입력받는 단계; 및
상기 쇼트 단위로 분할된 영상 중 상기 스토리 데이터에 해당하는 사용자 맞춤형 영상을 추천하는 단계를 포함하는 사용자 맞춤형 영상 추천 방법.

청구항 2

제1항에 있어서,
상기 사용자 맞춤형 영상을 추천하는 단계는,
상기 스토리 데이터를 기반으로 하는 스크립트 데이터에서 미리 지정된 특징점을 추출하는 단계;
추출한 상기 특징점에 기초하여 상기 스토리 데이터를 쇼트(shot) 단위로 분해하는 단계; 및
상기 쇼트 단위로 분할된 영상 중 상기 쇼트 단위로 분해된 스토리 데이터에 해당하는 사용자 맞춤형 영상을 결정하는 단계를 포함하는 사용자 맞춤형 영상 추천 방법.

청구항 3

제1항에 있어서,
상기 영상을 쇼트(shot) 단위로 분할하는 단계는,
상기 수집된 영상을 대상으로 추출된 장소, 시간, 인물, 감정, 행동 중 적어도 하나의 특징점을 상기 쇼트 단위로 분할된 해당 영상과 연관하여 저장하는 단계를 포함하는 사용자 맞춤형 영상 추천 방법.

청구항 4

제1항에 있어서,
상기 영상을 쇼트(shot) 단위로 분할하는 단계는,
상기 영상을 중간에 끊지 않고 촬영한 하나의 연속적인 화면을 나타내는 상기 쇼트, 상기 쇼트가 속하는 신(scene), 상기 신이 속하는 시퀀스(sequence) 단위로 구분하여 저장하는 단계를 포함하는 사용자 맞춤형 영상 추천 방법.

청구항 5

제1항에 있어서,
상기 수집된 영상을 대상으로 미리 정의된 특징점을 추출하는 단계는,
훈련 데이터를 대상으로 기계 학습(machine learning)에 기초하여 생성된 학습 모델에 상기 수집된 영상을 입력

값으로 설정하고, 상기 학습 모델의 출력값으로 상기 특징점을 추출하는 것을 특징으로 하는 사용자 맞춤형 영상 추천 방법.

청구항 6

제1항에 있어서,
상기 사용자 맞춤형 영상을 추천하는 단계는,
입력된 상기 스토리 데이터에 해당하는 입력 스크립트를 대상으로 자연어 처리 기반으로 데이터 분석을 수행하는 단계;
분석된 상기 데이터를 대상으로, 인물, 감정, 행동, 장소와 관련된 특징점을 추출하는 단계;
추출된 상기 특징점에 기초하여 상기 입력 스크립트를 쇼트 단위로 분해하는 단계;
상기 쇼트 단위로 분해된 입력 스크립트를 대상으로, 유사한 스토리끼리 하나의 그룹에 속하도록 그룹을 형성하는 단계; 및
형성된 상기 그룹을 대상으로, 사용자 맞춤형 영상을 추천하는 단계를 포함하는 사용자 맞춤형 영상 추천 방법.

청구항 7

제1항에 있어서,
상기 스토리 데이터를 입력받는 단계는,
상기 쇼트 단위로 분할된 영상 중 상기 스토리 데이터와 관련된 등장인물, 장소, 등장인물의 감정 및 행동 중 어느 하나를 기준으로 상기 사용자 맞춤형 영상을 추천하기 위한 우선순위를 입력받는 단계를 포함하는 사용자 맞춤형 영상 추천 방법.

청구항 8

영상과 관련된 정보를 수집하는 정보 수집부;
수집된 영상을 대상으로 미리 정의된 특징점을 추출하는 특징점 추출부;
추출된 상기 특징점에 기초하여 상기 영상을 쇼트(shot) 단위로 분할하는 분할부;
사용자가 제작하고자 하는 콘텐츠에 해당하는 스토리 데이터를 입력받는 정보 입력부; 및
상기 쇼트 단위로 분할된 영상 중 상기 스토리 데이터에 해당하는 사용자 맞춤형 영상을 추천하는 영상 추천부를 포함하는 사용자 맞춤형 영상 추천 시스템.

발명의 설명

기술 분야

[0001] 아래의 설명은 영상을 특정 단위로 분할하여 사용자의 스토리(story) 제작에 도움을 주는 영상 추천 기술에 관한 것이다.

배경 기술

[0002]영상의 콘텐츠가 다양해지면서 시청자가 기존 콘텐츠를 변경하여 새로운 내용이나 결말을 만들어내기 시작하였다.

[0003]그러나, 콘텐츠의 양이 방대하여 사용자가 원하는 콘텐츠를 검색하는 시간, 노력, 비용이 증가하게 되었다. 이에 따라, 사용자가 원하는 영상 콘텐츠를 제작할 수 있도록 영상을 검색하거나 추천하기 위한 다양한 방법들이 연구되어 있다.

[0004]새로운 콘텐츠를 제작하기 위해서는 영상을 분할하여 편집해야 한다. 영상의 편집으로 다양한 영상 데이터를

가지고 있을 수 있으며, 이러한 데이터를 활용하여 새로운 영상을 만드는데 도움을 준다.

[0005] 추천시스템 정보 필터링 기술은 특정 사용자가 관심을 가질만한 정보를 추천하는 기술로서, 접근방식에 따라 협업 필터링(collaborative filtering)과 콘텐츠 기반 필터링(content-based filtering)이 이용된다.

[0006] 한국등록특허 제10-1680367호는 시뮬레이터 카메라와 가상카메라에 의한 촬영 영상 정보(CG 적용 영상)이 영상 프레임 쇼트(shot) 형태로 실시간으로 시뮬레이터 카메라로 전송 및 출력되어 CG 적용 영상에서의 시뮬레이터 카메라 움직임이 촬영 현장에서 실시간으로 확인되도록 하는 시뮬레이터 카메라와 가상카메라 동기화에 의한 CG 적용 영상제작시스템이 개시되어 있다.

