



(19) 대한민국특허청(KR)
(12) 등록특허공보(B1)

(45) 공고일자 2024년02월13일
(11) 등록번호 10-2635351
(24) 등록일자 2024년02월05일

(51) 국제특허분류(Int. Cl.)

G06Q 50/26 (2024.01) G01D 21/02 (2006.01)
G06Q 50/10 (2012.01) G06T 13/40 (2011.01)
G06V 10/74 (2022.01) G06V 10/82 (2022.01)
G06V 40/16 (2022.01) G08B 21/02 (2006.01)
G10L 15/02 (2006.01)

(52) CPC특허분류

G06Q 50/26 (2024.01)
G01D 21/02 (2013.01)

(21) 출원번호 10-2022-0134005

(22) 출원일자 2022년10월18일

심사청구일자 2022년10월18일

(56) 선행기술조사문헌

JP2020013185 A*
KR1020180096038 A*
KR101377029 B1
KR101877294 B1

*는 심사관에 의하여 인용된 문헌

(73) 특허권자

(주)이미지데이

대전광역시 유성구 대덕대로512번길 20, 대전정보
문화산업진흥원 비동2층 1인 창조기업지원센터(도
룡동)

(72) 발명자

이찬우

경기도 용인시 기흥구 중부대로788번길 20 수원동
마을쌍용아파트 303동 1501호

류병용

경기도 화성시 영통로27번길 53 신영통현대타운
212동 204호

(74) 대리인

윤의상

전체 청구항 수 : 총 14 항

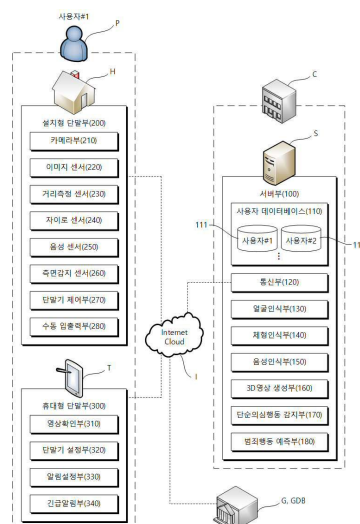
심사관 : 황운철

(54) 발명의 명칭 범죄 예방 시스템

(57) 요약

문 등에 부착식으로 설치하는 설치형 단말부와 서버, 휴대형 단말기를 이용하여 문에 다가오는 사람과, 그 사람의 얼굴, 체형, 행동을 인식하여 범죄를 예측하고 예방할 수 있도록 하는 범죄 예방 시스템을 개시한다. 본 발명의 범죄 예방 시스템은 하나 이상의 연산장치와 기억장치를 포함하는 서버부; 하나 이상의 설치형 단말부; 그리고 하나 이상의 휴대형 단말부를 포함한다.

대표도 - 도1



(52) CPC특허분류

G06Q 50/10 (2015.01)

G08B 21/02 (2013.01)

G08B 3/10 (2021.01)

G08B 6/00 (2021.01)

명세서

청구범위

청구항 1

범죄 예방 시스템으로서,

하나 이상의 연산장치와 기억장치를 포함하는 서버부; 하나 이상의 설치형 단말부; 그리고 하나 이상의 휴대형 단말부를 포함하고,

상기 서버부는 사용자 데이터베이스; 전기통신망을 통하여 상기 하나 이상의 설치형 단말부, 상기 하나 이상의 휴대형 단말부, 기관 및 기관 DB와 유선 또는 무선 중 어느 하나 이상의 방식으로 통신 가능하게 연결되는 통신부; 사람 객체의 얼굴을 인식하기 위한 얼굴인식부; 사람 객체의 체형을 인식하기 위한 체형인식부; 사람 객체의 음성을 인식하기 위한 음성인식부; 3D 영상을 생성하기 위한 3D영상 생성부; 영상에서 사람의 유무를 파악하고, 단순의심행동을 감지하는 단순의심행동 감지부; 그리고 범죄행동 예측부를 포함하며,

상기 설치형 단말부는 영상 및 이미지정보를 획득하기 위한 카메라부 및 이미지 센서; 사람 객체의 접근을 인식하여 거리정보를 생성하는 거리측정 센서; 사람 객체의 행동을 인식하여 자이로정보를 생성하는 자이로 센서; 사람 객체의 소리 및 음성을 인식하여 음성정보를 생성하는 음성 센서; 단말기 제어부; 그리고 입출력 장치 및 디스플레이를 포함하는 수동 입출력부를 포함하고,

상기 휴대형 단말부는 획득 및 분석, 생성된 영상을 하나 이상 확인할 수 있는 영상확인부; 상기 하나 이상의 설치형 단말부 및 휴대형 단말부를 설정하거나 설정상태를 변경하기 위한 단말기 설정부; 알림상태를 설정하기 위한 알림설정부; 긴급상황 발생 시 소리나 진동기능을 이용하여 사용자에게 긴급상황이 발생했음을 알리는 긴급알림부를 포함하고,

상기 사용자 데이터베이스는 하나 이상의 개별 데이터베이스를 포함하며, 상기 하나 이상의 개별 데이터베이스는 각각

해당 사용자의 인적사항 정보 및 로그인 정보를 포함하는 사용자 개인정보; 상기 거리측정 센서가 측정한 거리값정보를 하나 이상 저장하는 거리정보 저장 버퍼; 상기 카메라부 또는 이미지센서 중 어느 하나 이상을 통하여 생성된 이미지정보를 하나 이상 저장하는 이미지정보 저장 버퍼; 상기 자이로센서가 측정한 자이로정보를 하나 이상 저장하는 자이로정보 저장 버퍼; 상기 음성 센서가 녹음한 음성정보를 하나 이상 저장하는 음성정보 저장 버퍼; 하나 이상의 영상 데이터를 포함하는 영상 데이터베이스; 그리고 해당 사용자가 입력 또는 설정한 출동지점 주소를 포함하며,

상기 하나 이상의 거리값정보는 각각 구분정보, 측정한 날짜, 객체와의 측정된 거리 및 시간 데이터를 포함하고, 상기 하나 이상의 자이로정보는 각각 구분정보, 측정한 날짜, X, Y, Z축에 대한 각도 및 가속도 값 중 선택된 하나 이상의 값을 포함하며, 상기 하나 이상의 자이로정보는 각각 X축 각도값, Y축 각도값, Z축 각도값, X축 가속도값, Y축 가속도값, Z축 가속도값의 6개 정보를 포함하고, 상기 자이로센서는 초당 수집횟수 n 이 설정되어 있으며, 상기 하나 이상의 자이로정보는 각각 상기 6개의 자이로 정보값을 상기 초당 수집횟수 n 개 횟수번 연속 수집되어 시계열적으로 통합 나열한 자이로 시퀀스 데이터를 포함하며,

상기 얼굴인식부는 영상 또는 이미지정보에서 식별되는 얼굴 객체에 대한 특성을 추출하기 위한 얼굴 특성추출 신경망; 상기 얼굴 특성추출 신경망을 통해 추출된 얼굴 특성값을 구분 저장하는 얼굴인식 데이터베이스; 그리고 상기 얼굴 특성추출 신경망에서 추출된 얼굴 특성값이 상기 얼굴인식 데이터베이스 내 저장되어 있는지 비교 판단하는 얼굴 매칭 모델부를 포함하고,

상기 체형인식부는 영상 또는 이미지정보에서 식별되는 체형 객체에 대해, 체형을 작은 역삼각체형, 큰 사각체형, 역삼각체형, 작은 사각체형, 표준체형의 5가지로 구분하며,

상기 3D 영상 생성부는, 상기 영상 또는 이미지정보를 기반으로 얼굴 모델링 영상을 생성하는 3D 얼굴영상 생성 모델; 5개의 체형 모델링 영상을 포함하는 3D 체형 생성모델을 포함하는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 2

삭제

청구항 3

삭제

청구항 4

삭제

청구항 5

삭제

청구항 6

삭제

청구항 7

제 1항에 있어서,

상기 음성 정보는 반복적으로 들려오는 소리가 구분 표시되어 있는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 8

제 1항에 있어서,

상기 하나 이상의 영상 데이터는 영상 데이터 본체; 분석 플러그; 그리고 영상분석결과 데이터를 포함하는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 9

제 1항에 있어서,

상기 단말기 제어부는 거리정보 및 이미지정보가 입력(I1, I2)됨에 따라 상기 거리측정에서 측정된 객체와의 거리값(D)이 기 설정된 임계거리값(Dt) 이하인지 측정하는 거리값 판단단계(S1);

상기 단계(S1)에서 입력된 거리값(D)이 기 설정된 임계거리값(Dt) 이하라면, 입력된 거리값(D)이 0인지 판단하는 고장 판단단계(S2);

상기 단계(S2)에서 입력된 거리값(D)이 0이 아니라면, 상기 서버부의 단순의심행동 감지부를 통해 상기 이미지 정보(I2)에 사람이 포함되어 있는지 확인하는 객체 판단 단계(S4);

상기 단계(S4)에서 사람이 영상에 포함되어 있다고 판단되면, 기 설정된 최대 녹화시간(Rt)동안 영상녹화를 실시하여 녹화영상을 생성하는 녹화 단계(S5);

그리고 상기 최대 녹화시간(Rt)동안 녹화한 녹화영상, 상기 최대 녹화시간(Rt)동안 수집한 거리정보, 이미지정보, 자이로정보, 음성정보 중 선택된 어느 하나 이상을 상기 서버부로 전송하는 영상 전송 단계(S6)를 실시하는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 10

삭제

청구항 11

삭제

청구항 12

제 1항에 있어서,

상기 음성인식부는 음성인식 모듈; 기 설정된 하나 이상의 유의어가 저장되어 있는 유의어 데이터베이스; 그리고 영상과 연동된 음성정보에서 상기 유의어 데이터베이스 내 저장된 유의어가 몇 회 포함되는지 식별할 수 있는 자연어 처리부를 포함하고,

상기 자연어 처리부는 유의어 반복횟수 w 가 설정되어 있는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 13

제 1항에 있어서,

상기 범죄행동 예측부는 행동 해석 모델을 포함하여, 상기 행동 해석 모델을 통해 영상 또는 이미지정보에서 식별되는 객체의 문 열기, 물건 배송, 응시, 철사 구부리기 행동을 인식할 수 있는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 14

삭제

청구항 15

제 1항에 있어서,

상기 범죄 예방 시스템은 외부 영상 데이터베이스를 더 포함하고,

상기 외부 영상 데이터베이스는 하나 이상의 개별 영상 데이터베이스; 그리고 데이터베이스 운용부를 포함하며,

상기 하나 이상의 개별 영상 데이터베이스는 해당 사용자가 입력 또는 설정한 출동지점 주소를 포함하는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 16

삭제

청구항 17

제 1항에 있어서,

상기 얼굴인식부는 얼굴 비교 모듈부를 더 포함하고,

상기 얼굴 비교 모듈부는 기관 DB에 포함되어 있는 하나 이상의 수배자와 상기 얼굴 특성추출 신경망을 통해 추출된 특징 추출된 얼굴을 비교하여 기 설정된 임계 수준 유사도 $G\%$ 이상 유사한 지 여부를 판단하는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 18

제 1항에 있어서,

상기 얼굴인식부는 얼굴 비교 모듈부를 더 포함하고,

상기 서버부는 신원정보 DB를 더 포함하여,

상기 얼굴 비교 모듈부는 상기 신원정보 DB에 포함되어 있는 하나 이상의 수배자와 상기 얼굴 특성추출 신경망을 통해 추출된 특징 추출된 얼굴을 비교하여 기 설정된 임계 수준 유사도 $G\%$ 이상 유사한 지 여부를 판단하는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 19

제 1항에 있어서,

상기 얼굴인식부는 얼굴 비교 모델부를 더 포함하고,

상기 하나 이상의 개별 데이터베이스는 각각 등록얼굴 데이터베이스를 더 포함하여,

상기 얼굴 비교 모델부는 상기 등록얼굴 데이터베이스에 포함되어 있는 하나 이상의 경계인원과 상기 얼굴 특성 추출 신경망을 통해 추출된 특징 추출된 얼굴을 비교하여 기 설정된 임계 수준 유사도 G% 이상 유사한 지 여부를 판단하는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 20

제 12항에 있어서,

상기 하나 이상의 개별 데이터베이스는 각각 등록음성 데이터베이스를 더 포함하고,

상기 자연어 처리부는 유의어 반복횟수 w 를 측정할 때 상기 유의어 데이터베이스의 단어 수, 그리고 상기 등록 음성 데이터베이스의 단어 또는 문장의 수를 합하여 측정하는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 21

제 20항에 있어서,

상기 음성인식부는 상기 등록음성 데이터베이스에 저장된 목소리나 성문이 분석 중인 음성정보와 겹쳐지게 인식 되면, 상기 유의어 반복횟수 w 를 초 단위로 증가시키는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 22

제 1항에 있어서,

상기 설치형 단말부는 진동감지 센서를 더 포함하는 것을 특징으로 하는, 범죄 예방 시스템.

청구항 23

제 1항에 있어서,

상기 설치형 단말부는 측면감지 센서를 더 포함하는 것을 특징으로 하는, 범죄 예방 시스템.

발명의 설명

기술 분야

[0001] 본 발명은 범죄 예방 시스템에 대한 것으로서, 보다 상세하게는 문 등에 부착식으로 설치하는 설치형 단말부와 서버, 휴대형 단말기를 이용하여 문에 다가오는 사람과, 그 사람의 얼굴, 체형, 행동을 인식하여 범죄를 예측하고 예방할 수 있도록 하는 범죄 예방 시스템에 대한 것이다.

배경 기술

[0002] 가구원이 한 명인 가구를 의미하는 1인가구는 핵가족화 및 전통적 가족공동체의 붕괴 등 여러 이유로 인하여 2000년대 이후로 폭발적으로 증가하는 추세를 보이고 있다. 이에 따라 1인 가구를 대상으로 한 범죄 역시 증가하고 있는데, 특히 사회적 약자를 대상으로 한 1인 가구 대상 범죄가 크게 증가하고 있어 사회적인 문제가 되고 있다.

[0003] 상기한 사회적 약자들은 일반적으로 신체적 능력이 떨어져 자신을 방어하기 어려운 노약자, 장애인, 여성 등을 지칭하는데 이 중에서도 1인 가구의 대다수를 형성하는 것은 취업 등으로 인하여 독립하여 살고 있는 젊은 여성이며, 이들을 대상으로 한 범죄를 예방하는 것이 사회적 목표 중 하나가 되고 있다.

- [0004] 상기와 같은 1인 여성 가구에 대한 범죄로는, 대검찰청에서 발간하는 "범죄분석"에 따르면 2014년 기준 연인 대상 살인범죄가 861명 중 108명으로 12.6%, 상해는 92,600명 중 2,718명으로 2.9%, 폭행은 110,888명 중 2,967명으로 4.8%, 성폭력은 21,474명 중 686명으로 3.2% 등으로서 그 심각성을 알 수 있다. 상기한 통계는 연인 대상 범죄만 집계한 것으로 1인 여성 가구 대상으로 집계대상을 확대한다면 그 수치는 더 증가할 수밖에 없다.
- [0005] 상기와 같은 여성 1인 가구 대상 범죄를 예방하고자 하는 사회적 목표를 달성하기 위하여 다양한 정책들이 구상 및 실행되고 있지만 큰 효과를 보지 못하는 실정이다. 예를 들어 스토킹 방지법이 제정되어 실행 중이나 테이트 폭력이나 1인 여성 가구 대상 스토킹 범죄는 좀처럼 줄지 않고 있다.
- [0006] 또한 대도시의 경우 안심 서비스가 제공되는 곳이 있으나, 경찰 인력 부족으로 인하여 문제가 발생했을 경우 빠르게 대응하기 어려워 그 실효성에 의문이 들고 있다. 따라서 이러한 범죄를 사전에 예방하고 범죄의 발생을 미리 차단하는 시스템이 필요하게 되었다.
- [0007] 상기와 같은 범죄 발생을 미리 차단하는 시스템은, 영상처리 및 감시 시스템으로 구현될 수 있다. 현재 영상 감시 시스템 시장은 영상으로부터 사람을 탐지하고 추적하여 수상한 물체 및 행위에 따른 위험 방지 등을 사전에 대응할 수 있도록 하는 시스템으로서 아직까지는 CCTV를 통해 촬영된 영상을 사람(감시자)이 실시간으로 확인하여 상기 감시자가 영상을 눈으로 확인하여 문제 발생 여부를 판단하고 조치를 취하는 단순한 구조를 취하고 있다. 하지만 클라우드, 모바일, 영상 및 데이터 분석 기술의 향상으로 인하여 데이터 보호와 관리, 초고해상도 영상의 촬영 및 영상분석 기술이 융합된 클라우드 기반의 지능형 영상감시 시스템으로 발달하고 있다.
- [0008] 가장 유의미하게 발달하고 있는 분야는 영상처리 분야로서, 발달한 영상처리 분야를 범죄사건의 해결에 이용하고자 하는 많은 시도가 발생하고 있다. 예를 들어 등록특허 10-1564760호는 범죄사건의 예측을 위한 영상처리 장치 및 방법으로서 적어도 하나의 가시광선 카메라 및 적어도 하나의 열 카메라를 포함하는 영상생성부를 통하여 적어도 둘 이상의 가시광선 영상과 열 영상을 확보하여 대상체의 행위를 추적하고 얼굴 영역의 온도변화를 검출하여 대상의 행동을 추정하는 기술내용을 개시하고 있다.
- [0009] 그리고 등록특허 10-1647803호는 3차원 데이터베이스에 저장된 3차원 영상정보를 2차원 입력 영상내의 얼굴 따라 가공하여 얼굴 인식에 사용하는 3차원 얼굴모델 투영을 통한 얼굴 인식 방법을, 등록특허 10-1643573호는 3차원 얼굴 데이터베이스를 이용하여 입력되는 2차원 얼굴 영상을 3차원 얼굴 모델로 정합하고 무표정 파라미터를 이용하여 3차원 얼굴 모델을 무표정 보정한 뒤 이를 통해 다시 2차원의 무표정 얼굴 영상을 생성함으로써 얼굴 인식에 사용하는 것을 특징으로 하는 얼굴 인식 방법을 개시하고 있다.
- [0010] 또한 등록특허 10-1854316호는 사용자의 부분 얼굴만 입력되었어도 기 입력되어 있는 사용자 전체 얼굴 이미지와 비교하여 사용자 인증을 수행 및 갱신하여 수배자를 검출하는 데 사용할 수 있는 부분 얼굴 인식을 이용한 수배자 검출 시스템을 개시하고 있다.
- [0011] 상기와 같은 등록특허 기술내용은 주로 얼굴의 인식 형태변화 등에 초점이 맞춰져 있어 본 발명과는 구성 및 동작 방식이 완전히 다른 것이다.
- [0012] 또한 얼굴 인식 외에도 몸체에 대한 인식을 실시하는 기술 또한 개시되어 있는데, 예를 들어 공개특허 10-2021-0081223호는 2D영상 및 3D영상의 융합을 통해 공간상황을 인지하는 방법 및 장치 등에 대해 개시하고 있고, 공개특허 10-2018-0096038호는 행동 패턴 기반 범죄 예측 장치로서 카메라로부터 전송되어 오는 영상의 대상간 거리, 팔 및 손의 움직임을 파악하여 범죄행동을 예측하고자 하는 기술내용을 개시하고 있다.
- [0013] 상기의 공개특허 기술내용들은 융합영상정보의 생성에만 집중하거나 또는 행동내용만으로 판단하기 때문에 본 발명과는 구성 및 동작이 완전히 다른 것으로 판단된다.
- [0014] 이외에도 다수의 카메라가 아닌 최소 6메가픽셀 이상의 초고해상도 고정형 카메라 하나에서 촬영한 영상에 대하여 영상분할 가변인식 기술에 의한 소프트웨어적 방법으로 객체의 행동패턴의 이벤트 영상을 추출하여 미래의 범죄확률을 예측 표시하여 제공하는 영상분석 서버를 통한 미래 발생범죄 예측 시스템 또한 개시되어 있으며, 이 또한 그 구성 및 동작에 있어서 본 발명과는 완전히 다른 것이다.
- [0015] 그리고 스마트 범죄예방 시스템 관련 기술은, CCTV 등과 같은 영상획득 수단을 이용한 것으로는 등록특허 10-0995569호와 같이 지능형 CCTV를 이용하여 상황을 전달하는 범죄예방 시스템, 등록특허 10-1446143호에서와 같이 CCTV 환경에서의 얼굴 인식 기반으로 등록자/미등록자를 인식하고 유사도를 평가하는 관계 서버에 의해 인식 결과를 제공하는 보안 관리 시스템 및 방법, 등록특허 10-2149832호에서와 같이 CCTV에서 녹화한 관계영상을 딥

러닝 기반의 영상처리기법을 통하여 폭력 또는 이상행동을 감지하고 결과를 통지하는 딥러닝 기반 자동 폭력 감지 시스템 등이 있다.

