

## ‘인공지능 에이전트’의 자동화 사례 : 영상 자동화 사례를 중심으로

전준현

**To cite this article :** 전준현 (2025) ‘인공지능 에이전트’의 자동화 사례 : 영상 자동화 사례를 중심으로, 영상문화, 47, 149-178

① earticle에서 제공하는 모든 저작물의 저작권은 원저작자에게 있으며, 학술교육원은 각 저작물의 내용을 보증하거나 책임을 지지 않습니다.

② earticle에서 제공하는 콘텐츠를 무단 복제, 전송, 배포, 기타 저작권법에 위반되는 방법으로 이용할 경우, 관련 법령에 따라 민, 형사상의 책임을 질 수 있습니다.

[www.earticle.net](http://www.earticle.net)

## ‘인공지능 에이전트’의 자동화 사례: 영상 자동화 사례를 중심으로\*

전준현\*\* / 한성대학교

### [국문초록]

최근 생성형 인공지능의 발전은 디지털 영상 제작 전 과정의 자동화를 가속화하고 있으며, 특히 인공지능 에이전트를 활용한 뉴스·광고영상의 생산과 유통이 확산되고 있다. 본 논문은 생성형 AI 에이전트 기반 영상 자동화 사례를 중심으로, 에이전트 오케스트레이션 구조와 인간-AI 협업 거버넌스를 종합적으로 분석한다. 우선 n8n과 Make, Opal과 같은 로우코드 워크플로우 도구를 에이전트 기반 멀티에이전트 시스템의 오케스트레이션 계층으로 위치시키고, Anthropic의 모델 컨텍스트 프로토콜 MCP와 Google의 Agent-to-Agent A2A 프로토콜이 에이전트 내부·외부 통신을 어떻게 표준화하는지 이론적으로 정리하였다. 다음으로 버즈피드, CNET, 네이버의 텍스트·뉴스 자동화 사례와 MBN, 신화통신의 AI 아나운서 사례, 사례를 분석하여 스크립트 생성-팩트체크-영상 합성-편집-배포로 이어지는 공통적인 에이전트 워크플로우와 HITL Human In The Loop 구조를 도출하였다. 이를 바탕으로 n8n/Make 오케스트레이션, MCP/A2A 통신, 팩트체크·윤리 검토·편집 승인 등 인간 검증 모듈을 포함한 에이전트 기반 영상 자동화 통합 아키텍처를 제안하였다.

마지막으로 정확성·책임성, 투명성·편향, 노동 구조, 규제·표준 등 인간-AI 협업 거버넌스의 쟁점을 논의하고, 참조 아키텍처 수준에서 기술적 설계와 사회적 책임을 동시에 고려해야 함을 강조한다.

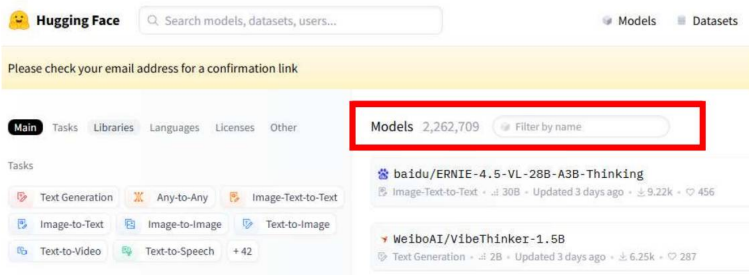
[주제어: 생성형 인공지능, 에이전트 자동화, 인간-AI 협업]

\* 본 연구는 한성대학교 교내학술연구비 지원과제임.

\*\* 한성대학교 창의융합대학 문학문화콘텐츠, naturaljeon@hansung.ac.kr

## 1. 서론

이제는 바야흐로 인공지능의 시대라고 할 수 있다. 2022년 10월 ChatGPT가 대중에 선보인 이후 불과 몇 년 되지 않은 지금, 뉴스에서는 매일 인공지능과 관련한 소식을 접할 수 있다. 인공지능 기술은 1차 산업혁명과도 비견되는데 증기기관 기반의 ‘기계화 혁명’을 통해 수공업 시대에서 증기기관을 활용한 기계가 물건을 생산하는 기계화 시대로 변화한 것과 같이 ‘인공지능 기술’을 통해 모든 지식 기반의 산업에 적용되면서 사람들의 생각 활동 및 데이터 처리를 클릭 한 번으로 처리하는 자동화 시대로 변화시키고 있다. 특히 단순 반복 작업이나 구조화되어 있는 작업에서 빠른 속도로 일을 처리할 수 있으며 대규모 데이터를 다루는 일에서는 오히려 사람보다 실수 없이 정확한 처리를 할 수 있다. 이런 특성을 통해 루틴화 routine化 되어 있는 작업은 인공지능에게 특화되어 있다고 해도 과언이 아니다. 여기에는 규칙적으로 작성해야 하는 블로그에서부터 매일 업무 메일을 확인하고 중요 클라이언트의 메시지를 확인하고 답신하는 일을 포함해 일정을 확인하고 작성하는 소소한 일상까지 모두 해당한다고 볼 수 있다. 이런 작업을 수행하는 생성형 AI 모델은 허깅페이스(Hugging face)에 등록된 수만 해도 2,262,709개이다.



[그림 1] 허깅페이스에 등록된 AI 모델 수(출처: 허깅페이스)

그야말로 AI 모델의 춘추 전국시대라 할 수 있다. 이제는 “어떤 AI가 글을 잘 쓰냐”가 아니라 “어떤 AI가 이 작업을 더 잘하는가?”가 생성형 AI 모델의 선택 기준이 되어버렸다. 다시 말해 단일 작업을 수행하는 생성형 AI 시대를 넘어 AI 에이전트의 시대로 넘어가고 있다. 여기서 AI 에이전트란 여러 AI가

유기적으로 협력하여 복잡한 업무를 수행하는 방식인 에이전틱 워크플로우와 다르게 자율적인 소프트웨어로 주어진 목표를 달성하기 위해 스스로 환경을 인식하고 행동하며 특정 작업을 독립적으로 빠르게 처리하는데 최적화 되어 있는 것을 말한다. 이런 AI 에이전트는 단순한 보고서 작성이나 영상 만들기 와 같은 하나의 작업을 수행하는 것이 아니라 특정 목적, 즉 영상을 만들어서 소셜 미디어에 등록하거나 글을 작성하여 특정 사람들을 선택하고 특정한 시간에 메일 발송하기 등 하나 이상의 작업을 자동화할 수 있다. 이는 비용적인 측면에서도 혁명을 가져왔다. 기존 영상 제작에 있어 가장 많은 시간과 비용이 발생했던 부분이 영상 편집과 CG, 특수효과와 같은 후반작업이었다. 이는 각 영역의 전문가가 투입되어 일주일 이상의 시간을 필요로 하는 것으로 높은 제작 비용과 긴 제작 기간으로 이어져 개인 창작자나 소규모 기업에게는 큰 진입장벽으로 작용했다. 하지만 최근 OpenAI의 Sora나 Google의 Veo를 통해 개인의 아이디어만 있으면 유튜브의 ‘덕빈이’나 ‘정서불안 김햄찌’와 같은 고품질의 동물 Vlog의 고품질 영상 콘텐츠를 제작할 수 있게 되었다. 이제는 단순히 작업 속도를 높이는 것을 넘어, 영상 제작의 민주화가 실현된 것이라 할 수 있다. 이러한 변화는 Netflix의 예고편 생성, BBC의 AI 기반의 자막 시스템, MBN의 AI 앵커 등 폭넓게 사용되고 있다.<sup>1)</sup> 하지만 아직 정확성과 책임 소재에 대한 문제가 남아있으며 CNET의 사례처럼 AI가 만들어낸 환각 Hallucination 현상이나 오류로 인해 언론의 신뢰도가 크게 흔들릴 수도 있다. 게다가 투명성과 편향 문제 또한 함께 도사리고 있다. 이는 AI가 작성했다는 사실을 숨기거나, 학습 데이터 속 편향이 반영되어 정보 생태계 전반이 왜곡되거나 위협해질 수 있다. 뿐만 아니라 AI 도입 이후 노동 구조의 재편 역시 전 세계 시장의 뜨거운 이슈로 자리 잡았으며, 자동화에 따른 대량 해고와 인력 감축은 꼭 풀어야 하는 문제로 떠오르고 있다. 버즈 피드에서 AI 도입 발표와 함께 대량 해고를 단행하면서 자동화에 따른 사회적 문제로 회자된 사건이 대표적 사례라 할 수 있다. 마지막으로 나라마다 다른 규제 환경과 표준의 부재로 인해 AI 생성 콘텐츠에 대한 일관된 거버넌스가 구축되어 있지 않은 것도 큰 문제이다.