발명의 내용

해결하려는 과제

[0007] 기존에 미리 제작된 콘텐츠(contents)를 기반으로 사용자가 원하는 새로운 콘텐츠를 제작하기 위하여 콘텐츠의 영상을 개별적으로 분류하고, 사용자 스토리(story) 흐름을 고려하여 사용자 맞춤형 영상을 추천하기 위한 것이다.

[0008] 기존에 미리 제작된 콘텐츠(contents)를 기반으로 사용자가 원하는 새로운 콘텐츠를 제작하기 위하여 콘텐츠의 영상을 개별적으로 분류하고, 사용자 스토리(story) 흐름을 고려하여 사용자 맞춤형 영상을 추천하기 위한 것이다.

과제의 해결 수단

[0009] 사용자 맞춤형 영상 추천 방법에 있어서, 영상과 관련된 정보를 수집하는 단계, 수집된 영상을 대상으로 미리 정의된 특징점을 추출하는 단계, 추출된 상기 특징점에 기초하여 상기 영상을 쇼트(shot) 단위로 분할하는 단계, 사용자가 제작하고자 하는 콘텐츠에 해당하는 스토리 데이터를 입력받는 단계, 및 상기 쇼트 단위로 분할된 영상 중 상기 스토리 데이터에 해당하는 사용자 맞춤형 영상을 추천하는 단계를 포함할 수 있다.

[0010] 일측면에 따르면, 상기 사용자 맞춤형 영상을 추천하는 단계는, 상기 스토리 데이터를 기반으로 하는 스크립트 데이터에서 미리 지정된 특징점을 추출하는 단계, 추출한 상기 특징점에 기초하여 상기 스토리 데이터를 쇼트(shot) 단위로 분해하는 단계, 및 상기 쇼트 단위로 분할된 영상 중 상기 쇼트 단위로 분해된 스토리 데이터에 해당하는 사용자 맞춤형 영상을 결정하는 단계를 포함할 수 있다.

[0011] 다른 측면에 따르면, 상기 영상을 쇼트(shot) 단위로 분할하는 단계는, 상기 수집된 영상을 대상으로 추출된 장소, 시간, 인물, 감정, 행동 중 적어도 하나의 특징점을 상기 쇼트 단위로 분할된 해당 영상과 연관하여 저장하는 단계를 포함할 수 있다.

[0012] 또 다른 측면에 따르면, 상기 영상을 쇼트(shot) 단위로 분할하는 단계는, 상기 영상을 중간에 끊지 않고 촬영한 하나의 연속적인 화면을 나타내는 상기 쇼트, 상기 쇼트가 속하는 신(scene), 상기 신이 속하는 시퀀스(sequence) 단위로 구분하여 저장하는 단계를 포함할 수 있다.

[0013] 또 다른 측면에 따르면, 상기 수집된 영상을 대상으로 미리 정의된 특징점을 추출하는 단계는, 훈련 데이터를 대상으로 기계 학습(machine learning)에 기초하여 생성된 학습 모델에 상기 수집된 영상을 입력값으로 설정하고, 상기 학습 모델의 출력값으로 상기 특징점을 추출할 수 있다.

[0014] 또 다른 측면에 따르면, 상기 사용자 맞춤형 영상을 추천하는 단계는, 입력된 상기 스토리 데이터에 해당하는 입력 스크립트를 대상으로 자연어 처리 기반으로 데이터 분석을 수행하는 단계, 분석된 상기 데이터를 대상으로, 인물, 감정, 행동, 장소와 관련된 특징점을 추출하는 단계, 추출된 상기 특징점에 기초하여 상기 입력 스크립트를 쇼트 단위로 분해하는 단계, 상기 쇼트 단위로 분해된 입력 스크립트를 대상으로, 유사한 스토리끼리 하나의 그룹에 속하도록 그룹을 형성하는 단계, 및 형성된 상기 그룹을 대상으로, 사용자 맞춤형 영상을 추천하는 단계를 포함할 수 있다.

[0015] 또 다른 측면에 따르면, 상기 스토리 데이터를 입력받는 단계는, 상기 쇼트 단위로 분할된 영상 중 상기 스토리 데이터와 관련된 등장인물, 장소, 등장인물의 감정 및 행동 중 어느 하나를 기준으로 상기 사용자 맞춤형 영상을 추천하기 위한 우선순위를 입력받는 단계를 포함할 수 있다.

[0016] 사용자 맞춤형 영상 추천 시스템은, 영상과 관련된 정보를 수집하는 정보 수집부, 수집된 영상을 대상으로 미리

정의된 특징점을 추출하는 특징점 추출부, 추출된 상기 특징점에 기초하여 상기 영상을 쇼트(shot) 단위로 분할하는 분할부, 사용자가 제작하고자 하는 콘텐츠에 해당하는 스토리 데이터 입력받는 정보 입력부, 및 상기 쇼트 단위로 분할된 영상 중 상기 스토리 데이터에 해당하는 사용자 맞춤형 영상을 추천하는 영상 추천부를 포함할 수 있다.

발명의 효과

[0017] 본 발명의 실시예에 따르면, 기존에 미리 제작된 콘텐츠(contents)를 기반으로 사용자가 원하는 새로운 콘텐츠를 제작하기 위하여 콘텐츠의 영상을 쇼트(shot) 단위로 분류하고, 사용자 스토리(story) 흐름을 고려하여 쇼트 단위로 분류된 영상 중 스토리에 해당하는 사용자 맞춤형 영상을 추천할 수 있다. 즉, 사용자가 기존에 미리 제작된 방대한 양의 콘텐츠 영상을 검색하지 않더라도, 입력한 스토리에 해당하는 영상을 추천함으로써, 콘텐츠 검색 시간, 콘텐츠 제작에 소요되는 노력 및 비용을 감소시킬 수 있다.