- [0016] 이외에도 셉티드 조명장치 기반 범죄행동인식을 이용한 상황 알림 서비스를 제공하는 등록특허 10-1972743호와 같은 기술내용도 개시되어 있다. 상기의 기술들은 모두 본 발명과는 그 구성 및 동작에 있어서 완전히 다른 것으로 판단된다.

발명의 내용

해결하려는 과제

- [0017] 본 발명은 상기와 같은 종래 기술과는 다른 방식으로, 능동적으로 문 등에 다가오는 사람과 그 사람의 얼굴, 체형, 행동을 인식하여 범죄를 예측할 수 있는 범죄 예방 시스템을 제공하는 데 그 목적이 있다.

과제의 해결 수단

- [0018] 본 발명은 상기와 같은 본 발명의 목적을 달성하기 위하여,
- [0019] 범죄 예방 시스템으로서, 하나 이상의 연산장치와 기억장치를 포함하는 서버부; 하나 이상의 설치형 단말부; 그리고 하나 이상의 휴대형 단말부를 포함하고,
- [0020] 상기 서버부는 사용자 데이터베이스; 전기통신망을 통하여 상기 하나 이상의 설치형 단말부, 상기 하나 이상의 휴대형 단말부, 기관 및 기관 DB와 유선 또는 무선 중 어느 하나 이상의 방식으로 통신 가능하게 연결되는 통신부; 사람 객체의 얼굴을 인식하기 위한 얼굴인식부; 사람 객체의 체형을 인식하기 위한 체형인식부; 사람 객체의 음성을 인식하기 위한 음성인식부; 3D 영상을 생성하기 위한 3D영상 생성부; 영상에서 사람의 유무를 파악하고, 단순의심행동을 감지하는 단순의심행동 감지부; 그리고 범죄행동 예측부를 포함하며,
- [0021] 상기 설치형 단말부는 영상 및 이미지정보를 획득하기 위한 카메라부 및 이미지 센서; 사람 객체의 접근을 인식하여 거리정보를 생성하는 거리측정 센서; 사람 객체의 행동을 인식하여 자이로정보를 생성하는 자이로 센서; 사람 객체의 소리 및 음성을 인식하여 음성정보를 생성하는 음성 센서; 단말기 제어부; 그리고 입출력 장치 및 디스플레이를 포함하는 수동 입출력부를 포함하고,
- [0022] 상기 휴대형 단말부는 획득 및 분석, 생성된 영상을 하나 이상 확인할 수 있는 영상확인부; 상기 하나 이상의 설치형 단말부 및 휴대형 단말부를 설정하거나 설정상태를 변경하기 위한 단말기 설정부; 알림상태를 설정하기 위한 알림설정부; 긴급상황 발생 시 소리나 진동기능을 이용하여 사용자에게 긴급상황이 발생했음을 알리는 긴급알림부를 포함하는 범죄 예방 시스템을 제공한다.
- [0023] 상기에서, 사용자 데이터베이스는 하나 이상의 개별 데이터베이스를 포함하고, 상기 하나 이상의 개별 데이터베이스는 각각 해당 사용자의 인적사항 정보 및 로그인 정보를 포함하는 사용자 개인정보; 상기 거리측정 센서가 측정한 거리값정보를 하나 이상 저장하는 거리정보 저장 버퍼; 상기 카메라부 또는 이미지센서 중 어느 하나 이상을 통하여 생성된 이미지정보를 하나 이상 저장하는 이미지정보 저장 버퍼; 상기 자이로센서가 측정한 자이로정보를 하나 이상 저장하는 자이로정보 저장 버퍼; 상기 음성 센서가 녹음한 음성정보를 하나 이상 저장하는 음성정보 저장 버퍼; 그리고 하나 이상의 영상 데이터를 포함하는 영상 데이터베이스를 포함한다.
- [0024] 상기에서, 하나 이상의 거리값정보는 각각 구분정보, 측정된 날짜, 객체와의 측정된 거리 및 시간 데이터를 포함하는 것이 바람직하다.
- [0025] 상기에서, 하나 이상의 자이로정보는 각각 구분정보, 측정된 날짜, X, Y, Z축에 대한 각도 및 가속도 값 중 선택된 하나 이상의 값을 포함하는 것이 바람직하다.
- [0026] 상기에서, 하나 이상의 자이로정보는 각각 X축 각도값, Y축 각도값, Z축 각도값, X축 가속도값, Y축 가속도값, Z축 가속도값의 6개 정보를 포함하는 것이 바람직하다.
- [0027] 상기에서, 자이로센서는 초당 수집횟수 n 이 설정되어 있으며, 상기 하나 이상의 자이로정보는 각각 상기 6개의 자이로 정보값을 상기 초당 수집횟수 n 개 횟수번 연속 수집되어 시계열적으로 통합 나열한 자이로 시퀀스 데이터를 포함하는 것이 바람직하다.
- [0028] 상기에서, 음성 정보는 반복적으로 들려오는 소리가 구분 표시되어 있는 것이 바람직하다.
- [0029] 상기에서, 하나 이상의 영상 데이터는 영상 데이터 본체; 분석 플러그; 그리고 영상분석결과 데이터를 포함하는

것이 바람직하다.

- [0030] 상기에서, 단말기 제어부는 거리정보 및 이미지정보가 입력(I1, I2)됨에 따라 상기 거리측정에서 측정된 객체와의 거리값(D)이 기 설정된 임계거리값(Dt) 이하인지 측정하는 거리값 판단단계(S1); 상기 단계(S1)에서 입력된 거리값(D)이 기 설정된 임계거리값(Dt) 이하라면, 입력된 거리값(D)이 0인지 판단하는 고장 판단단계(S2); 상기 단계(S2)에서 입력된 거리값(D)이 0이 아니라면, 상기 서버부의 단순의심행동 감지부를 통해 상기 이미지정보(I2)에 사람이 포함되어 있는지 확인하는 객체 판단 단계(S4); 상기 단계(S4)에서 사람이 영상에 포함되어 있다고 판단되면, 기 설정된 최대 녹화시간(Rt)동안 영상녹화를 실시하여 녹화영상을 생성하는 녹화 단계(S5); 그리고 상기 최대 녹화시간(Rt)동안 녹화한 녹화영상, 상기 최대 녹화시간(Rt)동안 수집한 거리정보, 이미지정보, 자이로정보, 음성정보 중 선택된 어느 하나 이상을 상기 서버부로 전송하는 영상 전송 단계(S6)를 실시한다.
- [0031] 상기에서, 얼굴인식부는 영상 또는 이미지정보에서 식별되는 얼굴 객체에 대한 특성을 추출하기 위한 얼굴 특성 추출 신경망; 상기 얼굴 특성추출 신경망을 통해 추출된 얼굴 특성값을 구분 저장하는 얼굴인식 데이터베이스; 그리고 상기 얼굴 특성추출 신경망에서 추출된 얼굴 특성값이 상기 얼굴인식 데이터베이스 내 저장되어 있는지 비교판단하는 얼굴 매칭 모델부를 포함하는 것이 바람직하다.
- [0032] 상기에서, 체형인식부는 영상 또는 이미지정보에서 식별되는 체형 객체에 대해, 체형을 작은 역삼각체형, 큰 사각체형, 역삼각체형, 작은 사각체형, 표준체형의 5가지로 구분하는 것이 바람직하다.
- [0033] 상기에서, 음성인식부는 음성인식 모듈; 기 설정된 하나 이상의 유의어가 저장되어 있는 유의어 데이터베이스; 그리고 영상과 연동된 음성정보에서 상기 유의어 데이터베이스 내 저장된 유의어가 몇 회 포함되는지 식별할 수 있는 자연어 처리부를 포함하고, 상기 자연어 처리부는 유의어 반복횟수 w가 설정되어 있는 것이 바람직하다.
- [0034] 상기에서, 범죄행동 예측부는 행동 해석 모델을 포함하여, 상기 행동 해석 모델을 통해 영상 또는 이미지정보에서 식별되는 객체의 문 열기, 물건 배송, 응시, 철사 구부리기 행동을 인식할 수 있는 것이 바람직하다.
- [0035] 상기에서, 3D 체형 생성모델은, 상기 영상 또는 이미지정보를 기반으로 얼굴 모델링 영상을 생성하는 3D 얼굴영상 생성 모델; 5개의 체형 모델링 영상을 포함하는 3D 체형 생성모델을 포함하는 것이 바람직하다.
- [0036] 상기에서, 범죄 예방 시스템은 외부 영상 데이터베이스를 더 포함하고, 상기 외부 영상 데이터베이스는 하나 이상의 개별 영상 데이터베이스; 그리고 데이터베이스 운용부를 포함하는 것이 바람직하다.
- [0037] 상기에서, 하나 이상의 개별 데이터베이스, 또는 하나 이상의 개별 영상 데이터베이스는 해당 사용자가 입력 또는 설정한 출동지점 주소를 포함하는 것이 바람직하다.
- [0038] 상기에서, 얼굴인식부는 얼굴 비교 모델부를 더 포함하고, 상기 얼굴 비교 모델부는 기관 DB에 포함되어 있는 하나 이상의 수배자와 상기 얼굴 특성추출 신경망을 통해 추출된 특징 추출된 얼굴을 비교하여 기 설정된 임계 수준 유사도 G% 이상 유사한 지 여부를 판단하는 것이 바람직하다.
- [0039] 상기에서, 얼굴인식부는 얼굴 비교 모델부를 더 포함하고, 상기 서버부는 신원정보 DB를 더 포함하여, 상기 얼굴 비교 모델부는 상기 신원정보 DB에 포함되어 있는 하나 이상의 수배자와 상기 얼굴 특성추출 신경망을 통해 추출된 특징 추출된 얼굴을 비교하여 기 설정된 임계 수준 유사도 G% 이상 유사한 지 여부를 판단하는 것이 바람직하다.
- [0040] 상기에서, 얼굴인식부는 얼굴 비교 모델부를 더 포함하고, 상기 하나 이상의 개별 데이터베이스는 각각 등록얼굴 데이터베이스를 더 포함하여, 상기 얼굴 비교 모델부는 상기 등록얼굴 데이터베이스에 포함되어 있는 하나 이상의 경계인원과 상기 얼굴 특성추출 신경망을 통해 추출된 특징 추출된 얼굴을 비교하여 기 설정된 임계 수준 유사도 G% 이상 유사한 지 여부를 판단하는 것이 바람직하다.
- [0041] 상기에서, 하나 이상의 개별 데이터베이스는 각각 등록음성 데이터베이스를 더 포함하고, 상기 자연어 처리부는 유의어 반복횟수 w를 측정할 때 상기 유의어 데이터베이스의 단어 수, 그리고 상기 등록음성 데이터베이스의 단어 또는 문장의 수를 합하여 측정하는 것이 바람직하다.
- [0042] 상기에서, 음성인식부는 상기 등록음성 데이터베이스에 저장된 목소리나 성문이 분석 중인 음성정보와 겹쳐지게 인식되면, 상기 유의어 반복횟수 w를 초 단위로 증가시키는 것이 바람직하다.
- [0043] 상기에서, 설치형 단말부는 진동감지 센서를 더 포함하는 것이 바람직하다.
- [0044] 상기에서, 설치형 단말부는 측면감지 센서를 더 포함하는 것이 바람직하다.

발명의 효과

- [0045] 본 발명에 의하면, 문으로 접근하는 사람과 그 사람의 얼굴, 체형을 용이하게 인식하고, 또한 그 사람의 행동을 인식하여 그 사람의 예상되는 행동을 예측함으로써 단순 의심행동과 범죄 의심 행동을 구분하고, 범죄 의심 행동이 예상되는 사람에 대하여 조치를 취할 수 있도록 함으로서 예방적인 범죄 대응이 가능해진다.

도면의 간단한 설명

- [0046] 도 1은 본 발명의 범죄 예방 시스템의 개략 구조도.
 도 2는 사용자 데이터베이스의 구조도.
 도 3은 설치형 단말부의 동작 순서도.
 도 3a 및 도 3b는 측면감지 센서의 동작 예시 및 순서도.
 도 4는 얼굴인식부의 구조도.
 도 5는 체형 구분 도식도.
 도 6은 체형인식부의 구조도.
 도 7은 음성인식부의 구조도.
 도 8 내지 도 10은 행동 인식 관련 도식도.
 도 11은 본 발명의 범죄 예방 시스템의 제2예시 개략 구조도.
 도 12는 외부 영상 데이터베이스의 구조도.
 도 13 내지 15는 본 발명의 범죄 예방 시스템의 동작 예시 구조도.
 도 16은 본 발명의 휴대형 단말기의 예시도.

발명을 실시하기 위한 구체적인 내용

- [0047] 이하에서는 바람직한 실시예와 첨부된 도면을 참조하여 본 발명을 보다 상세히 설명한다. 하기의 설명은 본 발명의 이해와 실시를 돕기 위한 것이지 본 발명을 이에 한정하려는 것은 아니다. 본 발명이 속한 기술분야에서 통상의 지식을 가진 자들은 이하의 청구범위에 기재된 본 발명의 사상 내에서 다양한 변형이나 수정 또는 변경이 있을 수 있음을 이해할 것이다.
- [0048] 설명에 앞서, 본 명세서에 있어서 '부(部)'란, 하나 이상의 물리적 하드웨어에 의해 실현되는 유닛(unit), 소프트웨어에 의해 실현되는 유닛, 양방을 이용하여 실현되는 유닛을 포함한다. 또한, 1 개의 유닛이 2 개 이상의 하드웨어를 이용하여 실현되어도 되고, 2개 이상의 유닛이 1 개의 하드웨어에 의해 실현되어도 된다. 한편, '~부'는 소프트웨어 또는 하드웨어에 한정되는 의미는 아니며, '~부'는 어드레싱 할 수 있는 저장 매체에 있도록 구성될 수도 있고 하나 또는 그 이상의 프로세서들을 재생시키도록 구성될 수도 있다.
- [0049] 따라서, 일 예로서 '~부'는 소프트웨어 구성요소들, 객체지향 소프트웨어 구성요소들, 클래스 구성요소들 및 태스크 구성요소들과 같은 구성요소들과, 프로세스들, 함수들, 속성들, 프로시저들, 서브루틴들, 프로그램 코드의 세그먼트들, 드라이버들, 펌웨어, 마이크로코드, 회로, 데이터, 데이터베이스, 데이터 구조들, 테이블들, 어레이들 및 변수들을 포함한다. 구성요소들과 '~부'들 안에서 제공되는 기능은 더 작은 수의 구성요소들 및 '~부'들로 결합되거나 추가적인 구성요소들과 '~부'들로 더 분리될 수 있다. 뿐만 아니라, 구성요소들 및 '~부'들은 디바이스 또는 보안 멀티미디어카드 내의 하나 또는 그 이상의 CPU들을 재생시키도록 구현될 수도 있다.
- [0050] 또한, 본 명세서에 있어서 단말, 장치 또는 디바이스가 수행하는 것으로 기술된 동작이나 기능 중 일부는 해당 단말, 장치 또는 디바이스와 연결된 서버에서 대신 수행될 수도 있다. 이와 마찬가지로, 서버가 수행하는 것으로 기술된 동작이나 기능 중 일부도 해당 서버와 연결된 단말, 장치 또는 디바이스에서 수행될 수도 있다.
- [0051] 도 1은 본 발명의 범죄 예방 시스템의 구조도이다. 이하에서는 도 1을 통하여 본 발명의 범죄 예방 시스템의 개략적인 구성 및 동작에 대하여 설명한다.
- [0052] 본 발명의 범죄 예방 시스템은 본 발명의 시스템을 구축하고 서비스를 제공하는 서비스 제공자(C)와 서비스를

제공받는 사용자(P)간 구축 및 사용되는 서비스이며, 따라서 상기 서비스 제공자(C)는 유, 무형의 서비스를 제공할 수 있도록 하나 이상의 연산장치 및 기억장치를 포함하는 서버부(100)를 포함한다.