1) Wood, C. (2024). “5 AI Case Studies in Customer Service and Support”. VKTR [On-line]. Available: <https://vktr.com/>

지금까지 AI와 관련한 많은 연구에서는 특정 AI 도구의 성능 분석이나 개별 사례 검토가 주를 이루었다.<sup>2)3)</sup> 따라서 본 논문은 여러 AI 에이전트가 함께 협업하여 작동하는 멀티에이전트 구조를 영상 콘텐츠 관점에서 자동화를 중심으로 구조적 분석을 통해 체계적으로 파헤쳐보고, 인간과 AI의 협업 거버넌스 관점까지 연결하고자 한다. 이를 위해 n8n이나 Make와 같은 로우코드 워크플로우 도구나 Google의 Opal과 같은 MCP, A2A 구조를 사용하는 복합 멀티에이전트 워크플로우 도구의 구조 분석과 동시에 사례 분석을 통해 반복적으로 나타나는 문제점을 파악하여 이를 해결할 통합 아키텍처 설계안을 제시하고 인간과 AI 협업의 쟁점과 거버넌스 방향을 탐색하고자 한다.

## 2. 이론적 배경

### (1) 영상 제작 파이프라인과 AI 에이전트의 역할

전통적인 영상 제작 파이프라인은 프리프로덕션(pre-production) → 프로덕션(production) → 포스트프로덕션(post-production)의 세 단계로 진행된다. 각 단계는 다시 세부 작업으로 나뉘지며 이러한 작업들이 순차적이면서도 병렬적으로 진행된다. AI 에이전트는 이러한 전통적 파이프라인의 근본적 구성을 재구성하고 있다.

#### 프리프로덕션 단계의 AI 에이전트 활용

프리프로덕션은 본격적인 촬영에 앞선 단계로 기획을 다듬기 위해 아이디어 발굴부터 시작해 시나리오 작성, 스토리보드 제작, 캐스팅, 로케이션 헌팅, 예산 수립 등을 진행하는 것으로 이 모든 과정을 창작자의 경험과 직관에 의해 진행되었다. 하지만 AI 에이전트가 들어오면서 데이터를 기반으로 작업이

2) 이선재(2025). 「AI 에이전트의 발전과 연구 동향」. 『한국통신학회지(정보와통신)』, 42권 11호, pp.43~52.

3) 정동균, 이종화(2025). 「설명 가능한 GPT 주식투자분석 시스템 연구: RAG 구조와 멀티 AI 에이전트 통합 설계」. 『경영정보학연구』, 27권 3호, pp.243~266.

진행되게 되었으며 ChatGPT나 Claude 같은 대형언어모델을 이용하여 단순한 초안 정도는 쉽게 뽑아낼 수 있으며 복잡한 서사 구조나 캐릭터 설계도 점차 가능해지고 있다. 대표적으로 Jasper AI는 마케팅 영상용 스크립트를 만들 때 브랜드의 톤앤매너를 맞춰 뽑아낼 수 있으며<sup>4)</sup> WriteSonic은 장르와 스타일을 바꿔가며 시나리오도 생성해준다.<sup>5)</sup> 이런 도구를 활용하여 작가들은 전체 시나리오의 결을 빠르게 잡거나 미처 떠올리지 못한 아이디어를 발견하기도 한다.

스토리보드 제작에서는 Midjourney나 DALL-E, Nano Banana와 같은 이미지 생성 AI가 텍스트 설명만으로도 고품질의 이미지를 생성할 수 있어 각 씬에 대한 시각적 구성이 가능하며 스토리보드 전문가 없이 비주얼 기획도 가능해졌다.

캐스팅과 관련해서는 Replica Studios나 Respeecher 등의 서비스를 이용하여 실제 배우의 목소리를 복제하거나 새로운 음성을 생성하는 등의 가상 캐스팅이 가능해졌다. 이를 통해 고인이 된 배우나 가수의 음성을 재현한다거나 다국어 더빙에 활발히 사용되고 있다.

### 프로덕션 단계의 변혁: 실제 촬영에서 AI 생성으로

프로덕션 단계는 전통적으로 실제 촬영이 이루어지는 핵심 단계였다. 하지만 AI 등장 이후 영상 생성 기술이 빠르게 발전하면서 이 핵심 단계 개념마저 변화하고 있다. 물리적 촬영 없이도 텍스트 설명이나 고도화된 명령 프롬프트를 이용하여 고품질 영상을 생성할 수 있게 되면서, 버추얼 프로덕션(Virtual Production)이라는 새로운 패러다임이 등장했다. 대표적인 예로 OpenAI의 Sora는 2024년 ‘걸고있는 여자’, ‘우주인’, ‘중국 축제’등 1분짜리 동영상 선보이며 사람들을 놀라게 한 적이 있다. 기존의 사람이 CG를 이용해 만들려면 각 분야의 수많은 전문 엔지니어가 붙어 오랜 시간 수작업으로 만들어야 하는 영상을 문장 몇 개로 만들 수 있게 되었다. 마치 실제 존재하는 인물이 일본의 밤거리를 걷는 것 같았으며 선글라스에 비친 네온사인의 디테일은 사람들을 놀라게 했다. 그 밖에도 Runway ML과 같은 AI 도구는 스타일화

4) Bain & Company. (2023). “Bain announces alliance with OpenAI; Coca-Cola first client”. *Bain & Company* [On-line]. Available: <https://www.bain.com/>

5) Coca-Cola Company. (2023). “Coca-Cola invites digital artists to ‘Create Real Magic’ using new AI platform”. *The Coca-Cola Company* [On-line]. Available: <https://www.coca-colacompany.com/>

된 영상을 생성할 수 있으며, 모션과 카메라 앵글을 세밀하게 제어할 수도 있다. 이런 버추얼 프로덕션은 AI 도구만으로 완성되는 것이 아니라 Unreal Engine이나 Unity와 같은 실시간 렌더링 엔진을 사용하기도 하고 After Effect와 같은 툴을 활용하여 정교하게 작업하여 내놓기도 한다. 이보다 이전에는 감독이 원하는 장면을 연출하기 위해 비슷한 장소를 찾아 해매거나 세트장을 통째로 만들어 제작하기도 했다. 그 만큼 많은 시간과 비용이 들었으나 실시간 렌더링 엔진을 활용하여 원하는 공간을 3D로 구현하고, LED 월을 만들어 촬영 카메라와 연동한 후 배우와 함께 촬영하는 인카메라 VFX 기술로 장소의 제약을 없애고 시간과 비용을 획기적으로 단축할 수 있게 됐다. 대표적인 사례로는 ‘The Mandalorian’이 있으며, 모션캡처 쪽에서도 AI 기반의 포즈 추정 기술을 이용하여 제작 단가를 낮추고 있다. 이런 이유로 최근 포스트프로덕션에서 AI를 더 반기는 추세다.