도면의 간단한 설명

[0018] 도 1은 본 발명의 일실시에 따른 영상을 구분하는 시퀀스, 신, 쇼트 구조를 나타내는 도면이다.
 도 2는 본 발명의 일실시에 있어서, 사용자 맞춤형 영상 추천 방법을 도시한 흐름도이다.
 도 3은 본 발명의 일 실시예에 있어서, 사용자 맞춤형 영상 추천 시스템의 내부구성을 도시한 블록도이다.
 도 4는 본 발명의 일실시에 있어서, 쇼트 영상에서 추출된 특징점들을 도시한 도면이다.
 도 5는 본 발명의 일실시에 있어서, 영상의 스크립트(script)를 도시한 도면이다.
 도 6은 본 발명의 일실시에 있어서, 영상을 분할하는 세부 동작을 도시한 흐름도이다.
 도 7은 본 발명의 일실시에 있어서, 영상을 추천하는 세부 동작을 도시한 흐름도이다.

발명을 실시하기 위한 구체적인 내용

[0019] 이하, 본 발명의 실시예를 첨부된 도면을 참조하여 상세하게 설명한다.

[0021] 본 실시예들은 영상의 쇼트(shot) 분류를 이용한 사용자 맞춤형 영상 추천 시스템 및 방법에 관한 것으로, 웹/모바일 웹 등 다양한 경로를 통해 수집된 콘텐츠의 영상을 대상으로 시퀀스(sequence), 신(scene), 및 쇼트(shot) 단위로 구분하고, 사용자가 제작하고자 하는 콘텐츠의 스토리 흐름을 고려하여 상기 구분된 영상을 추천하는 기술에 관한 것이다. 특히, 사용자가 제작하고자 하는 콘텐츠의 스토리에 해당하는 영상을 추천하기 위해, 상기 영상 및 스토리 데이터에서 특징점을 추출하고, 추출된 특징점에 기초하여 사용자 맞춤형 영상을 추천하는 기술에 관한 것이다.

[0022] 본 실시예들에서, "쇼트(shot)"는 영상 제작 시 중간에 끊지 않고 촬영한 하나의 연속적인 화면을 나타내는 것으로서, "레디~ 액션!"으로 촬영이 시작되어 "컷"하고 종료될 때 까지를 나타낼 수 있다. 쇼트는 감독의 "레디~ 액션!"과 함께 카메라를 켤 수는 없으므로, 감독이 "레디~"하는 중에 카메라는 켜 놓고 촬영할 준비를 하고, 중간에 NG 등과 같이 불필요한 부분을 제거하고 실제 영상에 반영되는 부분을 나타낼 수 있다. 쇼트는 영화의 가장 기본적인 단위로서, 통계적으로 하나의 쇼트는 10 내지 15초 정도이고, 보통 한 편의 영화는 900여 개의 쇼트로 구성될 수 있다.

[0023] 본 실시예들에서, "신(scene)"은 동일한 시간이나 장소에서 이루어지는 하나의 사건, 즉, 장면을 나타낼 수 있다. 일반적으로, 하나의 신(scene)에는 여러 쇼트가 포함되지만, 1개의 쇼트가 1개의 신을 구성할 수도 있다. 러닝 타임(running time)이 90분 정도인 영화의 경우, 120개 정도의 신으로 구성될 수 있다.

[0024] 본 실시예들에서, "시퀀스(sequence)"는 콘텐츠의 스토리(story) 흐름을 기준으로 나눈 신(scene)들의 묶음을 나타낼 수 있다. 예컨대, 첫만남까지, 만남에서 사랑하기까지, 사랑한 후 이별까지, 이별 후 재회까지 등과 같이 스토리 흐름에 따라 구분될 수 있다. 신(scene)은 시간과 장소의 제약을 받지만 시퀀스는 시간과 장소에 제약을 받지 않으므로, 스토리 흐름은 사용자에게 따라 다양하게 나올 수 있다. 즉, 시퀀스는 상황의 시작과 끝을 정하는 것이며, 영화 등의 콘텐츠를 시청하는 사용자마다 특정 상황의 시작과 끝이 다를 수 있기 때문에 스토리 흐름은 동일 콘텐츠를 대상으로 사용자에게 따라 다양할 수 있다.