- [0053] 이때 상기 서버부(100)는 하나 이상의 물리적 서버 컴퓨터(S)를 통해 구현 및 실현될 수 있으며, 상기 서버 컴퓨터(S)는 CPU와 같은 연산장치, 그리고 통상의 기억장치를 하나 이상 포함하는 통상의 서버용 컴퓨터 내지는 데스크톱 컴퓨터 등의 컴퓨팅 시스템을 사용하면 된다.
- [0054] 또한 상기 서버 컴퓨터(S)는 상기 서비스 제공자(C) 내 직원 등이 상기 서버부(100)에 제어 명령이나 정보를 입력하고 출력 내용을 인지하기 위한 통상의 입출력 장치 및 모니터와 같은 디스플레이를 모두 포함한다. 이러한 상기 서버 컴퓨터(S)의 구성은 통상의 서버 시스템을 사용하는 것이므로 이에 대한 설명은 생략한다.
- [0055] 그리고 서비스를 제공받는 사용자(P)의 거주지(H)에는 하나 이상의 설치형 단말부(200)가 설치된다. 이때 상기 사용자의 거주지(H)는 상기 사용자(P)가 지정한 거주지이며, 상기 설치형 단말부(200)는 상기 거주지(H)의 대문, 출입문, 창문 등과 같이 외부로 통하는 공간에 하나 이상 설치될 수 있다.
- [0056] 그리고 상기 설치형 단말부(200) 또한 하나 이상의 연산장치와 기억장치를 포함하여 내부의 기능부들을 유형 및 무형으로 구현하는 것이 바람직한데, 설치 형태 및 크기 등을 고려하여 상기 연산장치는 내부적인 자체 운영 프로그램과 CPU, 인터페이스 등을 제공하는 저전력 통합 임베디드(Embedded) 보드를 사용하는 것이 바람직하다. 상기와 같은 임베디드 보드 등을 통한 설치형 단말부(200) 내부 컴퓨팅 시스템의 구축은 통상의 방식으로 행해지면 되므로 이에 대한 설명은 생략한다.
- [0057] 또한 상기 설치형 단말부(200)의 형태는 상기 서비스 제공자(C)가 결정하여 제공할 수 있는데, 통상의 CCTV와 같은 형태이거나 문손잡이, 잠금장치 등에 부가하여 설치하는 형태, 초인종과 결합하여 설치하는 형태 등 다양한 형태로 제공될 수 있다. 이하에서는 그 중 상기 설치형 단말부(200)가 사용자의 거주지(H)의 대문 초인종에 결합되어 설치 및 제공되는 것을 일례시로서 설명하기로 한다.
- [0058] 그리고 본 발명의 시스템은, 상기 사용자(P)가 휴대하고 다니는 하나 이상의 휴대형 단말부(300)를 포함한다. 상기 휴대용 단말부(300)는 상기 사용자(P)가 휴대하고 다니거나 거치하여 사용 가능하며, 하나 이상의 연산장치와 기억장치, 입출력장치 및 디스플레이를 포함하고 프로그램을 설치할 수 있는 통상의 단말기 장치(T)를 통하여 구현되는데 상기 장치(T)는 예를 들어 스마트폰, PDA, 태블릿 PC, 랩톱 컴퓨터, 데스크톱 컴퓨터, 스마트워치 등의 종래의 기기를 사용할 수 있다.
- [0059] 상기와 같은 설치형 단말부(200) 및 휴대형 단말부(300)는 또한, 각각 하나 이상 포함되어 사용될 수 있는데 예를 들어 사용자는 자신의 스마트폰 및 랩톱 컴퓨터를 상기 휴대형 단말부(300)로서 포함시켜 독립적으로 사용할 수 있고, 상기 설치형 단말부(200) 또한 집의 대문에 설치하면서 동시에 창문에도 설치하여 2개로 구성할 수 있는 것이다. 각각의 구성요소들은 모두 동일하게 하면 된다.
- [0060] 그리고 상기 하나 이상의 설치형 단말부(200) 및 휴대형 단말부(300)는, 전기통신망(I)을 통해 유선 또는 무선 중 선택된 어느 하나 이상의 수단을 통하여 서버부(100)와 통신 가능하게 연결된다. 통신 방법 및 수단, 프로토콜 등은 종래의 장치 및 방식에 따라 구성 및 사용하면 된다.
- [0061] 또한 본 발명에서 사용되는 상기 전기통신망(I)은, 인터넷(Internet)과 같이 개방되어 있는 공중의 전기통신망 뿐만 아니라, 인트라넷(Intranet)과 같은 폐쇄형 전기통신망을 포함한다. 따라서 회사 사내망이나 보안망과 같이 접근이 제한된 통신 네트워크망도 본 발명에서의 통신망으로서 사용될 수 있는 것이다.
- [0062] 그리고 상기 서버부(100)는 상기와 같은 전기통신망(I)을 통해, 경찰서 또는 경찰서 내의 전산망, 정부조직의 행정전산망과 같은 기관(G)과 유선 또는 무선으로 통신 가능하게 연결되어 있다. 따라서 상기 서버부(100)는 필요에 따라 상기 기관(G)에게 통보하여, 정보의 획득 또는 제공, 출동신호의 요청을 할 수 있다. 따라서 예를 들어 사용자(P)의 설치형 단말부(200)가 촬영하고 분석된 영상에서 범죄의심 현상이 드러난다거나, 상기 휴대형 단말부(300)를 통해 긴급한 출동요청이 들어오게 되면 상기 서버부(100)는 이를 경찰서와 같은 기관(G)에 전달하여 상기 기관(G)에서 필요한 조치를 취할 수 있게 된다.
- [0063] 이하에서는, 경찰서나 지구대, 관할 관청과 같은 실체로서의 정부기관 내지는 지방자치기관, 공공기관, 공기업 기관 등과 같은 시설을 기관(G)으로 칭하며, 상기 기관(G)에서 운용하는 정부 데이터베이스 내지는 지방자치기관 데이터베이스, 공공기관, 공기업기관 등의 데이터베이스 통칭으로서 기관 데이터베이스(GDB)라 칭하기로 한다.
- [0064] 이하에서는 사용자(P)가 자신의 거주지(H) 대문에 초인종 기능을 포함하는 설치형 단말부(200) 하나를

설치하고, 자신의 스마트폰 기기(T)를 휴대형 단말부(300)로 사용하는 것을 일예시로 하여 설명하기로 한다.

- [0065] 상술한 바와 같이 본 발명에 포함되는 구성요소들에 대하여 이하에서는 각 구성요소들에 포함되는 세부 구성요소들에 대하여 설명한다.
- [0066] 먼저, 상기 서버부(100)는, 상기 사용자(P)의 개별 데이터베이스(111, 112...)를 포함하는 사용자 데이터베이스(110)와, 상기 하나 이상의 설치형 단말부(200) 및 휴대형 단말부(300)와 전기통신망(I)을 통하여 유선 또는 무선 중 선택된 어느 하나 이상의 수단을 통하여 통신하기 위한 통신부(120), 상기 설치형 단말부(200)에서 전송하는 영상에서 얼굴을 인식하기 위한 얼굴인식부(130), 체형을 인식하기 위한 체형인식부(140), 음성을 인식하기 위한 음성인식부(150), 3D영상을 생성하기 위한 3D영상 생성부(160), 상기 영상에서 사람의 유무를 파악하고 단순의심행동을 감지하는 단순의심행동 감지부(170), 범죄행동을 예측하는 범죄행동 예측부(180)를 포함한다.
- [0067] 그리고 상기 설치형 단말부(200)는 영상을 획득하기 위한 카메라부(210) 및 이미지 센서(220), 사람 객체의 접근을 인식하기 위한 거리측정 센서(230), 사람 객체의 행동을 인식하기 위한 자이로 센서(240), 주변 소리 및 음성을 인식할 수 있는 음성 센서(250), 해당 설치형 단말부(200)의 제어를 위한 단말기 제어부(270)와, 해당 설치형 단말부(200)를 수동으로 제어하기 위한 입출력 장치 및 디스플레이를 포함하는 수동 입출력부(280)를 포함한다.
- [0068] 그리고 상기 휴대형 단말부(300)는 획득 및 분석, 생성된 영상을 하나 이상 확인할 수 있는 영상확인부(310), 상기 설치형 단말부(200) 및 휴대형 단말부(300)를 설정하거나 설정상태를 변경하기 위한 단말기 설정부(320), 알람상태를 설정하기 위한 알람설정부(330), 긴급상황이 발생 시 상기 단말부(300)를 이루는 스마트폰 기기(T)의 소리나 진동기능을 이용하여 사용자(P)에게 긴급상황이 발생했음을 알리는 긴급알림부(340)를 포함한다.
- [0069] 이하에서는 상기의 세부 구성요소들에 대한 기능 및 동작에 대하여 설명한다.
- [0070] 먼저, 상기 서버부(100)의 세부 구성요소에 대하여 도 1 및 도 2를 통하여 설명한다. 도 2는 상기 사용자 데이터베이스(110) 내 개별 데이터베이스(111)의 구조도이다.
- [0071] 설명에 앞서, 상기 개별 데이터베이스(111)는 사용자 마다 모두 동일한 내용으로 구성된다. 따라서, 이하에서는 사용자 1(P)의 개별 데이터베이스(111)를 일예시로 하여 설명하지만, 나머지 개별 데이터베이스의 구성요소 또한 모두 동일하므로 그 설명은 생략된다.
- [0072] 상술한 바와 같이, 사용자 데이터베이스(110)는 본 발명의 서비스를 이용하는 사용자(P)의 개별 데이터베이스(111, 112...)로서 구분되어 저장공간에 형성된다.
- [0073] 상기 개별 데이터베이스(111)는 해당 사용자의 인적사항 정보 및 접속을 위한 ID 또는 식별번호, 비밀번호 등의 로그인 정보를 포함하는 사용자 개인정보(1110)를 포함한다. 상기 사용자 개인정보(1110)는 예를 들어, 로그인 정보로서 사용자의 구분 및 접속을 위한 ID와 패스워드, 사용자의 전화번호 및 거주지의 주소, 사용자가 사용하는 하나 이상의 설치형 단말부(200) 및 휴대형 단말부(300)에 대한 고유번호, UUID(Universally Unique Identifier; 범용 고유 식별자) 등의 구분정보를 포함한다.
- [0074] 그리고 상기 사용자 개인정보(1110)와 별개로, 상기 개별 데이터베이스(111)는 출동지점 주소(1111)를 포함할 수 있다. 상기 출동지점 주소(1111)는 상기 기관(G)에서 해당 사용자(P)에 대하여 범죄 등이 예상되어 출동 및 현장방문할 시 최우선 방문지역으로 지정된 곳으로서, 해당 사용자(P)의 상기 설치형 단말부(200)가 설치된 거주지 혹은 그 주변이 될 수도 있지만, 전혀 다른 제3의 장소가 될 수도 있다. 상기과 같은 출동지점 주소(1111)는 해당 사용자(P)가 지정하여 입력하도록 하는 것이 바람직하다.
- [0075] 그리고 상기 개별 데이터베이스(111)는 하나 이상의 데이터 버퍼를 포함하는데, 바람직하게는 상기 거리측정 센서(230)가 실시간으로 측정하는 거리값 정보를 하나 이상 저장하는 거리정보 저장 버퍼(1112), 상기 카메라부(210) 또는 이미지 센서(220) 중 어느 하나 이상을 통하여 실시간으로 촬영하는 영상의 이미지정보를 하나 이상 저장하는 이미지정보 저장 버퍼(1113), 상기 자이로 센서(240)가 실시간으로 측정하는 자이로정보를 하나 이상 포함하여 저장하는 자이로정보 저장 버퍼(1114)와, 상기 음성 센서(250)가 실시간으로 녹음하는 음성정보를 하나 이상 포함하여 저장할 수 있는 음성정보 저장 버퍼(1115)를 포함하는 것이 바람직하다.
- [0076] 그리고 상기 개별 데이터베이스(111)는 상기 설치형 단말부(200)의 카메라부(210) 및 이미지 센서(220)와 음성 센서(250)를 통해 녹화 및 획득할 수 있는, 하나 이상의 영상 데이터(V1, V2, V3, V4...)를 포함하는 영상 데이터베이스(1116)를 포함한다.

- [0077] 이하에서는 상기 버퍼(1112~1115)와, 상기 영상 데이터베이스(1116)에 저장되는 각종의 데이터 정보 및 영상정보의 획득 및 저장을 위한 구성 및 절차에 대하여 설명한다.
- [0078] 설명에 앞서, 상기 버퍼(1112~1115)에 저장되는 거리값정보, 이미지정보, 자이로정보, 음성정보는 모두 해당 데이터가 측정, 계산, 생성 또는 녹음된 해당 설치형 단말부(200)에 대한 고유번호, UUID(Universally Unique Identifier; 범용 고유 식별자) 등의 구분정보와, 해당 데이터가 측정, 계산, 생성 또는 녹음된 날짜를 포함한다.
- [0079] 먼저 상기 거리정보 저장 버퍼(1112)는 상술한 바와 같이 상기 거리측정 센서(230)가 실시간으로 측정하는 거리값 정보를 하나 이상 저장한다. 상기 거리측정 센서(230)는 따라서, 상기 설치형 단말부(200) 기준으로 거리 측정의 대상이 되는 객체와의 거리를 측정할 수 있어야 하는데, 바람직하게는 상기 거리측정 센서(230)를 통상의 초음파 센서로 구성하여 거리를 측정할 수 있다. 아니면 통상의 레이저 센서나 적외선 센서, TOF(Time of Flight)센서, 위상차 변위측정 센서 등 통상의 거리를 측정할 수 있는 센서들 중 하나 이상의 상기 서비스 제공자(C)가 선택하여 구성할 수 있다. 이하에서는 상기 거리측정 센서(230)가 초음파 센서인 것을 일례시로 하여 설명한다.
- [0080] 통상적인 I2C 인터페이스를 통해 사용되는 초음파 센서는 초음파 신호를 보내는 부분과 받는 부분으로 분리되어 있다. 따라서 상기 거리측정 센서(230)는 초음파 신호 출력부와 초음파 신호 수용부를 포함할 수 있으며, 상기 출력부에서 초음파 신호를 보내고, 거리측정 대상 객체와 충돌 후 반사된 초음파 신호가 다시 수용부를 통해 받을 때까지의 시간을 통해 아래 수학적 식 1과 같이 상기 객체와 단말부(200)간 거리를 계산한다.

수학적 식 1

$$D = \frac{340}{2} t$$

- [0081]
- [0082] 여기서 D는 상기 객체와 단말부 간 거리이고, 단위는 미터(m)를 쓰는 것이 바람직하다. 그리고 t는 초음파 신호가 출력 후 수용부를 통해 받을 때까지의 시간이며 단위는 초(s)를 쓰는 것이 바람직하며, 상기 측정된 거리(D) 및 시간(t)데이터는 상기한 거리값 정보에 포함된다.
- [0083] 또한 상기 거리측정 센서(230)는 하나 이상의 명령어를 내장하여 다양한 상황에 대비할 수 있도록 하는 것이 바람직한다. 예를 들어 거리 측정에 있어서 기 설정된 센서의 최대 측정 거리에 도달하고 돌아오는 시간보다 측정 시간이 오래 걸리면 측정에 실패한 것으로서 거리 측정을 중단하고 거리 측정의 실패를 상기 서버부(100) 또는 상기 휴대형 단말부(300)에 통보할 수 있다. 그리고 상기 거리측정 센서(230)를 구성하는 센서 자체의 오차가 주변 환경, 예를 들어 온도와 기압 등에 따라 초음파 매질의 성질 변동에 따른 오차가 발생할 수 있고, 센서 자체에도 시간 측정에 있어서 오차가 발생할 수 있다. 따라서 상기 거리측정 센서(230)는 실제 사용 전에 상기 설치형 단말부(200)가 설치되는 환경에 맞춰서 여러 번 시험 측정하여 거리 측정의 평균, 최소, 최대값을 설정하고 갱신할 수 있는 기능, 그리고 상기 수학적 식 1의 상수들을 보정할 수 있는 기능을 상기 거리측정 센서(230)에 내장된 프로그램 또는 소프트웨어에 포함하는 것이 바람직하다.
- [0084] 그리고 상기 이미지정보 저장 버퍼(1113)는 상술한 바와 같이 상기 카메라부(210)와 상기 이미지센서(220) 중 어느 하나에서 얻을 수 있는, 실시간으로 촬영된 하나 이상의 이미지를 저장한다.
- [0085] 여기서 상기 카메라부(210)와 상기 이미지센서(220)는 별개의 방식으로 촬영하는 두 개의 카메라, 즉 가시광선을 촬영하는 카메라부(210)와 적외선 영상을 촬영하는 이미지센서(220)로서 본 발명에 포함될 수도 있지만, 두 구성요소가 일체형으로 구성되어 상기 카메라부(210)가 하나 이상의 방식으로 촬영하는 촬영부분이고, 상기 이미지센서(220)가 상기 카메라부(210)가 촬영하는 이미지를 가공하는 부분으로서 본 발명에 포함될 수 있다. 본 발명에서는 상기 카메라부(210)가 촬영하고 상기 이미지센서(220)가 이미지를 가공하여 양 구성요소(210, 220)가 통합적으로 구성 및 기능하는 것을 일례시로 하여 설명한다.
- [0086] 상기 카메라부(210)는 상기 설치형 단말부(200)의 외부를 촬영할 수 있는 카메라로서, 통상적인 영상 촬영에 사용되는 적외선 카메라나 가시광선 카메라 중 선택된 하나 이상의 카메라를 사용할 수 있다. 바람직하게는, 신원 동작 인식과 야간의 어두운 환경에서의 원활한 촬영을 위해 적외선 카메라를 사용하는 것이 바람직하다.
- [0087] 그리고 상기 이미지센서(220)는, 하나 이상의 이미지 프로세싱용 연산장치를 포함하여 상기 카메라부(210)가 촬

영한 영상정보를 전달 받아 RGB 버퍼 형태의 데이터로 변환되어 상기 이미지정보 저장 버퍼(1113)에 저장된다.

[0088] 또한 상기 카메라부(210)의 원활한 촬영 및 동작을 위하여, 상기 단말기 제어부(270)는 상기 카메라부(210)의 촬영 카메라에 대한 제어기능으로서 사용권한의 요청(Camera Open), 사용권한의 해제(Camera Close), 카메라 상태 확인(Camera Status; 카메라의 상태, 사용가능 여부 정보 확인), 카메라의 영상 캡처 및 저장공간에 저장(Camera Capture), 그리고 카메라의 노출, 밝기, 컬러를 조절(Camera Adjust) 기능을 포함하고 있는 것이 바람직하다.

[0089] 그리고 상기 자이로정보 저장 버퍼(1114)에는 상술한 바와 같이 상기 자이로 센서(240)에서 측정하는 하나 이상의 자이로정보가 포함되어 저장된다.

[0090] 상기 자이로 센서(240)는 영상 속 객체의 행동을 인식 및 파악하기 위한 것으로서, 통상의 자이로센서를 포함하여 구성될 수 있다. 이에 따라 상기 자이로 센서(240)에서 측정되는 자이로정보는 상기 카메라부(210)를 통해 촬영되는 영상 속 객체의 X, Y, Z축에 대한 각도(Angle)값과 가속도(Velocity)값에 대하여 아래의 표 1에서와 같은 구성의 6개의 자이로 정보값을 포함한다.

표 1

자이로 정보값	설명
Angle X	중력가속도에 대한 X축 각도 크기
Angle Y	중력가속도에 대한 Y축 각도 크기
Angle Z	중력가속도에 대한 Z축 각도 크기
Velocity X	X축 회전 속도
Velocity Y	Y축 회전 속도
Velocity Z	Z축 회전 속도

[0092] 상기 자이로정보는 하나의 객체에 대하여 상기한 6개의 자이로 정보값을 포함하고, 상기 자이로정보 저장 버퍼(1114)에 상기한 자이로정보 값을 저장한다. 이때 자이로정보를 통해 구체적인 범죄 행동 예측을 위하여, 상기 자이로정보 내 6개의 자이로 정보값을 각각 초당 최소 n개, 예를 들어 10개 이상 수집하여 시계열적(Time Series)으로 통합 나열한 데이터를 자이로 시퀀스 데이터로 하여 상기 자이로정보 저장 버퍼(1114)에 저장시킬 수 있다. 상기와 같이 자이로 시퀀스 데이터가 상기 자이로정보 저장 버퍼(1114)에 저장된다면, 범죄행동의 예측에 상기 자이로 시퀀스 데이터를 사용하는 것이 바람직하다.