### 포스트프로덕션 단계: AI 자동화의 최전방

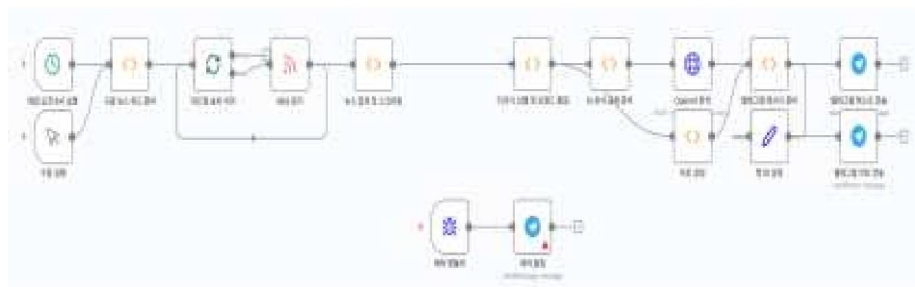
프로덕션 작업이 끝난 후 영상의 완성을 위한 포스트프로덕션 과정에서 특수효과나 CG 작업이 프리프로덕션과 프로덕션 단계보다 더 오랜 시간과 비용을 들이기도 한다. 이런 이유로 포스트프로덕션 단계에서 편집, 색보정, 시각효과, 사운드 디자인이나 자막 작업 등 AI 자동화가 가능한 모든 영역에서 적극적으로 사용하고 있으며 더 많은 영역까지 적용하려 하고 있다. 뿐만 아니라 AI는 단순한 반복 작업을 떠나 창의적 판단까지 보조하고 있다. Adobe Premiere Pro, Scene Edit Detection, DaVinci Resolve 같은 도구들이 편집·컷·색보정처럼 시간과 노력이 많이 드는 수작업 영역을 자동화하면서 제작비를 크게 줄여준 대표적 사례라 할 수 있다.

## (2) 워크플로우 자동화 도구와 영상 제작의 통합

n8n과 Make는 대표적인 워크플로우 자동화 도구로 영상 제작의 복잡한 프로세스를 체계화하고 효율화하는 핵심 인프라로 사용되고 있다. 최근에는 Google의 Opal 자동화 도구도 출시되어 코드 작성 없이 복잡한 자동화 워크플로우를 구축할 수 있게 해준다.<sup>6)</sup>

### n8n과 make의 영상 제작 워크플로우 구현

n8n과 Make는 유사한 구조를 사용하는데 오픈소스 워크플로우 자동화 플랫폼으로, 200개 이상의 노드(통합 모듈)를 제공하고 있다. 영상 제작에 있어 n8n과 Make의 주요 장점은 유연성과 확장성이다. 사용자는 시각적 인터페이스를 통해 복잡한 워크플로우를 설계할 수 있으며, 필요시 커스텀 노드를 추가할 수 있으며 영상 제작에 활용되는 주요 노드는 ‘HTTP Request 노드’, ‘FFmpeg 노드’, ‘YouTube/Vimeo 노드’, ‘AWS S3/Google Cloud Storage 노드’, ‘Webhook’노드이다. HTTP Request 노드는 다양한 AI API(OpenAI, Replicate, Hugging Face 등)를 호출하는 역할을 한다. FFmpeg 노드는 영상 변환과 자르기, 합성 등의 기본 편집 작업을 수행하며 YouTube/Vimeo 노드는 영상 업로드, 메타데이터 관리, 분석 데이터를 수집한다. AWS S3/Google Cloud Storage 노드는 대용량 영상 파일을 저장 및 관리하며 Webhook 노드는 실시간 이벤트 기반의 워크플로우 트리거 역할을 담당한다. 실제 구현 예시로 일일 뉴스 요약 영상을 자동 생성하는 주요 워크플로우를 살펴보면 [그림 2]와 같이 만들 수 있다.



[그림 2] n8n 뉴스 생성 워크플로우(출처: 저자 작성)

이런 워크플로우는 사용자의 명령 없이 매일 정해진 시간에 스스로 작동하면서 뉴스 영상을 만들고 배포까지 자동으로 수행한다.

- 6) Dmitrievna, Y., & Parsadanyan, E. (2024). “AI agentic workflows: a practical guide for n8n automation”. *n8n Blog* [On-line]. Available: <https://n8n.io/blog/>

### 워크플로우 도구를 통한 영상 제작 파이프라인의 통합

전통적인 파이프라인은 프리프로덕션-프로덕션-포스트프로덕션 모두 각자 따로 돌아갔으며 각 단계 사이를 사람이 이어오면서 전달하는 과정에서 소통 착오나 파일 포맷의 불일치, 버전 혼선 등 휴먼에러 Human Error가 심심찮게 일어나며 문제를 일으켰으나 n8n이나 Make와 같은 워크플로우 자동화 도구가 도입되면서 영상 제작 파이프라인 구조 자체가 바뀌고 있다. MCP나 A2A 구조를 사용해 각 단계 간 자동 연계가 가능하며 스크립트 에이전트가 작업을 끝내면 자동으로 웹훅이 발동돼 영상 합성 에이전트로 바로 움직이는 식으로 사람없이 표준 포맷으로 데이터가 넘어가 일정의 지연이나 비용의 로스를 줄일 수 있다. 또한 MCP 덕분에 각 에이전트의 프로토콜 규약에 따라 프로젝트의 일관성을 유지할 수 있으며 A2A를 통해 협업 중인 에이전트들의 작업 상태를 실시간으로 연동하여 어느 한 에이전트가 지연되거나 오류를 일으켜도 전체 워크플로우에는 영향을 미치지 않게 작동할 수 있다.

이러한 통합 구조의 효과는 세 가지로 요약할 수 있다. **리드타임 단축**은 단계 간 대기 시간이 제거되어 전체 제작 기간이 획기적으로 줄어든다. **모듈 교체 용이성**은 표준화된 인터페이스(MCP, A2A) 덕분에 특정 AI 모델이나 도구가 더 나은 것으로 업그레이드될 경우 해당 모듈만 교체하면 된다. **병렬처리 가능성**은 독립적인 작업(예: 음성 합성과 배경 영상 생성)을 동시에 진행하여 효율을 극대화할 수 있다. 결과적으로 n8n/Make 기반의 오케스트레이션은 기존의 선형적·단절적 파이프라인을 유기적·통합적 파이프라인으로 전환시키는 핵심 기술이라 할 수 있다.

### Opal을 활용한 템플릿 기반 영상 제작

Opal은 워크플로우를 시각적으로 제공하여 프롬프트, AI 모델 호출 및 기타 도구 연결이 가능하며 다단계 앱을 제작할 수 있도록 지원하고 있다. Opal은 n8n과 Make와 유사한 모습이지만 JSON 구조를 이해하고 각 노드에 대한 세팅을 해줘야 하는 n8n과 Make와 다르게 대화형 자연어 명령어로 설정할 수 있으며 [그림 3]과 같이 비주얼 편집기를 통해 새 기능을 추가하거나 도구를 호출할 수 있다.



[그림 3] Opal의 동영상 제작 워크플로우(출처: 저자 작성)

### (3) 에이전트 간 통신 프로토콜(MCP와 A2A)

여러 AI 에이전트가 동시에 함께 작동하려면 공통 언어, 즉 표준 통신 프로토콜이 있어야 한다. 이를 제안한 것이 처음 Anthropic이 내놓은 MCP(Model Context Protocol)이며 더 나아가 에이전트끼리 연결한 것이 Google의 A2A(Agent-to-Agent) 프로토콜이다.<sup>7)8)</sup>

#### MCP(Model Context Protocol)의 구조와 활용

Anthropic의 MCP는 원래 LLM과 외부 데이터 소스 사이의 상호작용을 규격화하기 위해 만들어졌지만, 영상 제작 워크플로우에서 에이전트 간 협업에도 잘 맞는다. 이 MCP의 핵심 구성요소는 컨텍스트 관리와 도구의 선언 및 호출, 그리고 스트리밍 지원이다.