- [0025] 본 실시예들에서, "사용자 단말"은 스마트폰(smart phone), 휴대폰, 내비게이션, 컴퓨터, 노트북, 디지털방송용 단말, 태블릿 PC, 웨어러블 디바이스(wearable device) 등의 전자기기를 나타낼 수 있다.
- [0027] 도 1은 본 발명의 일실시에 따른 영상을 구분하는 시퀀스, 신, 쇼트 구조를 나타내는 도면이다.
- [0028] 도 1을 참고하면, 콘텐츠의 영상은 스토리(story)의 흐름에 따라 크게 시퀀스(sequence, 100)로 구분될 수 있다.
- [0029] 시퀀스(100)는 시간 및 장소 등의 특징점에 따라 복수개의 신(scene, 110)으로 구분될 수 있으며, 신(110)은 다시 복수개의 쇼트(shot, 101, 102, 103, 104)로 구분될 수 있다.
- [0030] 일례로, 두 남녀가 카페에서 앉아서 대화하는 모습의 경우, 남자가 말할 때는 남자의 얼굴이 화면에 디스플레이 되고, 여자가 말할 때는 여자의 얼굴이 디스플레이되다가 둘이 마주보고 대화하는 모습이 디스플레이될 수 있으며, 남자화면, 여자화면, 남녀화면 각각이 쇼트에 해당할 수 있다. 즉, 두 남녀가 카페에서 앉아서 대화하는 신(scene), 즉, 장면은 3개의 쇼트(shot, 남자화면, 여자화면, 남녀화면)로 구성될 수 있다.
- [0031] 그리고, 복수개의 신(scene)들이 모여 하나의 스토리 흐름을 이루면, 하나의 시퀀스(sequence, 100)가 구성될 수 있다.
- [0033] 도 2는 본 발명의 일실시에 있어서, 사용자 맞춤형 영상 추천 방법을 도시한 흐름도이고, 도 3은 본 발명의 일 실시예에 있어서, 사용자 맞춤형 영상 추천 시스템의 내부구성을 도시한 블록도이다.
- [0034] 도 3을 참고하면, 사용자 맞춤형 영상 추천 시스템(300)은 정보 수집부(310), 특징점 추출부(320), 분할부(330), 정보 입력부(340), 및 영상 추천부(350)를 포함할 수 있다. 그리고, 도 2의 각 단계들(210 내지 250 단계)은 도 3의 구성요소인 정보 수집부(310), 특징점 추출부(320), 분할부(330), 정보 입력부(340), 및 영상 추천부(350)에 의해 수행될 수 있다.
- [0035] 210 단계에서, 정보 수집부(310)는 다양한 콘텐츠 영상과 관련된 정보를 수집할 수 있다.
- [0036] 일례로, 정보 수집부(310)는 웹/모바일 웹 등을 통해 제공되는 다양한 콘텐츠(예컨대, 영화, 뮤직비디오, 드라마 등)의 영상과 관련된 정보를 수집할 수 있다. 예컨대, 정보 수집부(310)는 비디오 영상, 영상에 해당하는 대사, 자막, 대본 등의 스크립트(script) 등을 수집할 수 있다. 이외에, 해당 영상의 음향 정보, 장소 정보, 등장 인물 정보 등을 수집할 수 있다.
- [0037] 220 단계에서, 특징점 추출부(320)는 수집된 영상을 대상으로, 미리 정의된 특징점을 추출할 수 있다.
- [0038] 일례로, 특징점 추출부(320)는 영상에서 장소, 시간, 인물, 감정, 행동, 대사 등의 미리 정의된 특징점들을 추출할 수 있다. 이때, 특징점 추출부(320)는 영상처리(image processing) 및 컴퓨터 비전(computer vision)을 이용한 기계 학습(machine learning)에 기초하여 훈련 데이터 세트를 훈련시켜 생성한 학습 모델을 이용하여 특징점들을 추출할 수 있다. 예컨대, 특징점 추출부(320)는 수집된 영상을 학습 모델의 입력값으로 설정하고, 입력된 영상에서 예측된 등장인물(즉, 액터), 등장인물의 기쁨, 슬픔, 괴로움, 화남 등의 감정, 행동 중 적어도 하나의 특징점이 상기 학습 모델의 출력으로 출력될 수 있다. 이때, 등장인물의 표정을 기반으로 등장인물의 감정 상태(즉, 기쁨, 슬픔, 괴로움, 화남 등)의 특징점이 추출될 수 있으며, 등장인물의 몸에 해당하는 부분의 행동 분석을 기반으로 등장인물의 행동(예컨대, 오른팔을 들고 좌우로 흔들며 인사하는 동작인지, 몸앞으로 팔을 내밀고 위아래로 흔드는 악수 동작인지 등의 행동)을 나타내는 특징점이 추출될 수 있다. 여기서, 학습 모델의 출력값으로, 시간적, 공간적 배경을 포함하고 있는 장소 역시 특징점으로 출력될 수 있다.
- [0039] 230 단계에서, 분할부(330)는 추출된 특징점에 기초하여 영상을 쇼트(shot) 단위로 분할할 수 있다.
- [0040] 일례로, 분할부(330)는 수집된 영상들 각각을 대상으로 시퀀스(sequence), 신(scene), 쇼트(shot)로 분해할 수 있으며, 분해된 영상들 각각을 각 영상에서 추출된 특징점에 기초하여 쇼트 단위로 구분하여 데이터베이스에 저장 및 유지할 수 있다. 예컨대, 도 4와 같이 서로 다른 콘텐츠의 영상을 대상으로 시퀀스, 신, 쇼트로 구분되고, 각 쇼트 단위의 영상에서 추출된 특징점들이 서로 연관될 수 있다. 그러면, 분할부(330)는 서로 유사한 특징점에 해당하는 영상들을 그룹화하여 데이터베이스에 저장 및 유지할 수 있다. 예컨대, 저녁 시간, 집, 가족, 식사라는 특징점들에 해당하는 쇼트 단위 영상(예컨대, 서로 다른 콘텐츠를 대상으로 구분된 쇼트 영상 중 엄마

화면, 아빠화면, 아이화면, 가족화면 등)을 그룹화하여 해당 특징점과 연관하여 저장 및 유지할 수 있다.