[0093] 여기서 상기 초당 수집횟수 n은 상기 자이로센서에 기 설정되어 있는 값으로서 시스템 설계자가 설정하는 값이다.

[0094] 그리고 상기 음성정보 저장 버퍼(1115)에는 상술한 바와 같이 상기 음성 센서(250)에서 측정하는 하나 이상의 음성정보가 포함되어 저장된다.

[0095] 상기 음성 센서(250)는 통상의 음성 및 소리 녹음 센서 및 구성을 사용하여 달성할 수 있다. 이러한 상기 카메라부(210)를 통해 촬영되고 있는 객체가 내는 소리에 대한 음성정보를 획득하는 것이 바람직한데, 따라서 상기 거리측정 센서(230)에서 측정하는 거리정보에 따라 객체가 내는 소리를 증폭시켜 보다 선명하게 할 수 있도록 하는 소리 증폭 기능과, 상기 단말부(200) 주위의 잡음, 즉 노이즈를 제거할 수 있는 노이즈 제거 기능을 포함하는 것이 바람직하다.

[0096] 이때 상기 노이즈는 기 설정한 소리일 수도 있고, 또는 상기 단말부(200) 주위에서 패턴화된 채 연속적으로 나는 소리일 수도 있다. 예를 들어, 빗소리는 대표적인 노이즈로서 어느 정도 정형화, 패턴화시킬 수 있는 음향 정보이다. 따라서 이러한 빗소리는 미리 상기 음성 센서(250) 또는 단말기 제어부(270)상에 저장하여 음성 녹음 시 필터링에 이용될 수 있다.

[0097] 상기와 같이 기 설정된 노이즈 음향 정보가 음성 센서(250) 또는 단말기 제어부(270)상에 저장될 수 있으며, 또한 임시로 음성정보 획득을 위한 음성 녹음 시 반복적으로 들려오는 소리를 별도로 상기 음성정보 내에 구분 표시 저장하여 상기 서버부(100)에서 음성정보를 처리할 때 제거 여부를 결정하게 할 수도 있다.

[0098] 왜냐하면, 음성정보 획득 시 반복적으로 들려오는 소리는 노이즈일 수도 있지만 촬영 대상이 되는 객체가 발생시키는 소리일 수도 있기 때문에 일괄적으로 제거할 수 없기 때문이다. 따라서 상기 음성정보 저장 버퍼(1115)에 저장되는 음성정보 중 반복적으로 들려오는 소리는 구분 표시되어 있는 것이 바람직하고, 이는 상기 음성인

식부(150) 내에서 자동적으로 처리되거나 또는 상기 서비스 제공자(C)에 의해 처리될 수 있다.

- [0099] 그리고 상기 사용자 개별 데이터베이스(111)에서, 영상 데이터베이스(1116)는 실제로 사용자가 자신의 휴대형 단말부(300)를 통해 확인할 수 있는, 상기 카메라부(210)가 촬영한 영상과 음성 센서(250)를 통해 촬영된 음향이 결합된 영상 데이터(V1~V4...)로서, 하나 이상의 영상 데이터가 상기 영상 데이터베이스(1116)에 포함될 수 있다.
- [0100] 그리고 상기 영상 데이터(V1~V4)는 각각 영상 데이터 본체(V1a), 분석 플래그(V1b), 그리고 영상분석결과 데이터(V1c)를 포함하는 것이 바람직하다.
- [0101] 상기 영상 데이터 본체(V1a)는 상술한, 상기 설치형 단말부(200)의 구성요소들을 통해 획득할 수 있는 영상과 음성의 결합된 영상 데이터를 포함한다. 그리고 상기 분석 플래그(V1b)은 해당 영상 데이터 본체(V1a)가 상기 서버부(100)를 통해 분석되었는지 여부를 나타내는 표시자로서 분석이 완료된 데이터를 표시할 수 있다. 예를 들어 도 2에 도시된 바와 같이 분석 완료된 데이터의 경우 상기 분석 플래그(V1b)에 "1" 을 표시하고, 분석되지 않은 데이터의 경우 "0" 을 표시하여 구분할 수 있다.
- [0102] 그리고 상기 영상분석결과 데이터(V1c)는 해당 영상 데이터 본체(V1a)에 대하여 상기 서버부(100)를 통해 분석한 분석결과로서 사용자(P)가 휴대형 단말부(300)를 통해 확인할 수 있도록 가시적, 음향적으로 가공되어 해당 영상 데이터 본체(V1a)와 함께 제공되는 데이터이다. 여기서 해당 영상 데이터 본체(V1a)가 분석되지 않았다면, 해당 영상분석결과 데이터(V1c)는 아직 없는, 비어있는 상태(null)일 것이고, 해당 분석 플래그(V1b)는 "0" 이 기록되어 있을 것이다.
- [0103] 또한 상기 영상 데이터베이스(1116)에 저장된 각각의 영상 데이터(V1~V4...)와, 상기 버퍼(1112~1115)에 각각 저장되어 있는 정보 데이터가 서로 호환되어야 한다. 따라서 상기 버퍼(1112~1115)에 각각 저장되어 있는 정보 데이터들은 각각 자신이 어떤 영상 데이터에 관한 것인지를 표시하는 식별자(Vd)를 포함한다.
- [0104] 도 3은 본 발명의 시스템에서 설치형 단말부(200)가 자동 녹화를 시작하여 녹화한 영상 및 데이터를 전송하는 과정을 도시한 순서도이다. 이하에서는 도 1 내지 3을 통하여 본 발명의 시스템에서 설치형 단말부(200)의 동작 및 단말기 제어부(270)의 기능 및 동작에 대하여 설명한다.
- [0105] 상기 제어부(270)는 기기의 동작을 유지하기 위한 배터리 관리 소프트웨어와 블랙박스 녹화 소프트웨어를 포함한다. 상기 배터리 관리 소프트웨어는 상기 설치형 단말부(200)가 사용하는 통상의 내장 배터리 하드웨어를 관리하기 위한 시스템과 과충전 방지 회로를 관리하기 위한 프로그램으로서 배터리의 충방전 상태정보, 그리고 현재 충전상태에 따른 각 카메라부 및 센서부(210~250)의 계산된 동작예상 시간을 상기 사용자(P)의 휴대형 단말부(300)에게 제공하여 사용자(P)가 설치형 단말부(200)의 동작 상태를 알 수 있도록 한다.
- [0106] 그리고 상기 블랙박스 녹화 소프트웨어는, 상기 카메라부 및 센서부(210~250)를 동작시켜 각 버퍼(1112~1115)에 제공할 정보 데이터와 상기 영상 데이터베이스(1116)에 제공할 영상 데이터를 생성하기 위한 것으로서, 종래의 블랙박스는 아무런 이벤트가 없을 때도 항상 영상을 갱신 저장하기 때문에 하드웨어의 자원을 낭비하고 전력 소모를 많이 발생시키는 비효율적인 구조를 가진다는 단점을 가지고 있기에 이를 개선하고자 개발한 것이다.
- [0107] 따라서 본 발명의 시스템에서 블랙박스 녹화 소프트웨어는 사람, 즉 영상 녹화의 대상이 되는 객체가 일정 거리 이하로 접근하였을 때 녹화를 수행하고, 멀어졌을 때 녹화를 중지한다. 이에 대한 구체적인 순서가 도 3에 도시되어 있다.
- [0108] 도 3에 도시된 바와 같이, 객체가 상기 단말부(200)쪽으로 다가오기 시작하면, 실시간으로 항상 동작하고 있는 상기 거리측정 센서(230)에 거리정보가 입력(I1)되고, 따라서 카메라부(210)나 이미지센서(220) 또한 객체에 대한 이미지정보가 입력(I2)될 수 있다.
- [0109] 상기 입력(I1, I2)이 이루어짐에 따라 상기 제어부(260)는, 상기 거리측정 센서(230)에서 측정한 객체와의 거리값(D)이 기 설정된 임계거리값(Dt) 이하인지 측정하는 거리값 판단단계(S1)가 실시된다.
- [0110] 상기 단계(S1)에서, 상기 거리값(D)이 기 설정된 임계거리값(Dt)을 초과한다면, 이는 객체가 상기 단말부(200)에서 멀리 떨어져 있기 때문에 녹화의 대상이 되지 않는다는 것을 의미하므로, 처음으로 돌아가거나 단계를 종료할 수 있다.
- [0111] 그리고 상기 단계(S1)에서 입력된 거리값(D)이 기 설정된 임계거리값(Dt) 이하라면, 이는 객체가 상기 단말부(200)에서 기 설정된 거리 내로 들어와 녹화의 대상이 되었다는 것을 의미하므로, 다음 단계로 넘어갈 수 있다.

먼저, 입력된 거리값(D)이 0인지 먼저 판단하는 고장 판단단계(S2)가 실시된다.

- [0112] 만약 입력된 상기 거리값(D)이 0이면, 상기 거리측정 센서(230)가 오류이거나, 또는 객체가 거리측정 센서(230)가 설치된 부분에 완전히 밀착되어 있는 특이한 상황이 발생한 것이므로 오류 메시지를 출력하는 단계(S3)를 실시하여 상기 휴대형 단말부(300)의 알림설정부(330)에 상기 오류 메시지를 출력함으로써 사용자(P)가 이에 따른 조치를 취할 수 있도록 할 수 있다.
- [0113] 그리고 상기 입력된 거리값(D)이 기 설정된 임계거리값(Dt) 이하이고, 0이 아니라면, 객체가 기 설정된 임계거리 내에 들어와 상기 단말부(200) 근처에 있다는 것을 의미하는데, 먼저 해당 객체가 사람인지를 먼저 판단해야 한다.
- [0114] 따라서 현재 입력되고 있는 이미지정보(I2)에 사람이 포함되어 있는지 판단하는 객체 판단 단계(S4)가 실시된다.
- [0115] 여기서, 상기 객체 판단 단계(S4)는 단말기 제어부(270)가 실시할 수도 있지만, 바람직하게는 상기 단말기 제어부(270)가 해당 이미지정보(I2)를 서버부(100)에 제공하여, 상기 서버부(100)의 단순의심행동 감지부(170)가 수행 및 처리하도록 한다.
- [0116] 왜냐하면, 상기 단계(S4)는 객체가 사람인지를 판단하기 위하여 얼굴 및 체형을 인지해야 하기에 지능형 영상분석 기술이 필요하기 때문이다. 따라서 소형의 임베디드 시스템으로 구현되는 상기 단말기 제어부(270)가 처리하기에는 그 연산량이 많아 처리가 어렵기 때문에, 상기 서버부(100)가 처리하도록 하는 것이 바람직하다.
- [0117] 여기서 상기 단계(S4)를 처리하는 상기 단순의심행동 감지부(170)의 동작 방식 및 구성 등은 차후에 설명하기로 한다.
- [0118] 상기 단계(S4)의 실시에 따라, 객체가 사람이 아니라면 촬영의 필요가 없으므로 단계가 종료될 수 있다.
- [0119] 그리고 상기 단계(S4)의 실시에 따라 촬영된 이미지정보(I2)에 사람이 포함되어 있다고 판단되면, 상기 카메라부(210)는 영상녹화를 시작하여 녹화영상(V)을 만들어내는 녹화 단계(S5)를 실시한다.
- [0120] 이때 상기 단말기 제어부(270)에는, 최대 녹화시간(Rt)이 정해져 있으며 해당 녹화시간(Rt)동안 녹화를 한 촬영된 녹화영상(V1)을 서버로 전송하는 영상 전송 단계(S6)를 실시하여 단계를 마무리한다.
- [0121] 상기와 같이 최대 녹화시간(Rt)동안 녹화를 하는 이유는, 녹화시간이 길어지면 그만큼 상기 설치형 단말부(200) 내에 필요한 저장공간이 많아져 부담이 될 수 있고, 또한 녹화 시간이 긴 영상은 분석에도 어려움이 많기 때문이다. 따라서 상기 녹화시간(Rt) 단위로 분절하여 다수의 영상을 확보하는 것이 확인 및 분석이 용이하고 사용자(P)가 제공받기에도 편리하기 때문에 상기와 같이 한다.
- [0122] 그리고 상기 단계(S6)가 실시하여 단계가 종료하더라도, 사람이 계속 임계거리(Dt) 내에 들어와 있는 상태이면 지속적으로 거리정보(I1; D)가 입력되고 있으므로, 연속적으로 다음 녹화영상(V2, V3, V4...)을 촬영하여 실시간으로 연속되어 있지만 녹화시간(Rt) 단위로 분절되어 있는 다수의 녹화영상(V1, V2, V3, V4...)을 확보할 수 있다.
- [0123] 또한 상기 단계(S6)에서 서버로 전송하는 데이터는 상기 녹화영상(V1) 뿐 아니라, 녹화시간(Rt) 동안 수집한 버퍼정보를 포함한다. 상기 버퍼정보는 상기 버퍼(1112~1115)에 제공되기 위한 거리정보, 이미지정보, 자이로정보, 음성정보 중 선택된 어느 하나 이상을 포함한다.
- [0124] 도 3a는 본 발명의 설치형 단말부(200)에서 추가로 설치될 수 있는 측면감지 센서(260)의 동작을 도시한 개략도이다. 이하에서는 도 3a를 통하여 본 발명의 측면감지 센서(260)의 구성 및 동작에 대하여 설명한다.
- [0125] 본 발명의 측면감지 센서(260)는 상술한 바와 같이 설치형 단말부(200)에 포함될 수 있으며, 측면에서 다가오는 사람을 감지한다. 도 3a에 도시된 바와 같이, 상기 설치형 단말부(200)가 설치된 전면(Area 1)은 카메라부(210)나 이미지 센서(220) 중 선택된 어느 하나 이상의 구성에 의하여 용이하게 식별될 수 있으나, 카메라 등의 특성 상 자신의 바로 옆 부분은 별도의 구성이 없는 한 인식이 어려워 양 측면으로 음영지역(Area 2)을 생성할 수 있다.
- [0126] 따라서 상기 전면(Area 1)으로 다가오는 침입자(In1)의 경우 상기한 카메라부(210) 및 이미지 센서(220)를 이용한 도 3의 순서 절차에 따라 용이하게 인식할 수 있으나, 상기 측면 음영지역(Area 2)으로 접근하는 침입자(In2, In3)의 경우 상기 카메라부(210)나 이미지 센서(220)가 인식하기 어려워 보안의 공백이 생기는 문제가 발

생할 수 있다.

- [0127] 이를 위하여 상기 측면감지 센서(260)가 본 발명에 포함될 수 있다. 상기 측면감지 센서(260)는 통상의 적외선 센서 등으로 구성되어, 도 3b와 같은 방식을 통하여 측면으로 다가오는 사람을 인지한다. 상기 측면감지 센서(260)는 침입자(In2, In3)의 접근에 대하여 접근거리(Ds)가 입력되면(I3), 상기 침입자(In2)가 기 설정된 임계 접근거리(Dst) 내로 접근하였는지 확인하고(S11), 접근거리값(Ds)이 0인지 확인하여 오류를 인식한 후(S21), 상술한 도 3의 절차대로 진행(S4~S6)함으로서 측면 침입자(In2, In3)에 대한 식별 및 대응이 가능해진다.
- [0132] 이외에도 상기 단말기 제어부(270)는 상기 전기통신망(I)과의 통신을 위한 통신부를 더 포함할 수 있다. 이때 상기 통신부 및 통신 방식은 종래의 유무선 통신 방식을 따르면 되므로 이에 대한 설명은 생략한다.
- [0133] 그리고 상기 수동 입출력부(280)는 상기 설치형 단말부(270)에 부가될 수 있는 디스플레이 형태의 입출력부로서, 상기 사용자(P)가 설치형 단말부(200)의 제어부(260)에 대하여 각종 제어 설정 등을 변경할 수 있고, 또한 녹화된 영상을 확인할 수도 있다. 이러한 상기 수동 입출력부(280)는 종래의 일체식 디스플레이 및 모니터, 스위치의 구성으로 제공될 수도 있지만 터치스크린 방식의 일체형 입출력 디스플레이로 제공될 수도 있다.
- [0134] 또한 상기 수동 입출력부(280)는, 탈부착 가능한 형태로 제공되어 상기 휴대형 단말부(300) 중 어느 하나에 포함될 수도 있다. 이에 따르면, 사용자(P)는 상기 휴대형 단말부(300)로서 제공된 단말기를 상기 설치형 단말부(200)에 부착 연결하여 수동 입출력부(280)로 사용할 수도 있지만, 탈착하여 휴대형 단말부(300)로서 사용할 수도 있다.
- [0135] 상기와 같이 단말부(200)가 하나 이상의 녹화영상(V1) 및 버퍼정보가 서버에 제공될 수 있다. 이하에서는 상기 서버부(100)의 나머지 구성요소(120~180)에 대하여 설명한다.
- [0136] 상기 서버부(100)에서, 통신부(120)는 전기통신망(I)과의 통신을 위한 통상의 구성 및 방법을 통해 통신을 실시할 수 있다.
- [0137] 일례시로서, 상기 통신부(120)는 구글 프로토콜 버퍼(Google Protocol Buffer) 기반의 메시지 요청 응답 프로토콜이 포함되어 사용될 수 있다. 상기 구글 프로토콜 버퍼는 일반적인 JSON(JavaScript Object Notation) 형태 메시지보다 효율적이고 빠른 것으로 알려져 있으므로, 상기 통신부(120)에 포함되어 사용될 수 있다.
- [0138] 그리고 상기 서버부(100)에서, 이하에서 설명할 각 기능부(130~180)는 하나의 영상(V1)에 대하여 분석을 실시하여 영상 내 객체에 대한 범죄행동을 예측 및 판단하는 기능을 한다. 이하에서는 각 기능부(130~180)에 대한 개별 동작에 대하여 먼저 설명한다.
- [0139] 설명에 앞서, 상기 기능부(130~180)에서 분석하는 영상은 상기 영상 데이터베이스(1116) 내 제1영상(V1)을 예시로 하여 진행하며, 다른 영상(V2~V4...)에 대하여도 동일하게 진행된다.
- [0140] 상기 얼굴인식부(130)는 제1영상(V1)에 대하여 상기 이미지정보 저장 버퍼(1113)에 저장된 V1 이미지정보에 대하여 객체의 얼굴(V1f)을 식별 및 인식하고, 정확하게 인식하는지 평가하는 구성 및 동작, 기능을 포함한다. 상기와 같은 얼굴인식부(130)의 얼굴 인식 및 학습 방법에 있어서 잘 알려진 통상의 신경망(Neural Network) 얼굴 인식, 학습 방법 및 구성이 다수 개발되어 있는데, 통상의 방법 및 구성은 새로운 사람이 추가되었을 때 신경망 학습을 다시 해야 하여 과도한 데이터 사용 및 계산소요를 발생시킨다는 문제와, 얼굴 사진을 사용하였을 때 다른 사람으로 인식하는 문제가 발생하여 본 발명에서는 상기한 문제를 해결하면서 효과적인 얼굴 인식 및 학습을 위해 딥러닝 기반 객체 특성 추출 방법과, 딥러닝 기반 매칭 모델링 방법이 포함되어 사용된다.
- [0141] 상기한 딥러닝 기반 객체 특성 추출 및 매칭 모델링 방법에 대한 개념도가 도 4에 도시되어 있다. 도 4에 도시된 바와 같이, 상기 얼굴인식부(130)는 얼굴 특성추출 신경망(131)과, 얼굴 매칭 모델부(132), 그리고 얼굴인식 데이터베이스(133)를 포함한다.
- [0142] 따라서 상기 제1영상(V1) 내 객체의 얼굴(V1f) 또는 사진을 통한 얼굴 사진이 입력되었을 때, 상기 얼굴 특성추출 신경망(131)을 통하여 얼굴 객체에 대한 특성(V1ff)을 추출할 수 있는데, 추출한 특성값(V1ff)을 상기 얼굴 인식 데이터베이스(133)에 저장하고 해당 사람에 대한 ID를 부여한다. 만약 제1영상(V1)을 통해 객체의 얼굴(V1f)이 입력됨으로서 추출되는 얼굴 특성값(V1ff)에 대하여 상기 얼굴 매칭 모델부(132)는 상기 얼굴인식 데이터베이스(133)에 해당 얼굴 특성값(V1ff)과 맞는 ID의 사람이 있는지 판단하고, 없으면 상기 얼굴인식 데이터베이스(133)에 해당 얼굴 특성값(V1ff)을 저장하고 ID를 부여하며, 맞는 사람의 ID가 있으면 해당 ID의 사람과 동

일인임을 통보한다.