컨텍스트 관리는 이전 작업의 결과, 사용 가능한 도구와 제약 조건과 같은 정보를 구조화된 형태로 가지고 있으며 프로젝트의 설정이나 스타일 가이드,

7) Aldridge, N., Brooker, M., & Sivasubramanian, S. (2025). “Open Protocols for Agent Interoperability Part 1: Inter-Agent Communication on MCP”. *AWS Open Source Blog* [On-line]. Available: <https://aws.amazon.com/blogs/opensource/>

8) Arsanjani, A. (2025). “Complementary Protocols for Agentic Systems: Understanding Google's A2A & Anthropic's MCP”. *Medium* [On-line]. Available: <https://medium.com/>

타임라인 정보 등이 이에 해당한다.

도구의 선언 및 호출은 에이전트가 사용 가능한 API나 함수, 서비스 등을 미리 정의하고 정해진 방식대로 부르는 것을 말한다.<sup>9)</sup> 가령 영상 편집 에이전트가 색보정 도구를 호출할 때 MCP 규격을 따른다면 호환성 문제에 대한 걱정은 하지 않아도 된다.

### A2A(Agent-to-Agent) 프로토콜의 협업 메커니즘

Google의 A2A 프로토콜은 MCP와는 조금 다르게 에이전트 간에 안전하게 협력할 수 있도록 설계됐다. 이는 에이전트가 개별적인 프로그램이나 함수를 호출하는 대신 그것을 총괄적으로 관리하는 에이전트와 통신함으로써 개별적으로 호출해야 하는 MCP보다 안정적으로 통신할 수 있게 만든 것이다. 가령 영상 편집 에이전트가 특수효과 전문 에이전트를 찾아 작업을 맡기는 식이다. 작업 협상이나 계약에 있어 품질, 시간, 비용과 같은 조건을 에이전트끼리 조율하고 합의하는 식이다. 이 기능은 상업 영상을 만들 때 에이전트마다 소속 조직이 다를 때 특히 빛을 발할 수 있다. 에이전트와 에이전트간 협업을 진행하므로 에이전트에 속한 API나 함수, 서비스 개별로 호출할 필요가 없다는 의미로 한 에이전트가 늦어지거나 문제가 생겨도 전체 워크플로우에 미치는 영향을 최소화 할 수 있다는 의미이다. 이에 그와 관련된 사례 분석을 통해 좀 더 자세히 알아보도록 하겠다.

## 3. 생성형 AI 에이전트 활용 사례 분석

### (1) 콘텐츠 자동 생성 사례(버즈피드, CNET, 네이버)

디지털 미디어 업계에서는 기사나 콘텐츠를 생성형 AI로 자동 생산하려는 시도가 증가하고 있다. 여기서는 미국의 온라인 매체 버즈피드와 CNET, 그리고 국내 포털 네이버의 사례를 통해 뉴스/콘텐츠 자동화의 현실을 살펴본다.

9) 백창원(2025). 「AI 에이전트와 MCP를 활용한 데이터 보안 취약 대응 전략」. 『한국지능시스템학회 논문지』, 35권 5호, pp.465~474.



[그림 4] 버즈피드가 생성한 뉴스(출처:서터스톡)

버즈피드의 경우, 2023년 CEO 조나 페레티(Jonah Peretti)가 직원들에게 보낸 메모를 통해 생성형 AI 활용 전략을 공개했다. 이 메모에서 그는 “2023년에 OpenAI의 기술(ChatGPT)를 활용하여 퀴즈 경험을 향상시키고 콘텐츠를 개인화하는 R&D 단계의 콘텐츠를 핵심 비즈니스 일부로 옮겨갈 것”이라며 AI 도구 활용에 뜻이 있음을 밝혔다.<sup>10)</sup> 실제로 버즈피드는 사용자가 퀴즈 답변을 입력하면 그림 4와 같이 그 사람을 주인공으로 한 짤막한 맞춤형 스토리를 AI가 생성해 주는 등 인터랙티브 콘텐츠를 선보일 계획이었다. 이 발표가 알려지자 시장에서는 버즈피드에 큰 관심을 보였고, 주가가 하루 만에 150% 폭등하는 현상도 있었다. 또한 트렌드 분석부터 초안 작성까지 AI가 돕는다면 기자나 에디터들은 더 창의적인 작업에 집중할 수 있다는 긍정론도 당시 존재했다.

그러나 그런 편리함 뒤로 단점과 위험도 존재했다. 버즈피드의 시도는 곧바로 윤리적 논쟁을 불러일으켰으며 저널리즘의 신뢰성이 도마 위에 올랐다.

10) Paul, K., & agencies. (2023). “BuzzFeed to use AI to 'enhance' its content and quizzes - report”. *The Guardian* [On-line]. Available: <https://www.theguardian.com/>

AI가 작성한 콘텐츠가 과연 버즈피드의 품질 기준을 만족할지, 혹시 잘못된 정보가 포함되지는 않을지에 대한 우려가 제기된 것이다. 또한 “AI 도입이 인력 감축을 정당화하는 수단 아니냐”는 비판도 있었다. 실제로 버즈피드는 같은 시기 비용 절감을 위해 전체 직원의 12%를 해고한다고 발표했는데, AI 도입 소식과 맞물려 “인간 기사를 대체하려는 움직임”이라는 반발을 샀다. 요약하면, 버즈피드 사례는 콘텐츠 개인화·대량생산이라는 이점과 저널리즘 윤리·일자리 영향이라는 위험 요소가 혼재된 사례로 볼 수 있다.

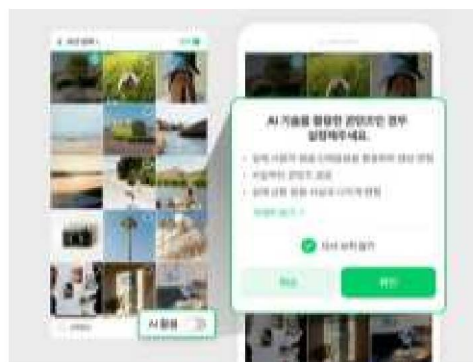
한편, 테크 전문 매체인 CNET의 사례는 AI 자동생성 콘텐츠의 한계를 극명하게 보여주는 예이다. CNET은 2022년 말부터 사람들 모르게 “CNET Money” 섹션에 AI가 작성한 금융 정보 기사 수십 편을 조용히 게시해 왔다. 2023년 1월 Futurism 등 외부 매체의 폭로로 이 사실이 알려졌고, 거센 논란이 일었다. CNET 편집장은 정정 공지문을 통해 AI 기사 77건 중 41건에 사실 오류가 있어 수정했다고 시인했다.<sup>11)</sup> 일부 기사에서는 표현이 다른 매체를 베낀 듯한 표절 시비까지 일어나, “고쳐 쓴 문장이 있다”는 문구를 달기도 했다. 게다가 CNET 내부 직원들도 어떤 기사가 AI가 작성한 것인지 모를 정도로 투명성이 없었던 것으로 밝혀졌다. 결국 CNET 모회사인 Red Ventures는 비난 여론 속에 “당분간 AI생성 콘텐츠를 중단한다”고 발표했다.

CNET 사례의 이점이라면, SEO(검색엔진 최적화)에 최적화된 금융 기사를 대량 생산하여 트래픽을 늘리고 광고 수익을 올리려는 비즈니스 모델이었다는 점이다. 실제로 모회사는 Bankrate, CreditCards.com 등 신용카드 추천 사이트도 운영하고 있었는데, AI 기사를 통해 이용자를 끌어들이려 신용카드 가입을 유도하려 한 것이다. 그러나 이러한 상업적 동기로 인한 품질 희생은 큰 역풍을 맞았다. 주요 단점은 사실 검증 부실로 인한 잘못된 정보 제공과 신뢰도 추락이었다. AI 모델 특성상 그럴듯하게 글은 쓰지만 환각이라고 하는 틀린 내용을 사실과 섞어 쓸 수 있는데, CNET의 초기 검증 프로세스는 소홀했고, 결국 브랜드 이미지에 타격을 입을 수 밖에 없었다. 이 사건 이후 위키백과 커뮤니티는 CNET을 신뢰할 수 없는 출처로 격하시키기도 했다. 결국 CNET 사례는 AI의 한계인 “환각(hallucination)” 문제와, 이를 관리·통제하지

11) Sato, M., & Roth, E. (2023). “CNET found errors in more than half of its AI-written stories”. *The Verge* [On-line]. Available: <https://www.theverge.com/>

않고 성급히 도입할 때 발생하는 위험을 보여준 사례로 남게 되었다.