- [0041] 240 단계에서, 정보 입력부(340)는 사용자가 제작하고자 하는 콘텐츠에 해당하는 스토리 데이터를 입력받을 수 있다.
- [0042] 예컨대, 정보 입력부(340)는 사용자 단말에 설치된 어플리케이션, 또는 사용자 단말을 통해 접속한 웹사이트의 웹페이지를 통해 사용자가 제작하고자 하는 콘텐츠관련 스토리를 입력받을 수 있다. 그러면, 입력된 스토리 데이터와 사용자 단말의 식별자 정보(예컨대, 사용자 ID 등)이 데이터베이스에 연관하여 저장 및 유지될 수 있다. 이때, 스토리 데이터를 처음부터 끝까지 한번에 입력하지 않고, 접속 종료 후 이후 다시 접속하여 기존에 입력한 스토리 데이터에 이어 스토리를 입력받을 수도 있다.
- [0043] 250 단계에서, 영상 추천부(350)는 쇼트 단위로 분할된 영상 중 스토리 데이터에 해당하는 사용자 맞춤형 영상을 추천할 수 있다.
- [0044] 일례로, 영상 추천부(350)는 입력된 스토리 데이터를 기반으로 하는 스크립트 데이터에서 미리 지정된 특징점을 추출할 수 있다. 그리고, 추출한 특징점에 기초하여 스토리 데이터를 쇼트(shot) 단위로 분해할 수 있다. 영상 추천부(350)는 쇼트 단위로 분할된 영상 중 쇼트 단위로 분해된 스토리 데이터에 해당하는 사용자 맞춤형 영상을 결정할 수 있다. 예컨대, 영상 추천부(350)는 쇼트 단위로 분할된 영상에 해당하는 특징점과 쇼트 단위로 분해된 스토리 데이터에 해당하는 특징점에 기초하여 유사도가 높은 일정 개수의 영상을 결정할 수 있다. 유사도가 높은 순서로 사용자 단말에 설치된 어플리케이션 등을 통해 입력된 스토리에 해당하는 쇼트 단위 영상을 제공할 수 있다. 즉, 영상 추천부(350)는 유사도가 높은 일정 개수의 후보 영상을 추천을 위해 제공할 수 있다. 이외에, 영상 추천부(350)는 쇼트 단위로 분할된 영상에 해당하는 쇼트 단위 스크립트를 대상으로 상기 쇼트 단위로 분해된 스토리 데이터와 일치하는 스크립트의 영상을 추천 영상으로 결정할 수도 있다.
- [0046] 도 5는 본 발명의 일실시예에 있어서, 영상의 스크립트(script)를 도시한 도면이다.
- [0047] 도 5에서, 영상의 스크립트는 도 1과 같이 영상을 시퀀스, 신, 쇼트로 분해하기 위한 지표로서, 장소, 시간, 인물, 감정, 행동, 대사 등의 정보를 포함할 수 있다.
- [0048] 영상과 관련된 정보를 수집 시 비디오 영상 및 해당 영상의 스크립트가 함께 수집될 수 있다. 스크립트를 대상으로, 장소, 시간, 인물, 감정, 행동, 대사 등의 특징점들이 추출될 수 있다. 그러면, 스크립트에 해당하는 영상을 대상으로 추출된 특징점 및 상기 스크립트에서 추출된 특징점에 기초하여, 스크립트는 쇼트 단위로 분할될 수 있다.
- [0049] 일례로, 쇼트 단위로 분할된 영상, 즉, 쇼트 영상에 해당하는 특징점과 일치하는 특징점에 해당하는 스크립트 부분이 쇼트 단위로 분할될 수 있다. 예컨대, 남자화면 쇼트 영상의 경우, 인물: 남자, 장소: 커피숍, 시간: 저녁, 감정: 즐거움, 행동: 커피잔을 들어올림 등의 특징점을 포함하는 경우, 스크립트에서 추출된 특징점 중 인물: 남자, 장소: 커피숍, 시간: 저녁, 감정: 즐거움, 행동: 커피잔을 들어올림 등의 특징점에 해당하는 스크립트 부분이 상기 남자 화면에 대응하도록 쇼트 단위로 분할될 수 있다. 동일한 방법으로, 여자 화면, 남녀화면 등의 쇼트 단위로 스크립트가 분할될 수 있으며, 분할된 쇼트 단위 스크립트, 쇼트 단위 영상, 해당하는 특징점이 연관하여 저장 및 유지될 수 있다. 이후, 스토리 데이터가 입력되면, 입력된 스토리 데이터에 해당하는 쇼트 단위 영상을 추천하기 위해 이용될 수 있다.
- [0050] 도 5를 참고하면, 호텔 안 계단/저녁 신(scene)은, 그래 화면, 조대리화면, 요르단 남자화면, 그래조대리요르단 남자 화면의 쇼트 영상 등으로 구성될 수 있으며, 해당 스크립트 역시 쇼트 단위로 구분될 수 있다.
- [0052] 도 6은 본 발명의 일실시예에 있어서, 영상을 분할하는 세부 동작을 도시한 흐름도이다.
- [0053] 610 단계를 참고하면, 정보 수집부(310)는 영상관련 정보를 수집하면, 수집된 영상과 관련된 스토리, 인물, 인물과의 관계, 줄거리들의 데이터 등을 함께 수집할 수 있으며, 수집된 데이터를 미리 정의된 분석 알고리즘에 기초하여 분석할 수 있다.
- [0054] 620 단계에서, 특징점 추출부(320)는 해당 영상에 포함된 배경 및 인물을 대상으로 특징점을 추출할 수 있다.