[0143] 상기와 같이 특성을 추출하여 저장함으로서, 신경망 학습을 다시 실시해야 함으로서 발생하는 여러 문제를 방지할 수 있다.

[0144] 또한 상기 얼굴 특성추출 신경망(131)은 입력되는 얼굴(V1f)이 사진인지 아닌지 판단하는 것도 중요한데, 사람의 얼굴은 항상 미세하게라도 움직이므로, 입력되는 얼굴(V1f)이 기 설정된 특정 시간동안 움직이지 않는다면 해당 얼굴(V1f)을 사진으로 판단하도록 한다.

[0145] 상기와 같은 얼굴 특성추출 신경망(131)의 예시로서, MTCNN(Multi-Task Cascaded Convolutional Neural Networks) 얼굴 검출 방법이 사용될 수 있다. 상기 MTCNN 얼굴 검출 방법은 아핀 변환(Affine Transform) 기반의 얼굴 특징점 검출(Face Alignment) 방식으로서 얼굴 영역 뿐 아니라 오른쪽 눈, 왼쪽 눈, 코, 입의 양쪽 점의 위치를 도 4a와 같은 형태로 제공한다.

[0146] 상기와 같이 서로 다른 위치의 5개 위치값과 얼굴 인식 표준 눈, 코, 입 위치는 서로 다를 수 있으므로, 두 위치의 정렬(Alignment)을 위하여 아래의 수식식 2와 같은 $Ax=b$ 형태의 아핀 변환 행렬이 상기 얼굴 특성추출 신경망(131)에 포함될 수 있다.

수학식 2

$$\begin{pmatrix} e_x^{il} & e_y^{il} & 1 \\ e_x^{ir} & e_y^{ir} & 1 \\ n_x^i & n_y^i & 1 \\ m_x^{il} & m_y^{il} & 1 \\ m_x^{ir} & m_y^{ir} & 1 \end{pmatrix} \begin{pmatrix} t_{11} & t_{12} \\ t_{21} & t_{22} \\ t_{31} & t_{32} \end{pmatrix} = \begin{pmatrix} e_x^l & e_y^l \\ e_x^r & e_y^r \\ n_x & n_y \\ m_x^l & m_y^l \\ m_x^r & m_y^r \end{pmatrix}$$

[0147]

[0148] 상기 수학식 2에서 $t_{11} \sim t_{33}$ 은 입력된 눈, 코, 입의 위치를 표준 위치로 변환해 주는 아핀 변환 행렬을 나타낸다. 입력된 눈, 코, 입의 위치와 표준 위치를 모두 알고 있으므로 $x=(AA^t)^{-1}AA^tB$ 형태의 LSM(Least Square Method; 최소제곱법) 풀이로 풀 수 있다. 이에 따라 입력 영상의 모든 픽셀에 아핀 변환 행렬을 적용하여 도 4b와 같은 정렬된 이미지를 생성할 수 있게 된다.

[0149] 그리고 상기 얼굴 매칭 모델부(132)의 예시로서, ArcFace 네트워크 기반 얼굴 인식 방법이 사용될 수 있다. 본 발명의 얼굴 매칭 모델부(132)에 사용하기 위한 네트워크 방식으로서 정확도와 속도를 고려하여 다수 평가하였고, 저해상도 및 적외선 영상에서도 사용 가능한지 여부를 고려하여 테스트한 결과 ArcFace 네트워크가 상기 얼굴 매칭 모델부(132)에 사용하기 가장 적합한 것으로 판단되었다.

[0150] 상기와 같이 하여, 상기 얼굴인식부(130)는 지속적으로 얼굴 사진 또는 얼굴이 포함된 영상을 제공받고 학습함으로서 제공된 영상에서 98% 이상의 확률로 얼굴을 인식하고 특정할 수 있게 된다.

[0151] [얼굴인식부(130) 성능 검증]

[0152] 상기와 같이 구성되는 얼굴인식부(130)의 얼굴인식 및 매칭 성능을 검증하기 위하여, 학습된 결과를 비교하여 자체적 성능검증을 실시하였다. 상기 얼굴 특성추출 신경망은 아핀 변환 기반 MTCNN 네트워크를, 얼굴 매칭 모델부(132)는 ArcFace 네트워크 기법을 적용하여 구현하였다.

[0153] 1차적으로 범용적 얼굴 특징을 추출하기 위해, 이미 얼굴 연구에 일반적으로 사용되고 있는 공개된 데이터베이스인 LFW, CFP-FP, AgeDB-30, IJB-C(E4) 얼굴 데이터베이스를 사용하여 1차 네트워크 학습을 수행하였다. LFW 13,000장, CFP-FP 500명의 정면사진 10장, 측면사진 4장의 총 7,000장, AgeDB-30의 경우 440명의 12,240장, IJB-C(E4)는 11,000명의 138,000장으로 구성된 총 170,240장의 얼굴이미지를 가지고 학습을 수행하였다.

[0154] 2차적으로 CCTV영상을 활용하여 네트워크 학습을 실시하였다. 공개된 CCTV영상 및 얼굴 클러스터링(Face Clustering) 기법을 이용하여 별도의 학습데이터 없이 얼굴영상을 추가 학습하였다. 이는 영상의 저해상도 얼굴에 대한 얼굴인식능력의 향상에 도움을 줄 것으로 판단하였다.

[0155] 그리고 3차로 흑백영상을 활용하여 네트워크를 학습시켰다. 본 발명의 시스템에서는 어두운 현관에서 여러 영상

처리를 하기 위해 적외선 카메라가 사용될 수 있으며, 가시광선 영역이 아닌 적외선 영역 영상의 경우 컬러 정보가 상대적으로 부족하므로 부족한 컬러 정보를 극복하기 위하여 네트워크에 컬러 정보 사용을 줄이도록 학습해야 할 필요가 있었다. 이를 위하여 영상의 밝기값만 사용하여 얼굴학습을 추가 진행할 수 있으며, 그레이스케일 변환 후 3차 학습을 진행하였다. 이에 따른 결과가 표 2에 도시되어 있는데, 모든 경우에서 95% 이상의 높은 식별 성능을 보였으며, 평균적으로 97.64%의 식별 성능을 확인하였다. 얼굴이 마스크와 같은 것으로 가려졌을 때와 같은 특수한 상황 외에는 성능이 저하되는 경우는 보이지 않았다.

표 2

Name	LFP	CFP-FP	AgeDB-30	IJB-C(E4)	CCTV	Grayscale Avg	Total Avg
AreFace ImageDei Train	99.83%	99.44%	99.31%	97.41%	95.10%	97.83%	97.64%

[0156]

그리고 상기 체형인식부(140)는 제1영상(V1)과 연동된 상기 이미지정보 저장 버퍼(1113)에 저장된 V1 이미지정보 내 객체의 체형을 인식하고 식별하기 위한 기능부로서, 우선 체형의 식별을 위하여 상기 체형인식부(140)는 도 5에 도시된 바와 같이 한국인 남성 체형을 5가지로 구분(작은 역삼각체형, 큰 사각체형, 역삼각체형, 작은 사각체형, 표준체형)하여, 상기 5가지 체형을 학습할 수 있도록 전신사진이 제공되어 레이블이 부여됨으로서 학습데이터를 구축시킬 수 있도록 한다.

[0157]

이때 상기 체형인식부(140)의 학습을 위하여 제공되는 데이터셋인 전신사진 학습데이터는 종래의 딥러닝 학습에서 사용하는 수십만장 단위의 사진 데이터셋을 제공하기는 어렵고, 수천장 단위의 사진 데이터셋을 사용하게 될 것으로 예상되는데 이에 따른 부족한 학습데이터를 극복하기 위하여, 상기 체형인식부(140)는 VGG16(16계층 심층신경망(VGGnet)) 기반의 전이학습(Transfer Learning) 방법이 적용된다.

[0158]

도 6에 VGG16 네트워크 구조에 대한 개념도가 도시되어 있다. 도 6를 통해 구체적으로 설명하면, 상기 체형인식부(140)는 하나 이상의 시각적 데이터셋, 예를 들어 이미지넷(ImageNet; 시각적 데이터베이스)을 학습시킨 VGG16 네트워크에서 기존 네트워크의 학습 가중치는 그대로 놓되 도 6에 도시된 바와 같이 마지막 노드 $1 \times 1 \times 1000$ 을 $1 \times 1 \times V_n$ 으로 변경하여 V_n 가지 체형을 분류할 수 있도록 변경한다.

[0159]

여기서 상기 상수 V_n 은 체형분류 가짓수로서 체형 분류를 위하여 상기 서버(S)내에 설정된 값이다. 상기 상수 V_n 값은 분류되는 체형의 구분수가 되는데, 예를 들어 상기와 같이 남성 체형을 5가지로 분류할 경우 상기 상수 V_n 값은 5가 된다. 마찬가지로 상기 남성 체형을 분류하는 수에 따라서 V_n 값은 변할 수 있는데, VGG16 네트워크의 일반적 구성의 마지막 노드인 1×1000 보다 적은 수의 노드 구성을 위하여, 상기 V_n 값은 1000 미만일 되도록 하는 것이 바람직하다.

[0160]

상기와 같이 변경된 VGG16 네트워크에 본 발명의 시스템에서 체형인식에 사용할 수 있는 준비된 체형 관련 데이터셋을 학습시킴으로서 상기 체형인식부(140)가 기능할 수 있도록 한다. VGG16은 상기한 전이학습이 잘 동작하는 네트워크로 이미 잘 알려져 있기 때문에 상기와 같이 하였으며, 본 발명의 출원인이 시험해본 결과 상기 체형 관련 데이터셋을 수천장 단위로 준비하여 학습시켰는데도 불구하고 95% 이상의 정확도로 체형을 분류할 수 있는 것이 확인되어 고성능의 체형학습 모델로서 기능할 수 있음을 확인하였다.

[0161]

[체형인식부 성능 검증]

[0162]

상기 체형인식부(140)의 성능검증을 위하여, 5가지 체형데이터 분류를 사용하고, VGG1-16BN 기반 체형인식 네트워크를 상기 체형인식부(140)에 적용하여 구현하였다. 성능검증에는 10,000장의 학습데이터를 사용하였는데 VGG-16BN 네트워크의 경우 이미 ImageNet의 5,000,000장의 이미지들이 시각적 중요 정보로서 학습되어 있으므로 이를 사용하되 네트워크 부분만 가중치 값을 0으로 두어 기존 학습 가중치 값을 활용하는 전이학습을 수행하였다.

[0163]

10,000장의 이미지를 학습하여 비교한 결과가 도 6a에 도시되어 있다. 도 6a에 도시된 바에 따르면 약 87.6%의 정확도를 가지는 식별 능력을 확인할 수 있었으며, 전처리 및 라벨링 오류를 보강하여 정확성을 더 높일 수 있기 때문에 본 발명에서 체형인식부(140)는 유효한 고성능의 체형학습 모델로서 본 발명에 포함될 수 있는 것을 확인하였다.

[0164]

- [0165] 그리고 상기 음성인식부(150)는 제1영상(V1)과 연동된 상기 음성정보 저장 버퍼(1115)에 저장된 V1 음성정보에서 객체가 내는 소리를 식별하기 위한 기능부로서, 객체가 내는 소리를 특정해 추출하기 위하여 종래의 잘 알려진 파이썬 리브로사(Python Librosa) 라이브러리를 사용하여 MFCC(Mel Frequency Cepstral Coefficient; 소리 데이터의 특징값)를 추출해 사용하거나, 또는 딥러닝 기반의 음성인식 모델을 준비하여 사용할 수 있다. 상기의 기법은 하나 이상 선택하거나, 또는 둘 다 사용할 수 있는데, 예를 들어 파이썬 리브로사 라이브러리를 데이터 셋으로 사용하는 딥러닝 음성인식 모델을 사용할 수도 있다. 상기한 파이썬 리브로사를 사용한 음성 데이터 분석기법은 잘 알려진 방법으로서 이에 대한 설명은 생략한다.
- [0166] 그리고 딥러닝 기반의 음성인식 모델 또한 종래의 잘 알려진 음성인식 모델 중 어느 하나 이상을 선택하여 사용하면 되는데, 음성 및 오디오 데이터가 1차원 시계열적 특성을 가지고 있다는 특징을 가짐에 따라 과거와 미래 데이터를 결합하여 Overlap-shift 방식으로 처리되거나 HMM(Hidden Markov Model; 은닉 마르코프 모델)기반의 DBN(Deep Belief Network; 심층신뢰망)을 사용하거나, 또는 은닉층에 루프를 추가한 RNN(Recurrent Neural Network; 순환신경망)이 사용될 수 있다. 이외에도 본 발명의 음성인식부(150)에서 사용될 수 있는 음성인식 방법 및 모델은 다양한 형태로 개발되어 있으므로 상기 서비스 제공자(C)는 종래의 방식 중 어느 하나 이상을 선택하여 사용할 수 있다.
- [0167] 또한 상기 음성인식부(150)는 상기와 같은 방식으로 수동적인 음성인식을 수행하는 음성인식 모듈(151) 뿐 아니라 내부적으로 능동적인 자연어 처리부(152)와, 유의어 데이터베이스(153)를 더 포함할 수 있다.
- [0168] 도 7은 상기 음성인식부(150)가 음성인식 모듈(151), 자연어 처리부(152), 그리고 유의어 데이터베이스(153)를 모두 포함하고 있을 때의 구조도이다. 이하에서는 도 7을 통하여 상기의 구성요소(151~153)에 대하여 설명한다.
- [0169] 만약 상기 음성인식부(150)에 음성인식 모듈(151)만 포함된다면, 입력되는 V1 음성정보에 대하여 인식 및 처리하여 다시 상기 음성정보 저장 버퍼(1115)에 갱신 저장하면 된다. 하지만 상기 자연어 처리부(152)가 상기 음성인식부(150)에 포함되면, 상기 음성인식 모듈(151)이 처리하여 특정한 V1 음성정보 중 객체의 소리에 대하여, 상기 자연어 처리부(152)가 상기 유의어 데이터베이스(153)를 참조하여 상기 객체의 소리 중 특정된 유의어가 있는지 검색한다.
- [0170] 상기 유의어 데이터베이스(153)에는 기 입력된 유의어가 하나 이상 저장되어 있는데, 상기 유의어는 범죄와의 연관성이 높을 것으로 판단되는 비속어, 고성 등이 포함된다. 만약 특정된 객체의 소리 중 상기 유의어가 하나 이상 포함되어 있으면, 범죄와의 연관성을 높은 확률로 생각하여 볼 수 있다.
- [0171] 따라서 예를 들어, 상기 자연어 처리부(152)는, 상기 음성인식 모듈(151)이 처리한 객체의 소리에 대하여 분석하였을 때 상기 유의어 데이터베이스(153)에 기 저장되어 있는 비속어 "@@" 와 "%%" 가 포함되었음을 알게 되었다면, 이를 상기 음성정보 저장 버퍼(1115)에 추가로 갱신 저장하면서 V1 영상에 대한 범죄 인식 판단에 사용될 수 있다.
- [0172] 이때 바람직하게는, 상기 자연어 처리부(152)에 기 설정된 유의어 반복횟수 w 가 설정되어 있어, 상기 V1 음성정보에 상기 유의어 데이터베이스(153)에 저장되어 있는 비속어 단어가 w 회 이상 포함되어 있다면, V1 영상에 대하여 범죄 의심 상황으로 판단할 수 있기에 상기 범죄행동 예측부(180)를 통하여 필요한 조치를 취하도록 한다.
- [0173] 상기와 같은 자연어 처리부는 통상의 자연어 처리(NLP) 기법 또는 학습모델 중 선택된 어느 하나 이상을 사용할 수 있다. 전통적으로 많이 쓰이는 통계적 언어 모델(SLM)이나 신경망 언어 모델(NLM), 순환 신경망, 장단기기억망 적용 순환 신경망(RNN-LSTM), 양방향 순환신경망(Bidirectional RNN), GRU(Gated Recurrent Unit; 게이트 순환 유닛), 어텐션 모델(Attention Model)과 같은 다양한 방법 중 선택된 어느 하나 이상의 방법을 포함하여 상기 자연어 처리부(152)가 구현될 수 있다.
- [0174] 그리고 서버부(100)의 나머지 구성요소들(160, 170, 180)은 상기 인식부(130~150)에서 처리한 각 버퍼(1112~1115) 및 영상 데이터베이스(1116)의 연동된 데이터 및 영상에 대하여 범죄 행위가 있는지 파악하고 데이터를 생성하는 부분이다.
- [0175] 먼저 상기 단순의심행동 감지부(170)는 상기 버퍼(1112~1115)에 포함된 하나 이상의 데이터를 가지고 영상에서 파악되는 객체가 행동을 하는 사람인지 여부를 판단하기 위한 것으로서, 상술한 바와 같이 상기 단계 S4 에서 상기 단순의심행동 감지부(170)가 관여하여 이미지정보에 사람이 포함되어 있는지를 판단한다. 이하에서는 상기 단순의심행동 감지부(170)에 대하여 설명한다.
- [0176] 상기 단순의심행동 감지부(170)는 실시간 객체탐지 방법을 사용하여 얼굴 및 체형을 객체로서 탐지한다. 이때

상기 단계 S4가 빠르게 진행되어야 할 필요가 있으므로, 상기 단순의심행동 감지부(170)는 정밀한 얼굴 및 객체 파악보다는 신속하고 효율적인 파악에 중점을 두어 실시간 객체 탐지 방법으로서 빠른 속도와 성능을 인정받고 있는 YOLO 네트워크를 사용한다. 본 발명의 출원인은 YOLOv5(YOLO Version 5)를 통하여 상기 단순의심행동 감지부(170)를 구현하여 테스트한 결과, 빠른 속도와 더불어 95% 이상의 확률로 머리와 신체를 탐지할 수 있는 것을 확인하였다.

