

인공지능, 가짜 뉴스를 만나다 - 딥페이크 영상과 인격권 침해



‘일본 피겨스케이팅 선수 아사다 마오 선수가 사망했다’, ‘최근 드라마로 인기를 얻고 있는 배우가 바이든 미국 대통령을 만났다’처럼 설마 사실일까 싶은 ‘뉴스 아닌 뉴스’들을 본 적 없으신가요? 요즘 유튜브에서는 이렇게 허위 사실들을 영상으로 만들어 제공하는 채널들이 늘고 있다고 합니다. JTBC **보도**에 따르면 이런 ‘가짜 뉴스’ 영상들은 채널 이름을 ‘OO TV’로 붙이거나, 실제 방송 뉴스와 유사한 형태의 자막을 넣기도 하고, 방송 앵커의 사진을 옆에 띄우는 등의 방식으로 진짜 언론사 보도인 것처럼 보이게 제작되었다고 하는데요.

‘가짜 뉴스’가 전 세계 화두로 떠오른 건 2016년 미국 대선 무렵이었습니다. 당시 ‘프란체스코 교황이 트럼프 후보를 지지했다’, ‘힐러리 클린턴 후보가 IS에 무기를 팔았다’와 같은 허위정보가 트위터, 페이스북 등 소셜 미디어를 통해 급속히 확산됐죠. 이후 팩트체크, 미디어 리터러시 교육 등 허위조작정보에 대응하기 위한 여러 방안들이 논의되기도 했습니다. 하지만 가짜 뉴스의 영역은 소셜 미디어에서 유튜브, 틱톡과 같은 온라인동영상플랫폼으로 확산된 모양새입니다. 특히 인공지능 기술로 가상의 인물을 만들거나 기존의 영상, 이미지를 쉽게 변형하고 합성할 수 있게 되면서 눈속임은 더 쉬워졌는데요. 실제로 영국 랭커스터대와 미국 버클리 캘리포니아 대학 연구진이 AI로 만든 가짜 얼굴과 실제 인물 사진을 사람들이 얼마나 잘 구별하는지 **실험한 결과**, 참여자의 절반가량이 실제 사진과 합성 사진을 구분하지 못했다고 합니다.

이렇게 인공지능을 통해 실체가 아닌 이미지나 오디오, 비디오 등을 만들어내는 것을 ‘딥페이크(deep fake)’라고 부르는데요. 알고리즘이 스스로 데이터를 학습해 새로운 결과물을 만들어내는 방식인 ‘딥 러닝(deep learning)’과 ‘가짜(fake)’의 합성어입니다. 허순철 경남대 법학과 교수는 논문 **<유튜브 딥페이크 영상과 허위사실공표>**에서 ‘가짜로 만든 테러나 미사일 공격 동영상은 심각한 전쟁을 촉발할 수 있고, 가짜 포르노그래피의 유포는 명예훼손과 인격권 침해를 야기할 수 있으며, 특히 선거 시기 유포되는 가짜 동영상들은 대의 민주주의의 핵심이라 할 수 있는 공정성을 해칠 우려가 있다’고 지적합니다.

*** 참조 : 허순철(2022). 유튜브 딥페이크 영상과 허위사실공표. <미디어와 인격권> 제8권 제1호.**



허위조작영상으로 인한 인격권 침해 증가

인공지능 기술을 활용해 만든 동영상 때문에 실제 피해를 입는 사례도 늘고 있는데요. 한겨레는 불법 합성 영상물 제작·유포를 처벌하도록 한 「성폭력범죄의 처벌 등에 관한 특례법」 개정안이 시행된 2020년 6월 25일 이후 관련 범죄 판결문을 수집, 분석해 **보도**했습니다. 기사에 따르면 분석 판결의 피고인(46명) 가운데 33명은 범죄를 저지른 적 없는 평범한 사람들이었고, 45.6%(21명)는 학교 선생님이나 동창생, 대학 동기와 같이 평소 알고 지내던 사람의 얼굴을 나체사진 등에 합성해 유포했다고 하는데요. 이렇게 보통 사람들이 가까운 지인을 상대로 범죄를 벌이게 된 데에는 모바일 앱 등으로 누구나 쉽게 영상이나 이미지를 합성, 변형할 수 있게 된 환경이 큰 영향을 미쳤다고 합니다.