- [0055] 예컨대, 학습 모델에 기초하여 배경 영상으로부터 장소, 시간(아침, 새벽, 저녁, 오후 등) 등의 특징점을 추출하고, 인물 영상으로부터 감정, 행동 등의 특징점 등을 추출할 수 있다. 이때, 영상만으로 인물의 감정이나 속마음을 정확하게 유추하기 어려우므로, 도 5와 같이 영상의 스크립트(script)에 포함된 내용이 세부적으로 분석하여 데이터로 정형화될 수 있으며, 스크립트는 시퀀스, 신, 쇼트로 구분될 수 있다. 그러면, 쇼트 단위로 구분된 스크립트는 입력된 스토리 데이터(즉, 스토리의 흐름)에 따라 유사한 영상을 추천하기 위해 이용될 수 있다.
- [0056] 630 단계에서, 분할부(330)는 영상을 시퀀스, 신, 쇼트 단위로 분할하여 저장할 수 있다. 이때, 분할부(330)는 스크립트를 기반으로 이전 프레임(previous frame)과 현재 프레임(frame)을 비교하여 영상을 병합 및 분할할 수 있다. 그러면, 쇼트 단위로 분할된 영상은 정형화되어 데이터베이스에 저장 및 유지될 수 있다.
- [0057] 이처럼, 수집된 영상 및 해당 영상의 스크립트가 존재하는 경우, 스크립트를 분석하여 사용자가 어플리케이션을 통해 입력한 스토리에 해당하는 적어도 하나의 후보 영상을 추천할 수 있으며, 스크립트가 존재하지 않는 경우에도 영상의 쇼트 자체를 분석하여 입력된 스토리에 해당하는 적어도 하나의 후보 영상을 사용자에게 추천할 수도 있다. 예컨대, 영상을 신 및 쇼트 단위로 구분하고, 구분된 쇼트 단위의 영상을 대상으로 등장인물에 해당하는 특징점을 추출하고, 등장인물의 감정 및 행동에 해당하는 특징점을 추출하고, 해당 영상의 배경을 나타내는 특징점을 추출함으로써, 상기 후보 영상을 추천할 수 있다. 예컨대, 쇼트 단위의 영상을 대상으로 얼굴에 해당하는 영역을 분석하여, 장그래, 오상식 등의 등장인물에 해당하는 특징점을 추출하여, 등장인물이 인식될 수 있다. 그리고, 얼굴 영역에 해당하는 표정을 기반으로 눈매, 입꼬리 등을 분석하여 특정 등장인물의 감정 상태(기쁨, 화남 등)를 인식할 수 있으며, 등장인물의 몸에 해당하는 영역의 영상 분석을 통해 행동이 인식될 수 있다. 그리고, 데이터베이스에 저장된 쇼트 단위의 영상들 중, 장그래 및 오상식이 포함되되, 인사하는 행동 및 배경이 사무실에 해당하는 적어도 하나의 쇼트 영상이 후보 영상으로 추천될 수 있다.
- [0059] 도 7은 본 발명의 일실시예에 있어서, 영상을 추천하는 세부 동작을 도시한 흐름도이다.
- [0060] 도 7을 참고하면, 710 단계에서, 영상 추천부(350)는 스토리 데이터들에 대한 데이터 분석을 수행할 수 있다.
- [0061] 일례로, 정보 입력부(340)는 스토리 데이터를 사용자 단말의 어플리케이션 등을 통해 입력받을 수 있다. 즉, 사용자가 제작하고자 하는 콘텐츠의 스토리 데이터를 입력받을 수 있다. 그러면, 영상 추천부(350)는 입력된 스토리 데이터를 대상으로, 자연어 처리를 통해 데이터를 분석 및 추출할 수 있다.
- [0062] 이외에, 정보 입력부(340)는 쇼트 단위로 분할된 영상 중 상기 스토리 데이터와 관련된 등장인물, 장소, 등장인물의 감정 및 행동 중 어느 하나를 기준으로 사용자 맞춤형 영상을 추천하기 위한 우선순위를 입력받을 수도 있다. 일례로, 데이터베이스에 저장 및 유지된 쇼트 단위로 분할된 영상들 중 추출된 특징점을 기반으로 등장인물과 유사도가 높은 일정개수의 후보 영상, 인식된 장소와 유사도가 높은 일정개수의 후보 영상, 인식된 감정 상태 및 행동과 유사도가 높은 일정개수의 후보 영상을 선별하여 제공함에 있어서, 장소, 등장인물, 등장인물의 감정이나 행동 중 어떤 요소에 가중치를 적용할지에 대한 우선순위를 입력받을 수 있다. 예컨대, 장소에 우선순위가 1순위로 입력된 경우, 장소가 가장 유사한 일정 개수의 후보 영상을 먼저 추출한 후, 추출한 쇼트 단위의 후보 영상들을 대상으로, 등장인물, 감정이나 행동이 유사한 후보 영상들을 추천하기 위해 상기 우선순위가 이용될 수 있다. 여기서, 우선순위는 장소가 1순위, 등장인물이 2순위, 감정상태가 3순위 등으로 지정되어 서로 다른 가중치가 적용되어 추천하고자 후보 영상들을 검색하기 위해 입력받을 수도 있고, 어느 하나의 요소(예컨대, 등장인물, 장소, 감정상태, 행동 등)에 대한 우선순위가 선택적으로 입력될 수도 있다.
- [0063] 720 단계에서, 영상 추천부(350)는 입력된 스토리 데이터에 해당하는 스크립트 데이터 내에서 미리 지정된 인물, 감정, 행동, 장소 등의 특징점을 추출할 수 있다. 입력된 스토리 데이터에 해당하는 스크립트, 즉, 입력 스크립트에서 추출된 특징점은 유사한 영상 데이터의 비교군을 추출하기 위해 이용될 수 있다.
- [0064] 730 단계에서, 영상 추천부(350)는 상기 입력 스크립트를 상기 추출된 특징점에 기초하여 쇼트 단위로 분해(즉, 분할)할 수 있다. 예컨대, 사용자의 스토리 데이터는 대부분 사용자가 새롭게 만들어 시퀀스(sequence)나 신(scene) 형태로 제작하므로 쇼트(shot) 단위로 분해될 수 있다.
- [0065] 740 단계에서, 영상 추천부(350)는 쇼트 단위로 분해된 입력 스크립트를 대상으로, 유사한 스토리끼리 하나의 그룹에 속하도록 그룹을 형성시킬 수 있다. 그러면, 쇼트 단위로 분해된 스토리에서 추출된 특징점과 스토리의 중요도 패턴을 분석하여 사용자의 스토리 흐름에 알맞은 영상을 추천 영상으로 결정할 수 있다. 예컨대, 사용

자의 스토리 흐름에 알맞은 쇼트 단위의 영상(즉, 쇼트 영상)이 속하는 그룹이 결정되고, 결정된 그룹에 포함된 쇼트 영상들이 사용자 단말로 어플리케이션을 통해 추천될 수 있다.