[0177] [단순의심행동 감지부 성능검증]

[0178] 상기 단순의심행동 감지부(170)의 성능검증을 위하여, YOLOv5 기반 단순의심행동 감지부(170)를 구현하되, 사람 영역만 필요하므로 상기 YOLOv5 네트워크의 일반적인 구성을 변형하여 사용하였다. 데이터셋은 일반적으로 제공되는 COCO 데이터셋 중 사람 영역이 포함된 66,808장의 이미지로 학습하였고, 학습 결과는 도 17에 도시하였다. 도 17에 도시된 바와 같이, 학습 결과 8,000번의 반복 비교에서 학습데이터 기준 mAP 0.5의 성능이 98.9%로 매우 높은 적중률을 보이는 것이 확인되어 높은 정확도로 머리와 신체를 탐지할 수 있는 것이 확인되었다.

[0179] 그리고 상기 범죄행동 예측부(180)는 상기의 인식부(130~150)에서 분석한 내용 및 상기 버퍼(1112~1115) 및 영상 데이터베이스(1116)에 포함되어 있는 정보 데이터 및 영상 데이터를 바탕으로 객체의 범죄행위를 판단하고 예측하기 위한 기능부이다. 이하에서는 상기 범죄행동 예측부(180)에 대하여 설명한다.

[0180] 상기 범죄행동 예측부(180)에서 먼저 선행되어야 하는 것은 어떤 행동을 범죄행동으로 설정할 지 결정하는 것이다. 예를 들어, 본 발명의 출원인은 상기 범죄행동 예측부(180)에 대하여 범죄행위를 학습시킴에 있어서 650,000개의 유튜브 영상을 대상으로 약 700개의 사람 행동을 분류한 Kinetics 700 데이터셋을 사용하였고, 이에 따라 도 8과 같은 4가지 행동(문 열기, 물건 배송, 응시, 철사 구부리기)을 범죄행동으로 설정하였다.

[0181] 상기의 4가지 행동 중 물건 배송의 경우, 충분히 의심스러운 행동일 수 있지만 일반적인 택배 배달 행위일 수도 있다. 하지만 택배 배달 행위 역시 긍정적인 의미에서 사용자(P)에게 빠르게 통보해 줄 필요가 있는 행위이므로, 상기 휴대형 단말부(300)를 통해 객체의 물건 배달 행위가 있었음을 통보하면 된다.

[0182] 그리고 상기 설치형 단말부(200)가 설치되어 있거나 또는 근접해 있는 문을 열거나 또는 문을 열려고 시도하는 행위, 문이 열려있는지 확인하는 행위(이하 문 열기), 문이나 문 주변을 지속적으로 확인하고 쳐다보는 행위(이하 응시), 그리고 근처에서 철사 등을 구부리거나 구부러진 철사 등을 열쇠 구멍 등에 넣는 행위(이하 철사 구부리기)는 충분히 의심스러운 범죄 예상 행동이므로 상기 범죄행동 예측부(180)는 상기한 4가지 동작이 감지되었을 경우 즉시 상기 휴대형 단말부(300)를 통해 상기 행동들 중 선택된 어느 하나 이상이 감지되었음을 통보한다.

[0183] 또한 상기 4가지 행동, 그리고 상기 4가지 범죄 예상 행동과 연관된 행동을 예측하는 모델을 더 포함하여 특정된 4가지 행동 외에 상기 4개 행동과 연관된 행동까지 파악할 수 있도록 하는 것이 바람직하다. 따라서 상기 범죄행동 예측부(180)는 이미 사람 행동 예측에 많이 사용되고 있는 ResNet 3D의 18 Layer를 가진 Pre-Trained Model에 기반한 행동예측 모델을 포함한다.

[0184] 그리고 상기 범죄행동 예측부(180)에 있어서 중요한 점은 상기 4가지 행동을 객체에게서 파악하는 것이다. 따라서 본 발명의 범죄행동 예측부(180)는 상기 버퍼(1112~1115)에 저장된 정보 데이터 중 선택된 어느 하나 이상을 이용하여 상기 4가지 행동의 포함 여부 및 범죄 의심 행동이 영상에 포함되는지 여부를 판단한다.

[0185] 여기서 상기 4가지 행동의 포함 여부는 상기 자이로정보 저장 버퍼(1114)에 저장되어 있는 자이로정보를 통해 판별할 수 있다.

[0186] 상술한 바와 같이 상기 자이로정보는 6개의 데이터가 초당 n번 이상 연속적으로 측정된 시계열적인 자이로 시퀀스 데이터로서 제공된다. 이를 해석하기 위하여, 상기 범죄행동 예측부(180)는 하나 이상의 행동 해석 모델이 포함된다.

[0187] 상기 행동 해석 모델로서 사용될 수 있는 방법 중 하나는 순환 신경망이다. 순환 신경망은 상술한 바와 같이 연속적인 행동을 해석함에 있어서 좋은 성능을 가지고 있기 때문에 선택하여 사용할 수 있는데, 상기 순환 신경망은 학습 데이터의 레이블을 선택하기 어렵거나, 또는 레이블 선택의 편의를 위하여 많은 연산량이 필요한 CTC(Connectionist Temporal Classification) loss를 사용해야 하는 문제점이 발견되었다. 또한 학습 때 사용한 Batch size나 시퀀스 길이에 따라 loss의 크기가 줄어들지 않고 요동쳐 학습에 시간이 오래 걸린다는 단점 또한 발견되었다.

[0188] 따라서 상기 행동 해석 모델로서, 상기 자이로 시퀀스 데이터를 일정 량 모아 정형화시킨 후 통상의

CNN(Convolutional Neural Network; 합성곱 신경망) 중 선택된 어느 하나 이상, 예를 들어 VGG8(VGG-8)을 사용하여 구현하는 것이 바람직하다.

- [0189] 도 9에 상기 행동 해석 모델로서, 통상의 CNN 중 VGG8과 정형화된 자이로 시퀀스 데이터를 사용하여 구현한 행동 해석 모델의 개념도가 도시되어 있다. 이를 통해 설명하면, 먼저 상기 자이로 시퀀스 데이터에서의 초당 연속 측정 횟수 n 이 10으로 설정되어 있고 3초간의 자이로 시퀀스 데이터를 해석한다고 가정하였을 경우 6개의 데이터를 사용하므로, $6(\text{데이터 수}) \times 3(\text{초}) \times 10(\text{횟수})$ 도합 180개의 특징을 가진 특징 벡터를 생성할 수 있다. 상기 특징 벡터는 도 9에 도시된 바와 같이 단위 시간 만큼 오른쪽으로 시프트하여 중첩된 특징 벡터를 단위 시간마다 생성할 수 있다.
- [0190] 상기 행동 해석 모델은 또한, 2가지 상황을 가정하여 일상 상황은 Class 0, 범죄 의심상황을 Class 1로 설정 구분하여 상기 중첩 생성된 특징 벡터에 대하여 상기한 VGG8 네트워크를 사용하여 학습 및 판단을 함으로서 상기 범죄행동 예측부(180)가 행동을 감지하고 감지된 행동이 상기 4개의 범죄 의심 행동에 부합하는지 여부를 판단하여, 감지된 행동이 4개의 범죄 의심 행동으로서 예상되면 해당 영상 V1을 Class 1으로 설정하고 상기 휴대형 단말부(300)에 통보하여 사용자(P)가 이를 인지할 수 있도록 한다.
- [0191] 상기와 같이 함으로서, 영상 내에서 4개의 범죄 의심 행동을 최소 80% 이상의 확률로 예측할 수 있는 것으로 파악하였다.
- [0192] 또한 상기 범죄행동 예측부(180)는 상기 통보를 휴대형 단말부(300)에 하는 것 뿐 아니라, 상기 기관(G)에 직접 통보하여 빠른 조치를 취할 수 있도록 할 수 있다. 상기와 같이 기관(G)에 직접 통보할 경우, Class 1으로 설정된 해당 영상 V1 뿐 아니라, 해당 사용자(P)의 출동지점 주소(1111)에 저장되어 있는 주소 지점을 함께 제공하여 특정 위치에 대하여 기관(G)이 방문 등의 조치를 취할 수 있게 해야 한다.
- [0193] 그리고 상기 범죄행동 예측부(180)를 통하여 Class 1으로 설정된 영상에 대하여, 상기 3D영상 생성부(160)가 객체, 즉 사람에 대한 3D 몽타주 영상을 생성한다. 이하에서는 상기 3D영상 생성부(160)의 기능 및 동작 방법에 대하여 설명한다.
- [0194] 상기 3D 영상 생성부(160)에 제공되는 영상 V1 및 이미지정보는 모두 2D 정보이므로, 이를 기반으로 하여 3D로 변환 생성해야 하는데 따라서 상기 3D 영상 생성부(160)는 얼굴 모델링 영상을 생성하기 위한 3D 얼굴영상 생성 모델과, 체형 모델링 영상을 제공하기 위한 3D 체형 생성모델을 포함한다.
- [0195] 상기 3D 얼굴영상 생성 모델은 종래의 인공지능 네트워크 중 선택된 어느 하나 이상을 사용하여 구현 및 학습시킬 수 있는데, 본 발명에서는 일례로서 Deep3DFace 데이터셋 및 모듈을 사용하여 상기 3D 얼굴영상 생성 모델을 구현하고 학습시켰다.
- [0196] 또한 상기 3D 체형 생성모델은, 상기한 5개의 표준 체형 모델에 맞춘 5개의 체형 모델링 영상이 포함된다. 따라서 상기 체형인식부(140)가 인식한 5개의 표준 체형 모델 중 어느 하나에 대하여 대응되는 체형 모델링 영상이 제공된다.
- [0197] 따라서, 상기 3D 얼굴영상 생성 모델에 따라 생성된 얼굴 모델링 영상과, 5개의 표준 체형 모델 중 어느 하나에 대응되는 체형 모델링 영상을 합성하여 도 10과 같은 형태의 완성된 3D 모델링 영상을 생성하여 상기 휴대형 단말부(300) 또는 기관(G) 중 선택된 어느 하나 이상에게 제공하거나 또는 참조할 수 있도록 상기 서버부(100)에 저장된다.
- [0198] 상기와 같은 구성을 포함하여 상기 서버부(100)의 각 구성요소들은 상기 설치형 단말부(200)에서 제공된 영상을 인식하고 범죄행동을 예측하여 제공할 수 있는데, 보다 신속한 인식 및 처리, 그리고 서버부(100) 내 포함되는 하나 이상의 연산장치의 처리 부하를 경감하기 위하여, 상기 각 기능부(130~180)에 포함된 하나 이상의 기능부 및 모듈, 모델들 중 어느 하나 이상은 통상의 소스 코드, 예를 들어 파이썬 코드 형태로 저장되는 대신 ONNX(Open Neural Network eXchange) 학습데이터 형태로 변환되어 저장될 수 있다. 이에 따라, ONNX 런타임 모듈을 활용하여 학습 시 사용된 프레임워크 없이 모든 변환된 기능부 및 모듈, 모델들을 원활히 실행할 수 있도록 상기 서버부(100)는 ONNX 런타임 라이브러리를 포함하는 것이 바람직하다.
- [0199] 도 11은 본 발명의 범죄 예방 시스템의 제2형태 구조도이다. 이하에서는 도 11을 통하여 본 발명의 범죄 예방 시스템의 제2형태 구성에 대하여 설명한다.
- [0200] 도 11에 도시된 범죄 예방 시스템의 제2형태는, 도 1의 범죄 예방 시스템에서 각 개인의 개별 데이터베이스(111...)마다 포함되어 있던 영상 데이터베이스(1116)가 외부로 나와 통합되어 외부 영상 데이터베이스(1115')

를 형성하고 있는 형태이다.

- [0201] 상기와 같이 하는 이유는 상기 서버부(100)와 다수의 사용자가 통신을 주고받음에 따른 상기 통신부(120)의 회선 과부하 등으로 인하여 처리속도가 지연될 수 있기 때문에, 고용량의 영상 데이터를 보관할 수 있는 별도의 영상 데이터베이스(1115')를 외부에 마련함으로써 회선 부하를 줄일 수 있기 때문이다. 게다가 현대에는 보안성과 성능, 사용성을 만족하고 있는 클라우드 서버 서비스가 다수 제공되고 있으므로, 상기와 같이 외부 영상 데이터베이스(1115')를 타사의 클라우드 서버 서비스를 통해 달성할 수도 있다.
- [0202] 상기와 같은 외부 영상 데이터베이스(1115')는 각각의 사용자에 대한 개별 영상 데이터베이스(P1V, P2V, P3V...)가 포함되며, 또한 상기와 같은 하나 이상의 데이터베이스들을 운용할 데이터베이스 운용부(0)가 포함된다. 상기 데이터베이스 운용부(0)는 통상의 데이터베이스 운영체제를 사용하면 되므로 이에 대한 설명은 생략한다.
- [0203] 그리고 각 사용자 데이터베이스(P1V, P2V, P3V...) 각각에 저장되는 영상 데이터(V1, V2, V3...)의 구성 및 내용은 상기한 영상 데이터베이스(1116)와 동일하므로 이에 대한 설명은 생략한다.
- [0204] 상기와 같은 외부 영상 데이터베이스(1115')는 상술한 바와 같이 외부 클라우드 서버 서비스를 사용하여 달성할 수도 있다. 예를 들어, 파일 공유 서비스로서 사용되고 있는 오픈소스인 넥스트클라우드 서버(NextCloud Server)의 웹 클라이언트 서버 서비스를 통해 상기 외부 영상 데이터베이스(1115')를 달성할 수 있다.
- [0205] 도 13은 상기 얼굴인식부(130)가 더 포함할 수 있는 얼굴 비교 모델(131)의 동작 구조도이다. 이하에서는 도 13을 통하여 상기 얼굴 비교 모델(131)의 기능 및 동작에 대하여 설명한다.
- [0206] 상술한 바와같이, 상기 얼굴인식부(130)는 얼굴 특성추출 신경망(131)과, 얼굴 매칭 모델부(132), 그리고 얼굴 인식 데이터베이스(133)를 포함하여 이미지정보 데이터 또는 영상 데이터 중 선택된 어느 하나로부터 이미지 또는 영상에 얼굴 객체가 있는지 찾아내고, 특성을 추출하여 그 추출된 특성을 얼굴인식 데이터베이스(133)에 구분 저장한다. 따라서 같은 인물이 영상에 다시 등장할 경우, 상기 얼굴인식 데이터베이스(133)를 통하여 동일 인물임을 상기 얼굴 매칭 모델부(132)가 인지하고 빠른 연산 및 판단이 가능해진다.
- [0207] 여기서 본 발명의 시스템은, 상기 얼굴인식부(130)에 얼굴 비교 모델부(134)를 더 추가하여, 범죄 상황을 더 정밀하게 예상할 수 있다. 구체적으로는, 도 13에 도시된 바와 같이, 상기 이미지정보 데이터 또는 영상 데이터로부터 추출된 얼굴이 잘 알려진 범죄자이거나 실종자 등이거나, 그 사람의 등장만으로도 범죄상황이 예상되거나, 또는 범죄상황이 아니더라도 상기 기관(G)에 통보해야 할 상황이 되는 것이다.
- [0208] 따라서 본 발명의 시스템은 상술한 바와 같이, 상기 서버부(100)가 상기 기관(G) 내 DB(GDB)에 전기통신망(I)을 경유하여 참조 가능하게 연결될 수 있다. 상기 기관 DB(GDB)는 공개수배중인 범죄자의 DB일 수도 있고, 또는 실종자의 DB일 수도 있다. 이는 정부에서 제공하는 DB의 성격에 따라 바뀔 수 있다.
- [0209] 또는 본 발명의 서버부(100)가 별도의 신원정보 DB(190)를 더 포함하여, 상기 기관 DB(GDB)로부터 정기적으로 공개수배중인 범죄자나 실종자(이하 수배자(C))의 얼굴 사진 및 인적사항 등의 정보를 갈무리하여 저장할 수 있다. 상기와 같이 신원정보 DB(190)가 상기 서버부(100)에 더 포함되어 구성될 경우, 상기 서버부(100)는 상기 기관 DB(GDB)로부터 정기적으로 상기 수배자(C)의 정보를 참조하여 갈무리해 저장할 수 있는 크롤링 프로그램을 더 포함하는 것이 바람직하다.
- [0210] 상기 기관 DB 또는 신원정보 DB(GDB / 190)에는 한 명 이상의 수배자(C)의 사진이 등록(C1~C4...) 저장되어 있다. 상기 얼굴비교 모델부(134)는, 상기 얼굴 특성추출 신경망(131)을 통해 현재 이미지정보 또는 영상 데이터에서 추출된 특징 추출된 얼굴(GF)과, 상기 수배자(C1~C4...)를 비교하여 유사도를 평가한다.
- [0211] 만약 상기 얼굴비교 모델부(134)가 유사도를 평가함에 있어서, 상기 얼굴 특성추출 신경망(131)을 통해 현재 추출된 특징 추출된 얼굴(GF)과, 상기 기관 DB 또는 신원정보 DB(GDB / 190)에 저장되어 있는 어느 수배자(C4)와 기 입력된 임계 수준 유사도 G% 이상으로 유사하다면, 상기 특징 추출된 얼굴(GF)의 해당 인원이 수배자 C4로 판단하고 즉시 범죄행동 예측부(180)를 통하여 연동된 영상 데이터 V1을 Class 1로 설정한 뒤 상기 휴대형 단말부(300)와 기관(G)에 통보함으로써 사용자(P)와 기관(G)이 필요한 조치를 취할 수 있도록 한다.
- [0212] 도 14는 본 발명의 시스템에서 더 추가될 수 있는 상기 얼굴비교 모델부(134)와 음성인식부(150)의 구성 및 동작에 대한 것이다. 이하에서는 도 14를 통하여 본 발명의 시스템에서 더 추가될 수 있는 개별 데이터베이스(111)의 구성요소와 얼굴인식부(130) 및 음성인식부(150)의 구성 및 기능에 대하여 설명한다.