박주일 MBC 기자는 **기고**에서 이미지를 조작한 사진 기사의 사례들을 제시하며 ‘기존의 이미지 조작이 편집자의 적극적 의지에 따른 것이었다면 최근의 가짜 뉴스는 커뮤니티 등에서 유통된 이미지가 검증 없이 전파되면서 피해를 발생시키는 특징을 보이고 있다’고 분석했는데요. 무엇보다 ‘짜

깍기로 만들어진 영상이나 이미지들이 워낙 교묘하고 정교할 뿐만 아니라 광범위하게 퍼져있어 자신도 모르는 사이 가짜 뉴스를 퍼뜨리는데 가담할 수 있다’는 점을 짚었습니다.

*** 참조 : 박주일(2020). 딥페이크 기술의 발전과 저널리즘의 새로운 위협 - 이미지공정성 확보를 위한 영상 팩트체크 필요성. 계간 <언론중재> 2020년 가을호.**



딥페이크에 대응하는 노력들

그렇다면 딥페이크 영상으로 피해를 입는 사례를 예방하기 위한 방법에는 어떤 것들이 있을까요? 먼저 조작 영상을 만들 때 사용되는 인공지능 기술을 딥페이크 영상을 감지하는 데 활용하는 연구가 활발히 이루어지고 있는데요. 미국 스탠퍼드 대학의 AI 연구팀은 조작된 영상 속 인물의 입 모양과 음성 소리가 일치하는지 여부를 81% 수준으로 감지하는 프로그램을 만들었고, 국내에서도 딥러닝으로 유해 영상의 진위여부를 감별하는 모델 개발 연구가 진행 중이라고 합니다. (**참조 기사: 국가미래연구원**)

저널리즘 영역과 법, 제도적 차원 규제 방안도 논의되고 있는데요. 프랑스 통신사 AFP와 독일의 언론사 도이치 벨레(Deutsch Well)는 대학, IT 보안업체들과 함께 온라인에 유통되는 영상물 가운데 딥페이크 영상을 탐지하는 프로그램을 개발해 웹사이트와 모바일 앱 등으로 제공했습니다(앞의 **박주일(2020) 글** 참조). 또 국내뿐만 아니라 세계적으로 활동하고 있는 팩트체크 전문 언론인들이 조작 영상을 비롯한 가짜 뉴스 근절을 위해 다방면의 노력을 기울이고 있지요.

법적 대처로는 미국에서 2018년 12월, 범죄 또는 불법행위를 야기하는(facilitate) 딥페이크를 제작하는 경우 벌금이나 2년 이하 징역 등에 처하도록 하는 「악의적 딥페이크 금지법(Malicious Deep Fake Prohibition Act of 2018)」이 발의되었고, 유럽연합은 2018년 허위정보를 식별할 수 있도록 출처 표기를 강화하고 신뢰도 여부를 파악하는 조치를 담은 ‘허위정보행동규범(Communication-Tackling online disinformation : a European Approach)’을 발표한 바 있습니다. 국내에서는 중앙선거관리위원회가 ‘딥페이크 영상 관련 법규운용기준’을 제정해, 선거운동을 할 때 딥페이크 영상 사용이라는 사실을 표시하지 않는 경우 「공직선거법」에 위반된다는 유권 해석을 내리기도 했었죠.

하지만 딥페이크 기술에 어두운 면만 존재하는 것은 아닙니다. 영화나 드라마에 오래 전 사망한 유명인물을 재현하거나 직접 촬영하기 어려운 위험한 장면에 딥페이크 영상을 접목해 창작의 영역을 넓힐 수 있게 되었습니다. 또 최근 글로벌 홍보 수단으로 각광받고 있는 AI 가상 모델은 지역과 언어의 장벽을 허물고 있죠. (**참조 기사: 경향신문**)

기술을 만들어 내는 것도 그 기술을 사용하는 것도 그 주체는 사람인데요. 기술 발전이 가져오는 사회적 편익을 높이고 부작용은 최소화하는 것 역시 인간의 몫 아닐까요?