마지막으로 네이버의 경우는 앞선 둘과 결은 다르지만 맥락상 흥미로운 예이다. 네이버 사례는 앞의 두 회사와 결이 다르지만 짚어볼 만하다. 네이버는 하이퍼클로바(HyperCLOVA) 같은 자체 대규모 언어모델을 갖고 있어서, 검색과 콘텐츠 영역에 생성형 AI를 적극 심을 계획을 세우고 있다. 이미 AI 브리핑 기능을 시범 운영 중인데, 금융·건강 등 분야별 뉴스를 AI가 요약해 보여주는 식이다. 웹소설이나 만화 같은 UGC 창작을 AI로 돕는 방안도 연구 중이다. 다만, 네이버는 플랫폼으로서의 책임을 강조하며 신중하게 움직이고 있다. 2024년 총선을 앞두고는 “AI가 자동으로 만든 기사”가 일반 뉴스 영역에 뜨지 않도록 막았고, 사람이 상당 부분 손댄 경우에만 노출을 허용했다. 선거 기간에 AI 허위 정보가 퍼지는 것을 막으려는 조치였다. 한 발 더 나아가 네이버 블로그, 카페, 밴드에서 콘텐츠를 올릴 때 AI 활용 여부를 직접 표시하도록 시스템을 바꿨다. 사용자가 얼굴 합성이나 가짜 음성 등 AI 기술을 썼다면 “AI 활용” 토글을 켜게 해둔 것이다. 다른 이용자는 그 표시만 보고 해당 게시물이 AI로 만들어졌거나 변형됐는지 바로 알 수 있다.



[그림 5] 네이버 밴드(Band) 앱에서 게시물 업로드 시 AI 활용 팝업  
(출처: 네이버)

이를 통해 다른 이용자들은 해당 게시물이 얼굴 합성이나 가짜 음성 등 AI 기술로 생성/변형되었는지 여부를 한눈에 알 수 있다. 네이버는 이처럼 AI 활용 콘텐츠에 대한 투명성 조치를 2025년부터 블로그 등 UGC 전반으로 확대

적용하고 있다.<sup>12)</sup> 네이버 사례가 의미 있는 이유는 AI를 콘텐츠 생산에 끌어들이면서도 투명성과 신뢰 장치를 먼저 깔았다는 데 있다. 최신 AI 기술로 서비스를 혁신하려는 시도가 장점이라면, 단점과 위험은 이를 제대로 관리하지 않았을 때 플랫폼 신뢰가 무너질 수 있다는 점이다. 네이버는 그 위험을 알고 선제적으로 움직였다. 즉, 장점은 최신 AI기술을 서비스 혁신에 활용하려는 시도이며, 단점/위험은 이를 관리하지 않을 경우 플랫폼 신뢰성이 훼손될 수 있음을 인식하고 선제적으로 대응하고 있다는 점에서 찾아볼 수 있다. 위의 세 사례를 표로 정리하면 다음과 같다:

<표 1> 미디어 콘텐츠 자동화 & 각 회사 사례 비교표

| 사례 및 분야                      | 주요 활용 내용                                                                                                  | 이점(Advantages)                                                                                                                      | 한계 및 위험<br>(Drawbacks & Risks)                                                                                                                         |
|------------------------------|-----------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------------------------------------------------------|
| <b>버즈피드</b><br>(디지털 미디어, 미국) | <ul style="list-style-type: none"><li>• 퀴즈 및 기사 콘텐츠에 생성형 AI 도입</li><li>• ChatGPT 활용 개인맞춤 스토리 생성</li></ul> | <ul style="list-style-type: none"><li>• 콘텐츠 대량 생산으로 생산성 증대</li><li>• 독자 <b>개인화 콘텐츠</b>로 참여도 향상</li><li>• 인간 기자는 고부가 작업 집중</li></ul> | <ul style="list-style-type: none"><li>• 정확성·품질 저하 우려 (사실 검증 필요)</li><li>• <b>저널리즘 윤리</b> 논란 (AI로 기사 작성)</li><li>• 기자 직업 대체 불안 및 내부 반발</li></ul>        |
| <b>CNET</b><br>(IT 전문매체, 미국) | <ul style="list-style-type: none"><li>• 금융정보 기사 자동작성</li><li>• AI 비공개 활용: 77편 중 41편 오류 발견</li></ul>       | <ul style="list-style-type: none"><li>• 반복적인 설명기사 자동화로 <b>속도·비용 절감</b></li><li>• SEO 최적화 콘텐츠로 트래픽 증대 기대</li></ul>                   | <ul style="list-style-type: none"><li>• 사실 오류 다수 및 문장 표절 시비</li><li>• <b>신뢰도 실추</b> 및 정정 작업 증가</li><li>• 투명성 부족 (독자에 AI 활용 미공개)</li></ul>              |
| <b>네이버</b><br>(포털 플랫폼, 한국)   | <ul style="list-style-type: none"><li>• 뉴스 AI요약 등 콘텐츠 실험</li><li>• UGC에 ‘AI 활용’ 여부 표기 제도화</li></ul>       | <ul style="list-style-type: none"><li>• 검색서비스 혁신 (AI 브리핑으로 정보 접근성 향상)</li><li>• <b>콘텐츠 투명성</b> 제고 (AI 제작물 명시로 신뢰도 관리)</li></ul>     | <ul style="list-style-type: none"><li>• 선거 등 민감시기 AI기사 제한 필요 (허위정보 확산 우려)</li><li>• AI 생성물에 대한 <b>품질·사실성 검증</b> 부담</li><li>• 일관된 가이드라인 마련 과제</li></ul> |

(출처: 연구자 구성)

12) 정유림(2025. 5. 20). 「“AI로 만든 콘텐츠입니다”...네이버 ‘AI 활용 콘텐츠’ 표시 도입」. 『아이뉴스24』 [On-line]. Available: <https://www.inews24.com/>

## (2) 영상 콘텐츠 자동화 및 AI 아나운서 사례(MBN, 신화통신)

영상 콘텐츠 분야에서도 생성형 AI는 영상 편집 자동화나 버추얼 휴먼(가상 인간) 기술과 결합하며 새로운 실험들을 진행하고 있다. 특히 뉴스 진행자의 영상과 음성을 본떠 만든 디지털 복제 AI 아나운서는 24시간 뉴스 전달을 가능케 하는 기술로 주목받았다. 한국의 MBN과 중국 신화통신(Xinhua)이 선도적인 사례로 꼽힌다.