- [0213] 도 14에 도시된 바와 같이, 상기 개별 데이터베이스(111)는 등록얼굴 데이터베이스(1117)와, 등록음성 데이터베이스(1118)가 더 추가될 수 있다.
- [0214] 상기 등록얼굴 데이터베이스(1117)는, 해당 사용자(P)가 제공하는 한 명 이상의 경계인원(CP1, CP2...)의 사진 및 인적사항 등의 정보가 저장되는 저장공간이다.
- [0215] 상기 경계인원(CP1, CP2...)은 상기한 수배자(C)와 같이 객관적으로 위험하거나 유의해야 할 사람은 아니지만, 해당 사용자(P)에게 위협이 될 수 있는 사람으로서 상술한 바와 같이 상기 경계인원(CP1, CP2...)의 사진 및 인적사항은 상기 사용자(P)가 등록하여 자신의 개별 데이터베이스(111) 내 등록얼굴 데이터베이스(1117)에 저장시킨다. 따라서 상기 등록얼굴 데이터베이스(1117)가 할당되어 있다고 하더라도, 상기 사용자(P)가 경계인원(CP1, CP2...)을 등록시키지 않는다면 상기 등록얼굴 데이터베이스(1117)는 공란이 될 것이다.
- [0216] 그리고 상기 얼굴 인식 모델(131)은 상기한 얼굴 특성추출 신경망(131)을 통해 추출된 특징 추출된 얼굴(GF)과 상기 등록얼굴 데이터베이스(1117)의 경계인원(CP1, CP2...)의 어느 특정 경계인원 CP2가 임계 수준 유사도 G% 이상으로 유사하다면, 상기 특정 추출된 얼굴(GF)의 해당 인원이 경계인원 CP2로 판단하고 즉시 범죄행동 예측부(180)를 통하여 연동된 영상 데이터 V1을 Class 1로 설정한 뒤 상기 휴대형 단말부(300)와 기관(G)에 통보함으로써 사용자(P)와 기관(G)이 필요한 조치를 취할 수 있도록 한다.
- [0217] 또한 상기 등록음성 데이터베이스(1118)는 상기 사용자(P)가 등록한 하나 이상의 음성파형(W1, W2...)이 구분 저장되는 저장공간이다.
- [0218] 상기 음성파형(W1, W2...)은 상기한 유의어 데이터베이스(153)에 등록된 비속어 등과 같이 공통적으로 범죄와 연관되었으리라 예상되는 단어들은 아니지만, 해당 사용자(P)가 위협으로 느낄 수 있는 단어이거나 또는 목소리, 성문(聲紋)을 추출한 것이다. 이는 해당 사용자(P)가 위협으로 느낄 수 있는 특정한 단어이거나 문장일 수도 있고, 위협으로 느낄 수 있는 특정인의 목소리나 성문일 수도 있다.
- [0219] 상기 음성인식부(150) 내 포함될 수 있는 자연어 처리부(152)는 상기 영상 V1의 음성정보에 대하여 유의어가 w회 반복되는지에 있어서, 상기 유의어 데이터베이스(153) 뿐 아니라 상기 등록음성 데이터베이스(1118)를 같이 참조하여, 상기 음성정보 V1을 분석함에 있어 상기 유의어 데이터베이스(153)에 등록된 단어와 상기 등록음성 데이터베이스(1118)에 등록된 단어 또는 문장의 합이 w회 이상 드러난다면, 이를 범죄 예상 상황으로 판단하고 해당 영상 데이터를 Class 1로 설정하여 상기 휴대형 단말부(300)와 기관(G)에 통보함으로써 사용자(P)와 기관(G)이 필요한 조치를 취할 수 있도록 한다.
- [0220] 그리고 상기 음성인식부(150)가 상기 음성정보 V1을 분석함에 있어서, 특정 단어나 문장이 아니라 등록음성 데이터베이스(1118)에 등록된 목소리나 성문이 상기 음성정보 V1과 겹쳐지게 인식한다면, 상기 횟수 w를 초 단위로 실시간 증가시켜 일정 길이 이상이 충족된다면 이를 범죄 예상 상황으로 판단하여 상기한 조치를 취할 수 있도록 한다.
- [0221] 도 15는 상기 설치형 단말부(200)에 더 포함될 수 있는 진동감지 센서(280)에 대한 구조도이다. 이하에서는 도 15를 통하여 상기 진동감지 센서(280)의 구성 및 기능에 대하여 설명한다.
- [0222] 상기 진동감지 센서(280)는 설치형 단말부(200)에 설치되거나, 또는 상기 사용자(P)가 출입하는 대문 또는 창문에 설치되며 상기 진동감지 센서(280)와 통신 가능하게 연결된 센서로써 부착 설치된 해당 대문 또는 창문의 진동 또는 충격을 감지한다. 상기한 진동감지 센서(280)는 통상의 진동감지 센서 또는 충격감지 센서를 사용하여 달성하되, 그 진동 감도를 조절할 수 있도록 하는 것이 바람직하다.
- [0223] 상기 진동감지 센서(280)는 연동된 대문 또는 창문의 진동 또는 충격을 감지한다. 이는 상기 카메라부(210)에 사람이 드러나지 않더라도 상기 카메라부(210) 시야 바깥에서 침입자가 기구나 기계, 로봇 등을 사용하여 대문이나 창문을 여는 것을 시도할 수 있기 때문이며, 또는 상기 카메라부(210)가 사람을 인식하더라도 그 사람이 창문이나 대문을 두들기거나 흔드는 행동은 충분히 위협적이므로 별도의 범죄 예상 판단 없이 즉각적인 통보를 실시하도록 하기 위함이다.
- [0224] 상기 2가지 경우(기구나 기구를 통한 침입, 문에 충격을 주는 행위) 모두 범죄행위와 밀접성이 매우 높은 위험한 상황이므로 별도의 판단 없이 상기 설치형 단말부(200)가 바로 휴대형 단말부(300)에게 알림 신호를 제공하여 사용자(P)가 해당 행위를 인지할 수 있도록 하며, 또한 상기 서버부(100)에 상기 진동감지 센서(280)의 작동을 알리게 되면, 상기 서버부(100)는 별도의 판단 없이 즉시 기관(G)에게 해당 사실을 알려 조치를 취하게 할 수 있도록 한다.

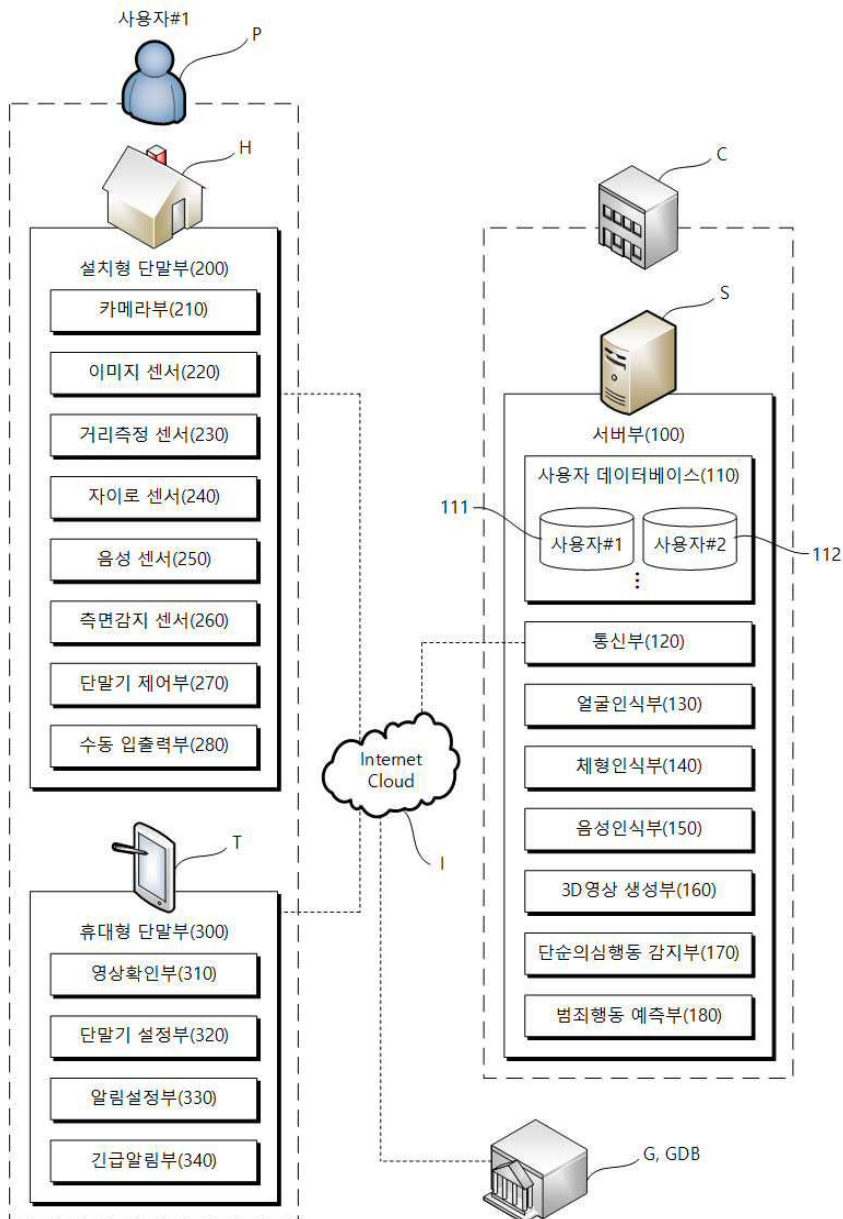
- [0225] 상기와 같이 정보 데이터 및 영상 데이터를 분석하고 예측한 모든 결과물, 예를 들어 상기 버퍼(1112~1115)에 저장된 해당 영상에 대한 각종 정보의 분석 결과, 상기 3D영상 생성부(160)가 생성한 얼굴 모델링 영상 및 체형 모델링 영상, 완성된 3D 모델링 영상, 상기 범죄행동 예측부(180)가 분석 및 예측한 객체의 행동 유형 데이터, 분석 데이터, 예측 데이터 중 설정된 하나 이상은 상기 사용자(P)가 언제든지 자신의 휴대형 단말부(300)를 통해 확인할 수 있도록 해당 영상(V1)의 분석결과 데이터(V1c)에 갱신 저장한다. 이에 따라 상기 사용자(P)는 필요한 경우, 해당 영상(V1) 및 제공되는 분석결과 데이터(V1c)를 확인할 수 있다.
- [0226] 도 16은 상기 휴대형 단말부(300)의 인터페이스 창 화면에 대한 구조도이다. 이하에서는 도 16을 통하여 상기 휴대형 단말부(300)의 구성 및 동작에 대하여 설명한다.
- [0227] 상기 휴대형 단말부(300)는 상술한 바와 같이, 상기 영상 데이터(V1, V2, V3, V4...)를 확인할 수 있는 영상확인부, 상기 설치형 단말부(200) 및 휴대형 단말부(300)를 설정할 수 있는 단말기 설정부(320)를 포함한다.
- [0228] 그리고 상기 알림설정부(330)는 상기 설치형 단말부(200)에서 감지하는 각종의 정보 데이터에 대한 알림 수신 여부를 결정한다. 예를 들어, 상기 알림설정부(330)에서는 단순의심 행동 발생, 영상에 따른 범죄 행동 예측 알림, 진동에 따른 범죄 행동 예측 알림, 택배 배송 예측 알림, 등록된 얼굴의 접근 알림, 영상 업로드 알림에 대한 수신 ON/OFF 기능을 포함한다.
- [0229] 또한 상기 휴대형 단말부(300)는 무선 벨 기능부(350)가 더 포함될 수 있다. 이는 상기 설치형 단말부(200)가 초인종의 형태로서 구현되어 설치되었을 때 추가될 수 있는 것으로, 원거리에서 문 앞에 있는 사람에 대하여 상기 카메라부(210)를 통해 실시간으로 영상을 제공 받고 상기 휴대형 단말부(300)가 설치되어 있는 단말기(T)의 마이크 등을 통하여 응답할 수 있도록 하는 기능을 포함한다.

부호의 설명

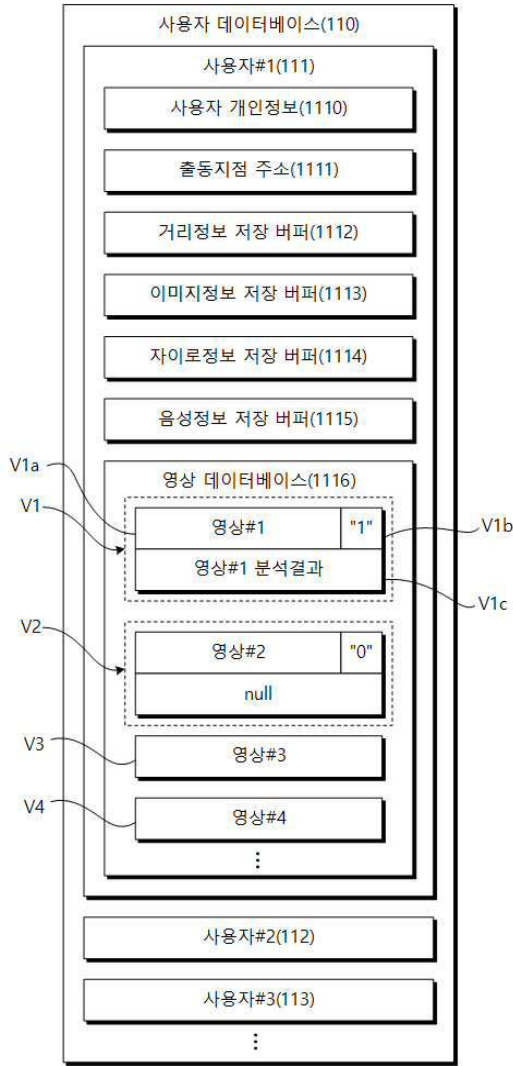
- [0231]
- | | |
|-----------------|-------------------|
| 100 : 서버부. | 110 : 사용자 데이터베이스 |
| 120 : 통신부. | 130 : 얼굴인식부. |
| 140 : 체형인식부. | 150 : 음성인식부. |
| 160 : 3D영상 생성부. | 170 : 단순의심행동 감지부. |
| 180 : 범죄행동 예측부. | 200 : 설치형 단말부. |
| 210 : 카메라부. | 220 : 이미지 센서. |
| 230 : 거리측정 센서. | 240 : 자이로 센서. |
| 250 : 음성 센서. | 260 : 측면감지 센서. |
| 270 : 단말기 제어부. | 280 : 수동 입출력부. |
| 300 : 휴대형 단말부. | 310 : 영상확인부. |
| 320 : 단말기 설정부. | 330 : 알림설정부. |
| 340 : 긴급알림부. | |