[0067] 이상에서 설명된 장치는 하드웨어 구성요소, 소프트웨어 구성요소, 및/또는 하드웨어 구성요소 및 소프트웨어 구성요소의 조합으로 구현될 수 있다. 예를 들어, 실시예들에서 설명된 장치 및 구성요소는, 프로세서, 콘트롤러, ALU(arithmetic logic unit), 디지털 신호 프로세서(digital signal processor), 마이크로컴퓨터, FPGA(field programmable gate array), PLU(programmable logic unit), 마이크로프로세서, 또는 명령(instruction)을 실행하고 응답할 수 있는 다른 어떠한 장치와 같이, 하나 이상의 범용 컴퓨터 또는 특수 목적 컴퓨터를 이용하여 구현될 수 있다. 처리 장치는 운영 체제(OS) 및 상기 운영 체제 상에서 수행되는 하나 이상의 소프트웨어 어플리케이션을 수행할 수 있다. 또한, 처리 장치는 소프트웨어의 실행에 응답하여, 데이터를 접근, 저장, 조작, 처리 및 생성할 수도 있다. 이해의 편의를 위하여, 처리 장치는 하나가 사용되는 것으로 설명된 경우도 있지만, 해당 기술분야에서 통상의 지식을 가진 자는, 처리 장치가 복수 개의 처리 요소(processing element) 및/또는 복수 유형의 처리 요소를 포함할 수 있음을 알 수 있다. 예를 들어, 처리 장치는 복수 개의 프로세서 또는 하나의 프로세서 및 하나의 콘트롤러를 포함할 수 있다. 또한, 병렬 프로세서(parallel processor)와 같은, 다른 처리 구성(processing configuration)도 가능하다.

[0068] 소프트웨어는 컴퓨터 프로그램(computer program), 코드(code), 명령(instruction), 또는 이들 중 하나 이상의 조합을 포함할 수 있으며, 원하는 대로 동작하도록 처리 장치를 구성하거나 독립적으로 또는 결합적으로(collectively) 처리 장치를 명령할 수 있다. 소프트웨어는 네트워크로 연결된 컴퓨터 시스템 상에 분산되어서, 분산된 방법으로 저장되거나 실행될 수도 있다. 소프트웨어 및 데이터는 하나 이상의 컴퓨터 판독 가능 기록 매체에 저장될 수 있다.

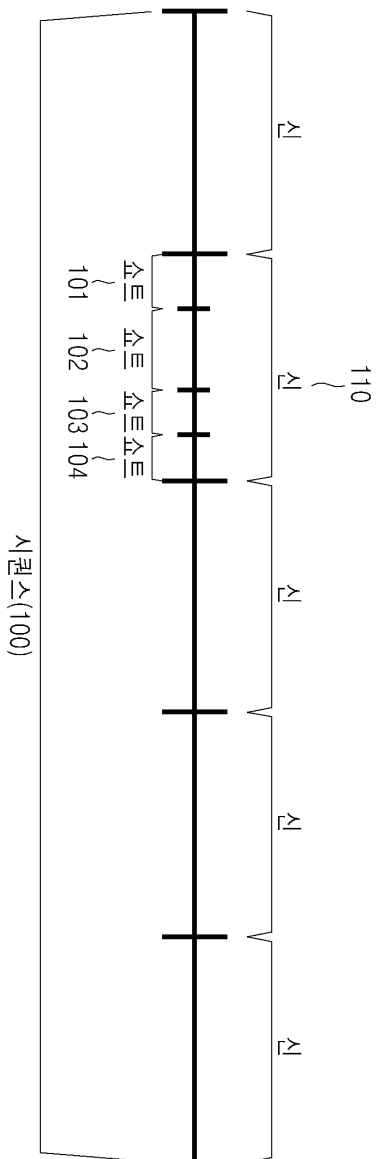
[0069] 실시예에 따른 방법은 다양한 컴퓨터 수단을 통하여 수행될 수 있는 프로그램 명령 형태로 구현되어 컴퓨터 판독 가능 매체에 기록될 수 있다. 상기 컴퓨터 판독 가능 매체는 프로그램 명령, 데이터 파일, 데이터 구조 등을 단독으로 또는 조합하여 포함할 수 있다. 상기 매체에 기록되는 프로그램 명령은 실시예를 위하여 특별히 설계되고 구성된 것들이거나 컴퓨터 소프트웨어 당업자에게 공지되어 사용 가능한 것일 수도 있다. 컴퓨터 판독 가능 기록 매체의 예에는 하드 디스크, 플로피 디스크 및 자기 테이프와 같은 자기 매체(magnetic media), CD-ROM, DVD와 같은 광기록 매체(optical media), 플롭티컬 디스크(floptical disk)와 같은 자기-광 매체(magneto-optical media), 및 롬(ROM), 램(RAM), 플래시 메모리 등과 같은 프로그램 명령을 저장하고 수행하도록 특별히 구성된 하드웨어 장치가 포함된다. 프로그램 명령의 예에는 컴파일러에 의해 만들어지는 것과 같은 기계어 코드뿐만 아니라 인터프리터 등을 사용해서 컴퓨터에 의해서 실행될 수 있는 고급 언어 코드를 포함한다.