[그림 6] 인간 앵커와 AI 앵커 비교(출처: MBN 종합뉴스)

[그림 6]의 MBN 뉴스 앵커 김주하(왼쪽)와 그의 AI 복제판(오른쪽)이 방송 대화를 나누는 장면으로 MBN은 AI 기업 머니브레인<sup>13)</sup>과 협업하여 실제 앵커와 외모·목소리가 거의 같은 AI 아나운서 “AI 김주하”를 개발, 2020년 11월 최초 도입했다.<sup>13)</sup> 머니 브레인은 실제 MBN 간판 앵커인 김주하 아나운서의 10시간 분량 방송 영상을 딥러닝으로 학습시켜, 얼굴 표정, 입술 동작, 말투까지 그대로 모사한 가상 김주하를 만든 것이다. 이 AI 앵커는 MBN 종합뉴스에서 실제 김주하 앵커와 함께 등장하여 대화를 나누며 음성 톤을 비교하는 퍼포먼스도 선보였다. MBN은 이 AI 앵커 도입을 통해 “자연재해 등 비상

13) “MBN introduces Korea's first AI news anchor”. (2020). *JoongAng Daily* [On-line]. Available: <https://koreajoongangdaily.joins.com/>

상황에서 신속한 보도가 가능하고, 새벽 시간 등에도 24시간 뉴스 진행이 가능”하다는 장점을 피력했다. 또한 장기적으로는 인건비 절감 등 제작 비용을 낮추는 효과도 거둘 수 있을 것으로 예상했다. 하지만 단점과 논란도 있었는데 시청자들 중 일부는 “진짜 앵커와 너무 비슷하다 보니 오싹하다”는 반응을 보였다. 이는 앞서 언급한 언캐니 밸리 효과의 일례라 할 수 있다.<sup>14)</sup> 사람들은 현실과 약간 다른 정도의 가상 캐릭터에는 호감을 느낄 수 있어도, 거의 똑같아지면 미묘한 불쾌감을 느낄 수 있다는 것이다. 실제 방송 후 온라인 댓글에는 “알고 보니 가짜라서 소름 끼친다”는 반응과 “신기하다”는 반응이 엇갈렸다. 또한 AI 앵커가 인간 앵커를 대체할 수 있는가에 대한 의문도 제기됐다. 보도 내용 자체가 사람이 작성한 원고를 AI가 읽는 것이기에, 창의적 질문이나 돌발상황에 대한 대응 등이 어려울 수밖에 없다. 이 때문에 전문가들은 “AI로 완전 대체하기는 힘들고, 보조적 역할에만 머물 것”이라는 전망을 내놓았다. 요컨대 MBN AI 앵커는 기술 시연 측면에서는 성공적이었으나, 대중 수용성과 콘텐츠 다양성 측면의 한계가 여실히 드러난 사례라고 할 수 있다.

중국 신화통신(Xinhua)은 MBN보다 앞선 2018~2019년에 이미 AI 뉴스 앵커를 선보였다.<sup>15)</sup> 2019년 3월에는 세계 최초의 여성 AI 앵커인 신 샤오명(Xin Xiaomeng)을 공개했는데, 이는 실제 신화통신 앵커인 추 멩(Qu Meng)을 모델로 한 것이었다. 신 샤오명은 중국 최대 정치행사인 양회(兩會) 뉴스를 전하며 데뷔했고, 같은 해 영어로 뉴스를 전하는 남성 AI 앵커도 추가 공개되었다. 신화통신의 AI 앵커들은 중국 IT기업 소고우(Sogou)와의 협업으로 개발되었으며, AI 합성 엔진이 텍스트 원고를 입력받으면 사람과 흡사한 얼굴 영상과 음성을 합성하여 출력하는 형태이다. 이 방식의 장점은 다국어 뉴스를 동시에 제작하거나, 일기예보 등과 같이 단순하고 반복적인 뉴스는 사람의 손을 거치지 않고 생산할 수 있다는 점이다. 실제로 신화통신은 영어, 러시아어 등 다국어 AI 앵커를 활용해 국제뉴스 콘텐츠를 늘렸다. 또한 국가 주도로 AI 기술력을 과시하는 측면도 있었다. 반면 단점으로는 MBN 사례와 유사하

14) 전혜리, 윤명세, 김성희(2025). 「전문가형 AI 에이전트의 프로필 이미지 의인화 정도에 따른 사용자 경험 연구」. 한국HCI학회 학술대회. 강원.

15) “China's Xinhua presents news using robot news anchor”. (2019). Reuters [On-line]. Available: <https://www.reuters.com/>

게, 표정이나 어조에서 어색함이 완전히 해소되지 않아 자연스러움의 한계가 있었다는 점이다. 처음 공개 당시엔 기술 시연이라는 호평이 있었지만, 시간이 지나면서 화제성 감소와 함께 실제 현업 기자들을 대체하지는 못했다는 평가가 나왔다. 이밖에 가짜뉴스 합성에 대한 우려나 정보의 편향 가능성이 제기되기도 한다.

전반적으로 AI 아나운서 사례들을 통해 얻은 교훈은 다음과 같다. 장점은 대량의 영상 콘텐츠를 신속히 만들어낼 수 있다는 것이며, 노동집약적 방송 제작의 효율화 가능성을 보여줬다는 것이다. 또한 긴급 상황 발생 시 새벽에도 바로 속보 영상을 내보낼 수 있고, 동시에 여러 채널의 다국어 뉴스를 송출할 수 있다는 확장성을 얻을 수 있다. 반면 단점과 위험으로는, 인간다운 영상일수록 발생하는 미묘한 거부감 문제와 창의적 소통 부재로 발생하는 콘텐츠 획일화 가능성 등이 지적되며 시청자가 AI 앵커를 진짜 사람으로 오인하거나, 가짜 뉴스 영상에 악용되는 등의 사회적 영향 관리도 중요한 쟁점 중 하나이다. 이러한 이점과 한계를 요약하면 <표 2>와 같다.

<표 2> 영상 콘텐츠 자동화 & AI 아나운서 사례 비교표

| 사례              | 주요 활용 내용                                                                                                                   | 이점 (Advantages)                                                                                        | 한계 및 위험 (Drawbacks & Risks)                                                                                                                                     |
|-----------------|----------------------------------------------------------------------------------------------------------------------------|--------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------|
| MBN AI 앵커 (한국)  | <ul style="list-style-type: none"><li>실제 앵커(김주하) 음성·표정·제스처 딥러닝 합성</li><li>실제 앵커와 동시 방송 시연</li><li>24시간 반복 방송 자동화</li></ul> | <ul style="list-style-type: none"><li>재난·야간 뉴스 즉시 대응</li><li>제작 비용 절감</li><li>다국어·멀티채널 확장 용이</li></ul> | <ul style="list-style-type: none"><li>언캐니 밸리로 인한 시청자 거부감 발생</li><li>돌발상황·인터뷰·긴급 대응 불가</li><li>AI 영상이 딥페이크로 악용될 가능성</li><li>AI 오작동·편향 정보 → 대량 허위 보도 위험</li></ul> |
| 신화통신 AI 앵커 (중국) | <ul style="list-style-type: none"><li>세계 최초 AI 여성 앵커 공개</li><li>영어·중국어 등 다국어 버전 확장</li><li>일상 반복 뉴스 자동화</li></ul>          | <ul style="list-style-type: none"><li>빠른 뉴스 생성</li><li>다언어 콘텐츠 확장성 증가</li><li>국가 기술력 홍보 효과</li></ul>   | <ul style="list-style-type: none"><li>표정·억양 부자연스러움 지속</li><li>콘텐츠 획일화·창의성 부족</li><li>가짜뉴스 합성에 대한 악용 가능성</li><li>사회·정치적 편향 증폭 위험</li></ul>                       |

(출처: 연구자 구성)

### (3) 사례 분석의 종합: 공통 시사점과 구조적 문제점

앞서 분석한 5개 사례(버즈피드, CNET, 네이버, MBN, 신화통신)는 서로 다른 국가와 맥락에서 진행되었으나, 몇 가지 공통된 시사점과 구조적 문제점을 도출할 수 있다.