도면

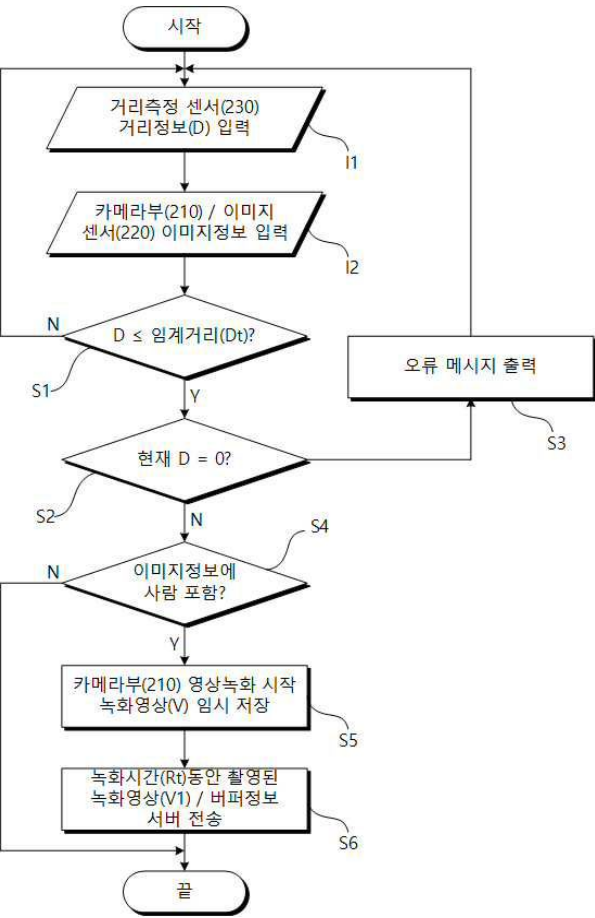
도면1



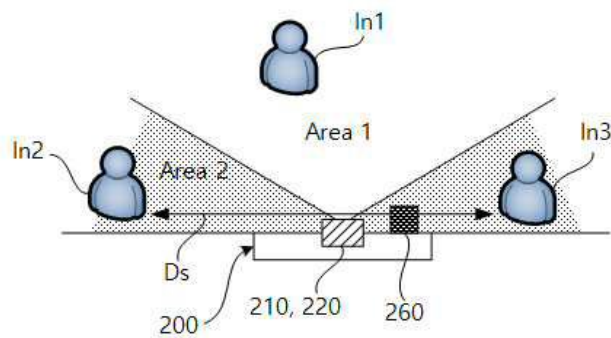
도면2



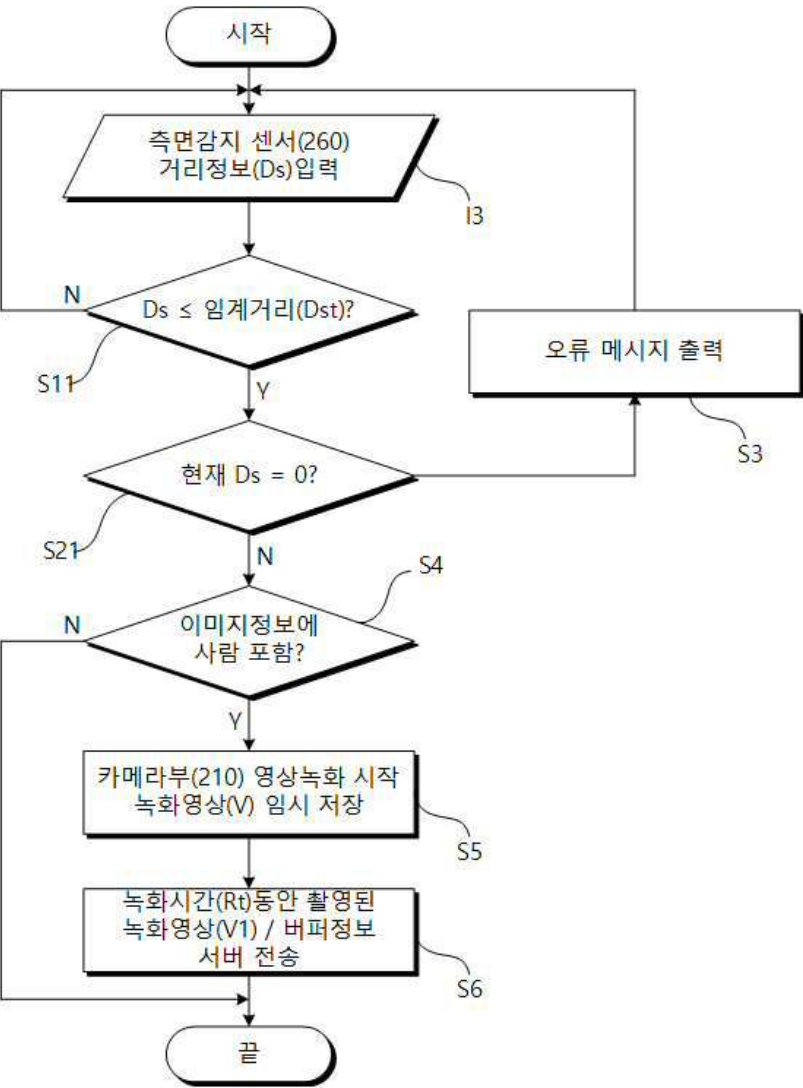
도면3



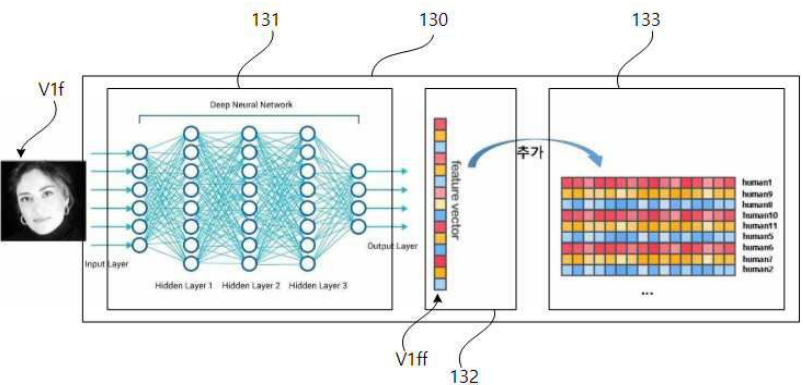
도면3a



도면3b



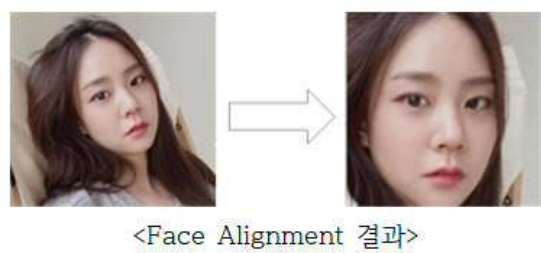
도면4



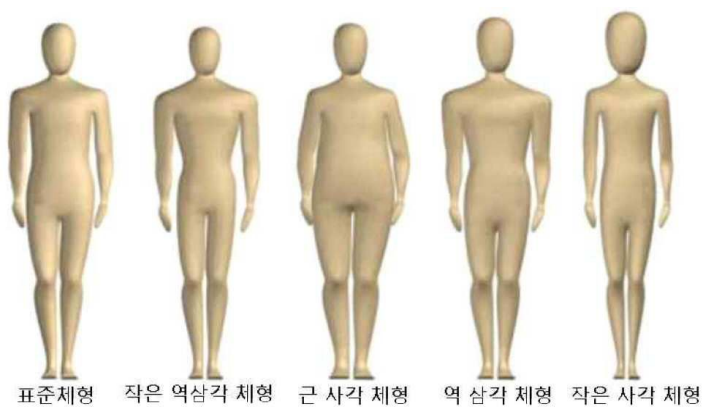
도면4a



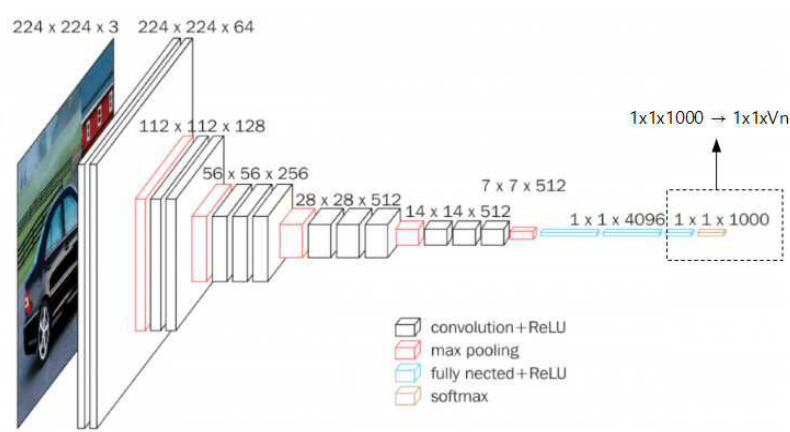
도면4b



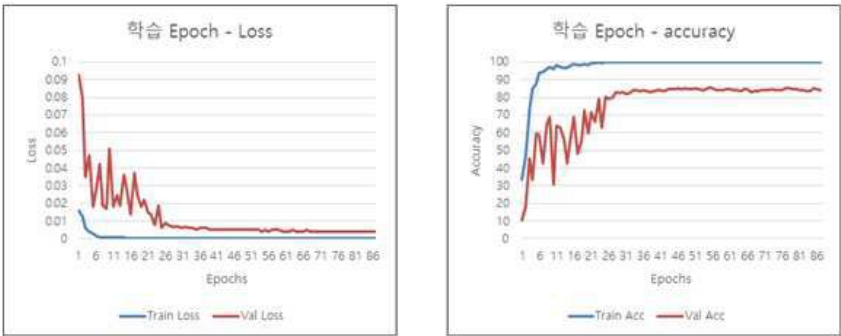
도면5



도면6

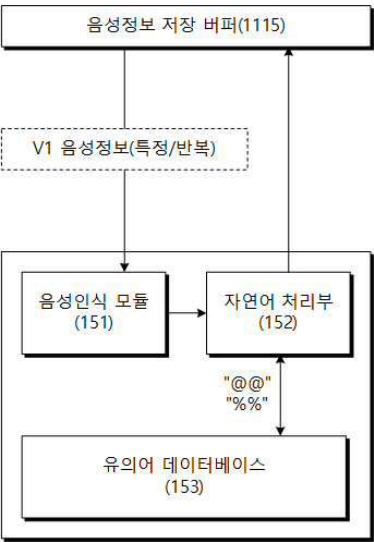


도면6a



<체형 인식 모델 학습 및 검증 결과>

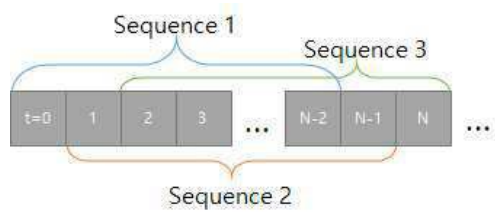
도면7



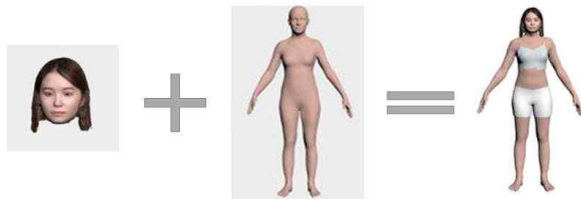
도면8

레이블	예시 및 범죄 연관	데이터 영상
opening door	문 열기, 문이 열려있는지 확인하는 동작	
delivering mail	물건 배달, 택배 기사를 범죄 행위로 오인할 수 있음	
staring	응시, 문의 상태를 확인하는 동작	
bending metal	쇠를 구부리는 행위, 철사 같은 쇠를 구부려서 열쇠 구멍에 넣어 문을 열리는 시도	

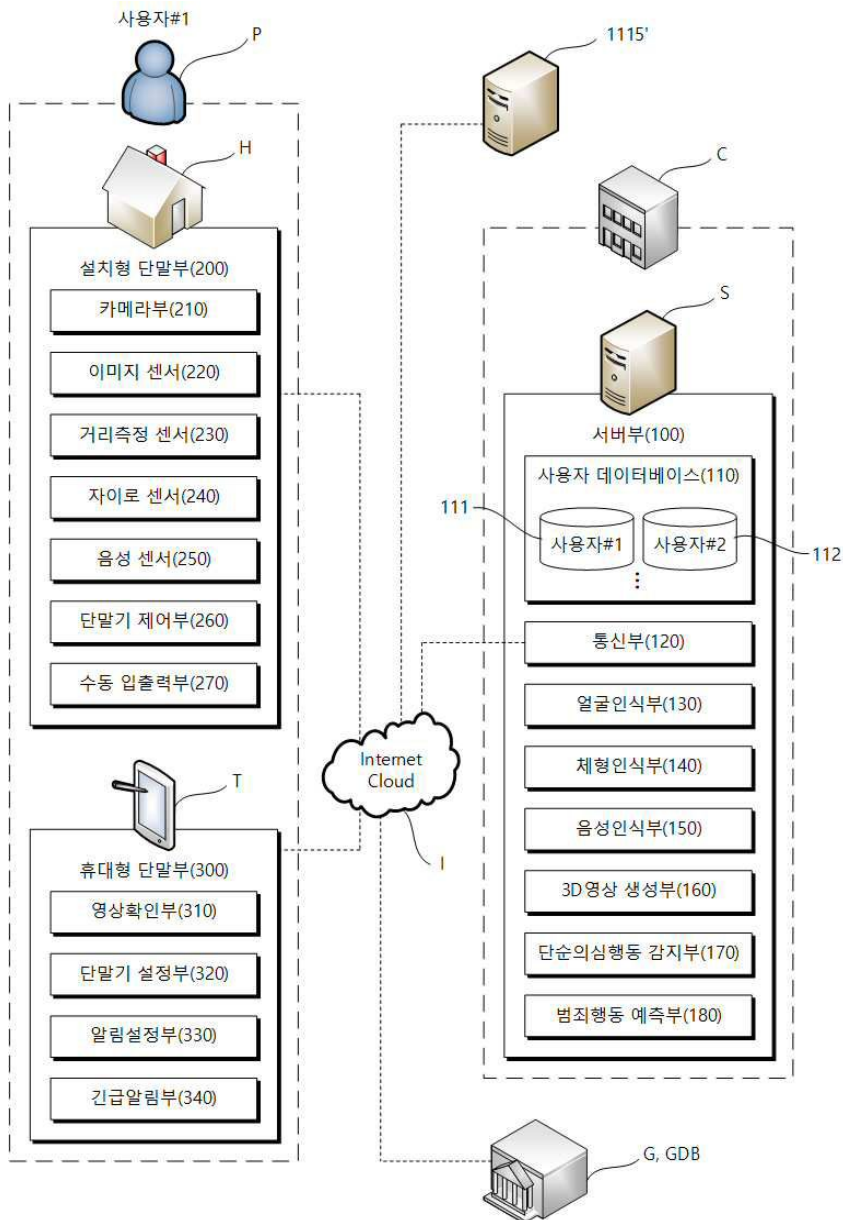
도면9



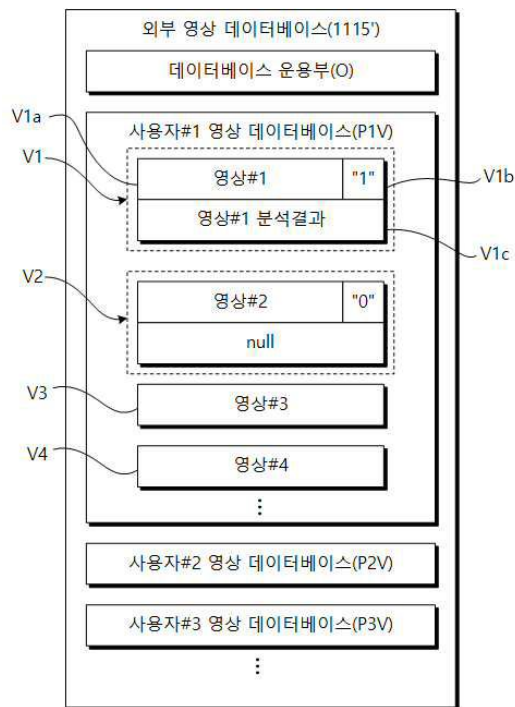
도면10



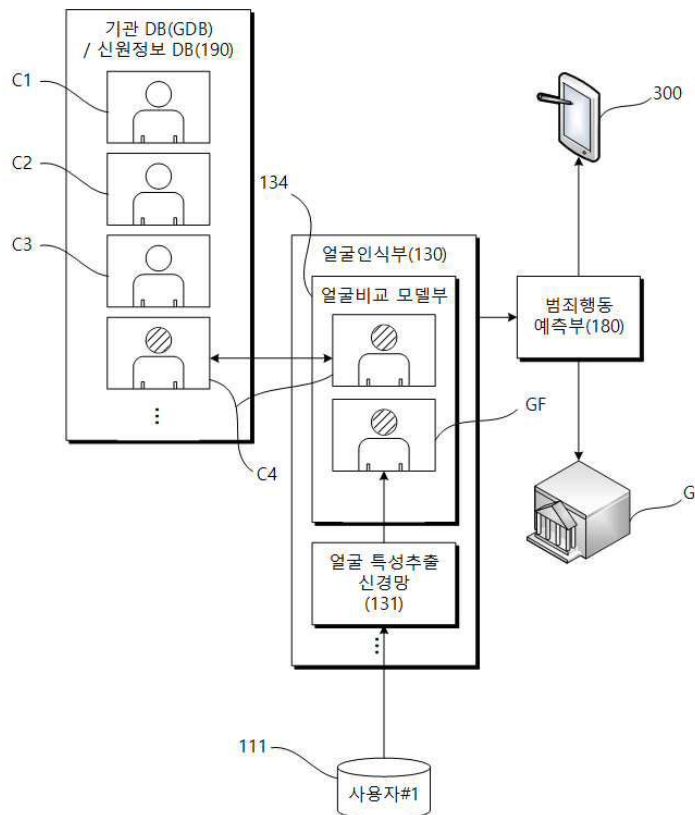
도면11



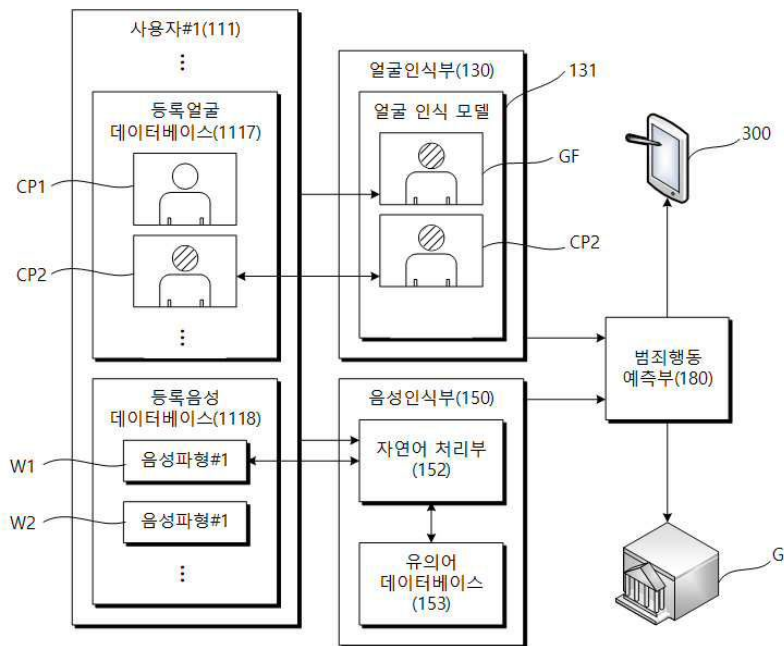
도면12



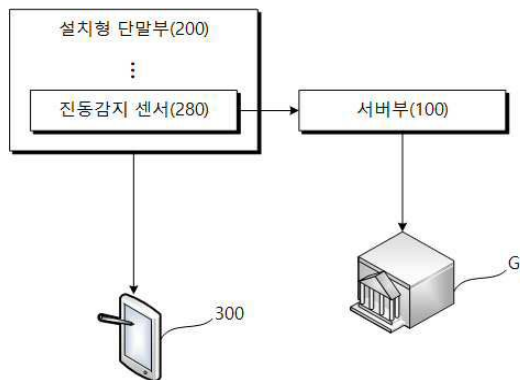
도면13



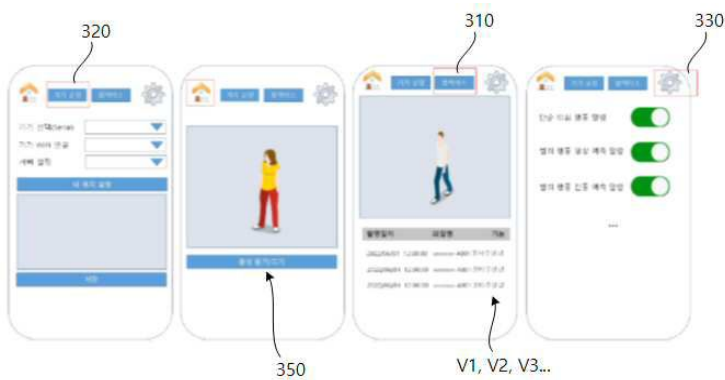
도면14



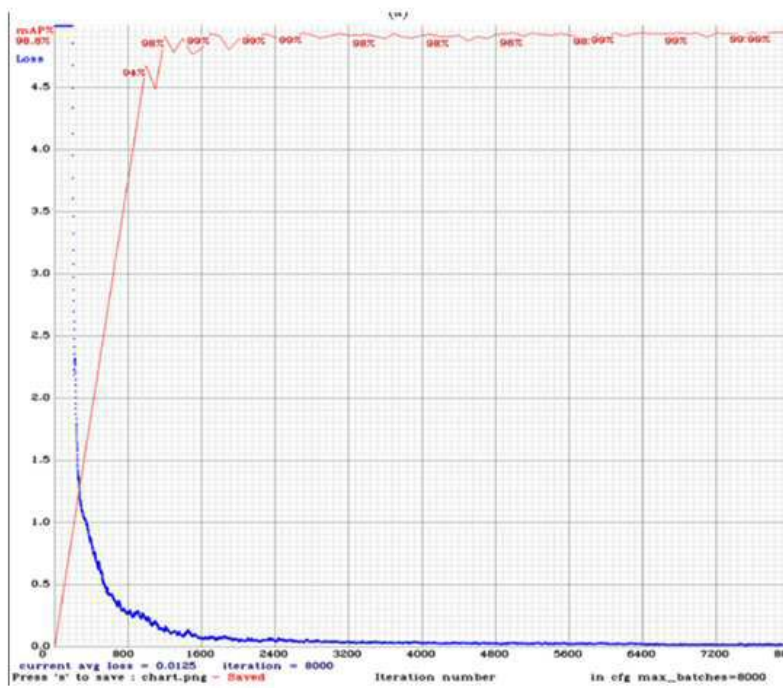
도면15



도면16



도면17



<YOLOv5 사람 검출 학습 결과>