[0070] 이상과 같이 실시예들이 비록 한정된 실시예와 도면에 의해 설명되었으나, 해당 기술분야에서 통상의 지식을 가진 자라면 상기의 기재로부터 다양한 수정 및 변형이 가능하다. 예를 들어, 설명된 기술들이 설명된 방법과 다른 순서로 수행되거나, 및/또는 설명된 시스템, 구조, 장치, 회로 등의 구성요소들이 설명된 방법과 다른 형태로 결합 또는 조합되거나, 다른 구성요소 또는 균등물에 의하여 대치되거나 치환되더라도 적절한 결과가 달성될 수 있다.

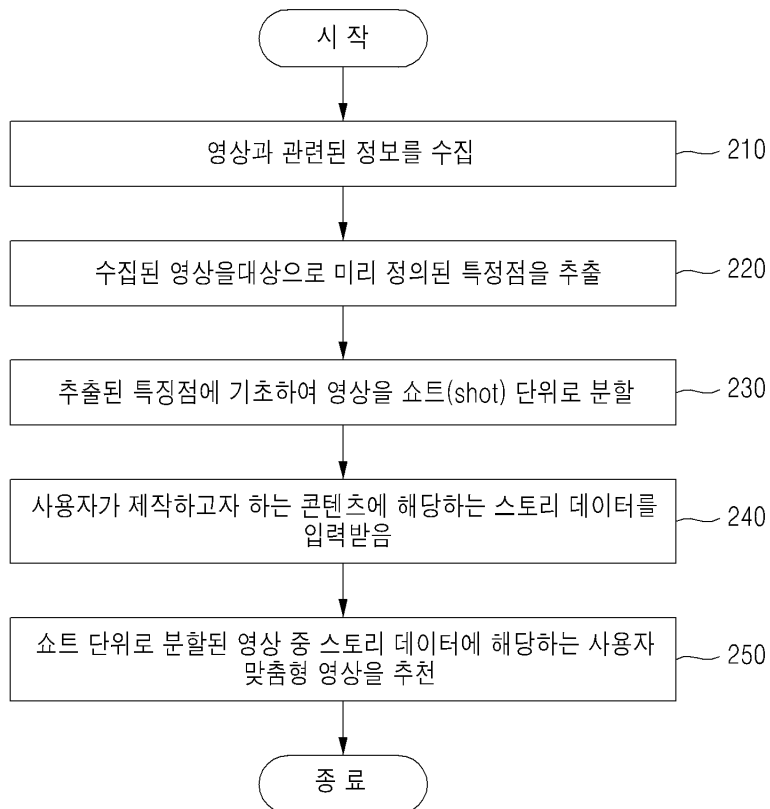
[0071] 그러므로, 다른 구현들, 다른 실시예들 및 특허청구범위와 균등한 것들도 후술하는 특허청구범위의 범위에 속한다.

도면

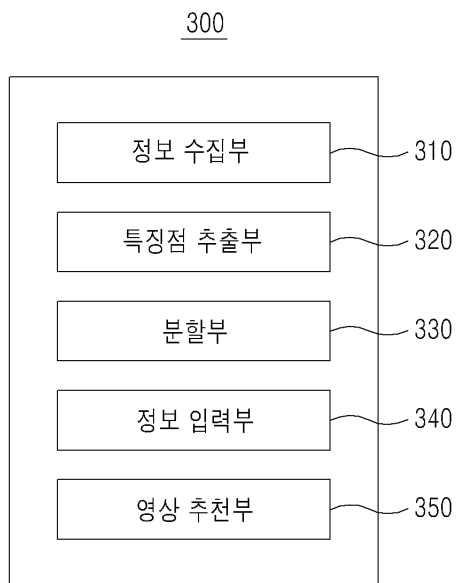
도면1



도면2



도면3



도면4



도면5

5. 호텔 안 계단/ 저녁

긴장한 얼굴로 입구를 살피고는 다급하고 조용한 발걸음으로 안으로 들어서는 그래와 조대리.
좁은 계단을 급히 올라가는데 요르단 남자 하나가 내려오며 흘끔거린다.
한 쪽으로 몸을 피해 주며 올라가는 그래와 조대리.

6. 호텔 안 2층 카페/ 저녁

다급히 그러나 조용히 안으로 들어서는 그래와 조대리
천정의 조각한 붉은색 조명등 장식과 낡은 이슬람식 모자이크 벽면이 눈에 들어오는 허름한 이슬람식 카페다.
벽면 한 쪽에서는 큰 장식용 유리 물 담뱃대가 있고, 그 옆으로 실제 사용하는 청동 물 담뱃대들이 수십 개 진열되어 있다.
여기저기서 물 담배를 피우며 중동식 마작을 하고 있는 남자들도 보인다.
두리번거리던 조대리는 발코니에서 물 담배를 피우고 있는 남자를 본다.

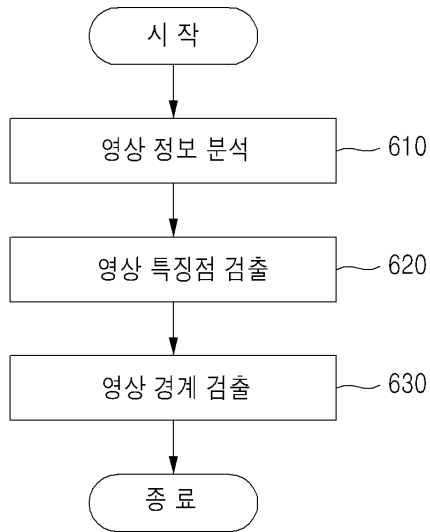
조대리 잠시만요.

다급히 남자에게 가서 인사를 건넌 후 짧은 얘기를 나누고 다시 그래에게 돌아오는 조대리

조대리 502호랍니다.

그래 (끄덕하면)

도면6



도면7