먼저, 공통적인 위험 요소로는 팩트체크와 인간 검증 부재라고 할 수 있다. CNET에서 발견된 77건 중 41건에 오류는 AI가 만든 콘텐츠에 대해 체계적인 사실 확인 절차가 없었기 때문에 발생한 사건이다. 버즈피드 또한 AI 콘텐츠에 대해 품질 기준 검증에 대한 우려가 나왔으며 신화통신과 MBN의 AI 앵커에 대해서는 돌발 상황 대처나 복잡한 맥락에서의 대응이 사람과 같지 않다는 한계를 보이는 예시이다. 이는 결국 HITL(Human-In-The-Loop) 검증 체계의 필요성을 부각시키는 것이라 할 수 있다.<sup>16)</sup>

다음의 투명성 부재는 신뢰도 추락의 대표적인 사례라 할 수 있다. CNET의 경우 AI를 쓴다는 사실을 독자에게도, 내부 직원에게도 알리지 않아 회사의 평판에 큰 타격을 입은 반면, 네이버는 ‘AI 활용’ 표기 제도를 통해 투명성을 확보하려 움직임을 보여 모범 사례로 꼽히기도 한다. 이는 AI 생성 콘텐츠는 결국 장기적 신뢰 구축이 필수라는 점을 보여주는 중요한 사례라고 할 수 있다.

세 번째로 분석한 사례들은 비슷한 자동화 파이프라인 구조를 갖는데 ‘스크립트·콘텐츠 생성 → 검증·팩트체크 → 영상·음성 합성 → 편집 → 배포’라는 반복적 흐름으로 구성된다. 버즈피드와 CNET은 텍스트 생성에, MBN과 신화통신은 영상·음성 합성에 초점을 맞추고 있지만, 모두 이 파이프라인의 어느 한 구간을 자동화한 것이다. 하지만 문제가 됐던 부분은 ‘검증·팩트체크’ 단계를 건너뛰거나 대충 처리한 경우 문제가 됐다는 점이다.

네 번째로, 노동 구조 재편을 둘러싼 사회적 갈등인데 버즈피드가 AI 도입을 알리면서 동시에 전체 인력의 12%를 해고한다고 발표하자, AI 자동화가 비용 절감 수단으로만 비춰지며 기존 인력에 대한 반발을 야기하는 동시에 “인간 기사를 대체하려는 것 아니냐”는 불안이 함께 퍼졌다. 이는 곧 사회적

16) 유호현, 유재연, 노가빈, 황혜민(2025). 「인간-AI 결합 에이전트의 시너지틱 인텔리전스 거버넌스 모델 제안」. 한국HCI학회 학술대회. 강원.

충돌로 붙어졌으며 AI를 도입할 때 재교육(reskilling)과 역할 재설계 계획을 함께 세워야 한다는 교훈을 남기기도 했다.

다섯 번째로, 주관적으로 보일 수 있지만 나라마다 규제 환경과 문화적 맥락이 다르다는 것이 구조적 문제점을 만들어 내기도 한다. 각 국가나 단체 간 AI 콘텐츠 거버넌스가 다르게 작동해 버즈피드와 CNET사태가 일어난 미국의 경우 Section 230 면책 조항과 자율규제가 지배적이어서 사전 규제 없이 AI 콘텐츠가 쏟아지다가 문제가 터지면 스스로 멈추는 사후 대응 패턴이 나타난다. 한국의 경우 선거 때, AI 생성 콘텐츠를 제한하거나, 플랫폼이 자율적으로 AI 활용 표기 제도를 두는 식의 국가 협력적 규제가 특징이라 할 수 있다. 중국의 경우 신화통신은 국가가 AI 기술 개발과 활용을 주도하는 경향을 보이며 딥페이크 규제처럼 정부 주도 규제가 시행된다. 하지만 이는 국가 선전 도구로 악용될 수 있다는 우려도 함께 나오고 있다.

정리하면, AI 영상 자동화 시스템을 문제없이 작동시키려면 몇 가지 설계 원칙이 필요해 보인다. 첫째, 팩트 체크 및 윤리 검토와 편집 승인 같이 인간 검증 모듈이 반드시 필요하다는 점이며 AI가 만든 콘텐츠에 대해서는 투명성을 담보하기 위한 표시가 반드시 뒤따라야 한다. 또한 표준화된 오케스트레이션 도구와 통신 프로토콜을 매끄럽게 연결하기 위해 각 나라마다 다른 규제 환경을 유연하게 대처할 수 있는 거버넌스 계층을 포함해야 한다는 것이다. 이는 4장의 통합 아키텍처의 뼈대가 될 것이다.

## 4. 통합 아키텍처 설계 및 구현 방안

앞서 사례 분석에서 드러난 팩트 체크의 부재, 투명성 미흡, 인간 검증의 생략과 같은 문제를 해결하기 위해 이 장에서는 멀티에이전트 기반 통합 아키텍처를 제안하고자 한다. 이는 로우코드 워크플로우 플랫폼과 MCP와 A2A 프로토콜, HITL 검증 계층을 하나로 엮는 구조로 구현하여 제안하고자 한다.

**1) 에이전트 오케스트레이션 계층:** 이 계층에서는 n8n이나 Make, Opal과 같은 자동화 도구를 활용하여 전체 프로세스를 관리한다.<sup>17)18)</sup> 예컨대 콘텐츠

생성 사례의 경우, 워크플로우는 “기사 주제 입력 → 관련 데이터 수집 → AI 초안 작성 → 인간 편집 검수 → 게시”의 단계를 포함할 수 있다. 오케스트레이터는 각 단계별로 적절한 모듈을 연결해주는 것으로 데이터 수집에는 API 호출 모듈, AI 초안 작성에는 LLM 호출 모듈(OpenAI 노드 등), 인간 검수 단계에는 승인 대기 모듈 등을 배치하여 시각적 플로우를 구성한다. 이렇게 오케스트레이션을 구성하면 복잡한 멀티에이전트 흐름도 한 눈에 관리할 수 있고, 예외적인 상황(예: AI 응답 저하나 오류 발생 시 대체 경로)에서도 어떤 모듈을 수정해야 하는지 쉽게 발견하고 설정할 수 있다. 이 계층은 비즈니스 로직을 코드 결합하는 것이 아니라 프로세스 제어면(Control Plane)에서 선(先)모델링 한다. 여기서 n8n이나 make, Opal은 트리거·분기·예외 복구·병렬화·승인 대기를 선언형으로 기술하는 상층 엔진이다.

**2) 도메인별 에이전트 구성:** 에이전트 계층 안에서는 작업 영역별로 특화된 서브에이전트를 두어 필요에 따라 A2A 프로토콜로 상태 정보를 주고 받으며 필요 작업을 상황에 맞게 생성하고 연결하여 작업을 진행한다. 예를 들어 영상 AI 앵커는 “스크립트 작성 에이전트”와 “영상 합성 에이전트” 두 개를 이용하여 스크립트 에이전트가 생성형 AI로 뉴스 원고를 쓰고, 영상 합성 에이전트가 TTS로 원고를 읽어준 뒤 CG 아나운서 얼굴에 입히는 식이다. 이들은 오케스트레이터 안에서 직렬로 A2A 프로토콜 형식으로 정보를 주고받는다. 마케팅 캠페인처럼 이미지와 텍스트가 모두 필요한 경우라면 “카피라이팅 에이전트”와 “이미지 생성 에이전트”를 병렬로 놓고 협력 절차를 진행할 수도 있다. 상황에 따라 도메인별 에이전트를 얼마든지 추가할 수 있다.

**3) MCP를 통한 도구 연계 및 컨텍스트 공유:** 이 작업은 각 에이전트가 LLM을 사용할 때 일관된 형식으로 맥락을 유지하기 위해 MCP로 도구를 연결하고 컨텍스트를 공유하는 방식이다.<sup>19)</sup> 가령 영상 에이전트의 MCP 세션이

17) 최승혁, 송민(2025). 「데이터 거버넌스 프레임워크 및 다중 생성형 AI 에이전트 오케스트레이션 프레임워크를 통한 빅데이터 자동 분석 시스템」. 한국경영정보학회 춘계통합학술대회. pp.185~191.

18) Dmitrievna, Y., & Parsadanyan, E. (2024). “AI agentic workflows: a practical guide for n8n automation”. *n8n Blog* [On-line]. Available: <https://n8n.io/blog/>

이전 인물 자료와 새로 등장하는 인물 자료를 톨로 정의하고, 영상이 진행되는 동안 해당 컨텍스트를 유지시켜 각 씬에 나오는 인물의 일관성을 가져가는 식이다. 이렇듯 여러 에이전트가 협업할 때, 앞선 에이전트의 최종 결과를 MCP 규격에 맞춰 구조화해 넘기면 다음 에이전트가 곧바로 연결해서 사용하는 것으로 워크플로우 전체의 모듈화와 유연성을 높이는 방법으로 사용될 수 있다.

**4) A2A를 통한 플랫폼 경계 넘나들기:** A2A는 서로 다른 플랫폼의 에이전트 간 협업을 위해 경계를 넘나들 때 사용하는 것으로 예를 들어 한 에이전트가 사내 데이터베이스 접근 권한이 없으면, 권한을 가진 외부 에이전트에 A2A 요청을 보내 필요한 정보를 받아오는 식으로 작동한다.<sup>20)</sup> 이에 대해 콘텐츠 생성 사례로 살펴보면, 사실 검증 에이전트를 따로 두고, 작성 에이전트가 기사를 쓴 뒤 사실 검증 에이전트에 A2A로 검증 요청을 보내 지식 베이스를 조회한 후 사실 관계를 확인해 응답으로 돌려줘 이를 반영한 기사를 고치는 식이다. 이런 교차 에이전트 통신은 팀원들이 이메일이나 API로 협업하는 것과 비슷한 원리로 이해하면 쉽다. 다만 A2A 통신을 도입할 때는 인증서 교환이나 권한 토큰으로 허가된 에이전트만 통신하게 하거나 민감 정보는 암호화하는 등의 안전장치도 함께 고려해야 한다.

**5) 인간 검증 및 거버넌스 계층:** 이 계층이 통합 아키텍처에서 가장 중요하다고 할 수 있는데, 인간검증 HITL 원칙에 따라 인간 검증 및 거버넌스 계층을 통합 아키텍처에 필수 단계로 두어 신뢰성과 투명성을 확보하는 방안이다. 아키텍처 안에 사람이 개입할 지점을 명확히 설정하여 최종 출고 전에 인간의 승인 노드를 넣어 내용을 확인하게 하고, 필요하면 피드백을 AI에게 다시 넘기는 루프를 완성하는 것이다. AI 뉴스 앵커라면 사람이 원고를 검토한 뒤 AI 합성 단계로 넘어가게 하거나, 실시간 방송 중에도 모니터링 인력이 개입해 중단할 수 있는 장치를 두는 방식으로 설계할 수 있다. 이런 거버

19) 백창원 (2025). 「AI 에이전트와 MCP를 활용한 데이터 보안 취약 대응 전략」. 『한국지능시스템학회 논문지』, 35권 5호, pp.465~474.

20) Arsanjani, A. (2025). “Complementary Protocols for Agentic Systems: Understanding Google's A2A & Anthropic's MCP”. *Medium* [On-line]. Available: <https://medium.com/>

년스 모듈은 워크플로우 엔진에서 조건 분기(if-else)나 수동 승인 블록으로 구현할 수 있으며 시스템 차원에서 로그를 남겨 추후 문제가 생겼을 때 어느 단계에서 어떤 결정을 거쳤는지 추적할 수 있게 하는 것이다.

위 요소들을 모아 본 논문에서 다룬 사례들을 지원할 수 있는 통합 아키텍처 개념도를 [그림 7]과 같이 완성할 수 있다. 실제로 해당 개념을 구현할 때 현실적으로 고려해야 할 점은 시스템 통합성과 성능 최적화에 중점을 두어야 할 것이다. 이 아키텍처의 강점은 재사용성과 확장성으로 한번 짜놓으면 멀티에이전트 파이프라인을 활용해 비슷한 작업에 변형해서 사용 가능하다는 것이다. 또한 새 AI 모델이나 API가 나오면 해당 부분만 새 에이전트로 바꾸거나 추가만 하면 되므로, 기술이 바뀌어도 유연하게 대응할 수 있는 것이 장점이라 하겠다.



[그림 7] 통합 아키텍처 다이어그램(출처: 저자 작성)

끝으로, 이런 시스템을 운영하기 위해서는 지속적인 모니터링과 피드백 루프가 필요하며, 실제로 해당 아키텍처를 돌려본 뒤 수집된 AI 응답 정확도나 처리량, 사용자 만족도와 같은 성능 데이터를 바탕으로 에이전트 프롬프트나 전략을 추가·수정하거나 규칙 기반 모듈과 AI 모듈 사이 균형을 조정하는 등 시스템 튜닝을 반복적으로 거쳐야 한다. 이는 본 아키텍처는 한 번 개발하고 끝나는 것이 아니라 운영 단계까지 관리하는 아키텍처라는 뜻으로 해석하면 좋을 것 같다.

## 5. 논의: 인간-AI 협업의 쟁점과 거버넌스

이 장에서는 본 논문이 제안한 아키텍처가 기술적인 것 뿐만 아니라 사회적 신뢰까지 얻으려면 어떤 설계 원칙과 제도적 장치가 필요한지 좀 더 논의해 보고자 한다. 논의 축은 정확성·책임성, 투명성·신뢰, 편향·검열, 노동과 역할 재편, 표준·규제, 인간 중심 설계로 분류해 보고자 한다.

정확성·책임성 부분에서 먼저 보자면, CNET 사례는 팩트체크와 편집 승인 없이 ‘그럴듯한 오류’가 쏟아지면 어떤 평판 비용을 치르게 되는지 적나라하게 보여준 사례다. 이에 제안 아키텍처에는 팩트체크 에이전트와 인간검증 노드를 필수 경유지로 넣었다. 이는 승인 로그, 근거 링크, 버전 이력으로 결정 과정을 기록하여 오류가 터졌을 때 어떤 부분을 정정하거나 추가해야 하는지를 알 수 있으며, 책임의 추적과 피해 최소화를 함께 달성할 수 있다.

투명성·신뢰 측면에서는 네이버의 “AI 활용” 표기가 플랫폼 수준의 표준으로 자리잡을 수 있는 가능성을 보여준 사례이며 다층 라벨링의 필요성을 부각시키는 동시에 이를 숨기면 장기적으로 더 큰 불신을 산다는 것을 보여주며 사전 공개가 원칙이어야 하는 이유를 제공해 주었다.<sup>21)</sup>

편향·검열·선전 문제에 있어서는 특히 영상 뉴스에서서 파급력이 큰 것을 볼 수 있다. 신화통신 AI 앵커 사례는 정책 집행 에이전트가 콘텐츠에 얼마나 큰 영향을 미칠 수 있는지 일깨워주며 이에 대응하기 위해 정기적으로 편

21) 임가영, 김애경(2025). 「지각된 신뢰성의 역할: 인간과 AI 에이전트 선택에서 서비스 영역의 상호작용 효과」. 『지식경영연구』, 26권 3호, pp.151~174.

향 여부를 점검하는 외부 감사와, 규칙이 필요함을 보여준 사례이다.

노동과 역할 재편 관점에서는 자동화가 일자리를 없애기도 하고 새로 만들기도 하기 때문에 조직은 재교육과 업무 재설계를 통해<sup>22)</sup> 인간의 전문성을 AI-증강(augmentation) 형태로 재배치해야 함을 보여준다<sup>23)</sup>.

표준·규제 관점에서는 MCP나 A2A 같은 공개 표준을 쓰면 서로 다른 플랫폼이나 시스템에 쉽게 연결되어 나중에 감사나 검증이 수월해질 수 있음을 확인할 수 있다. 이런 표준화는 내부 통제를 넘어 외부 규제 준수 비용을 줄이는 효과도 가져올 수 있다.

끝으로 **인간 중심 설계**의 원칙을 다시 강조한다. MBN의 AI 앵커 사례는 **언캐니 밸리**의 용어를 사용하며 우리의 목표가 ‘사람과 완전히 구분되지 않게 만드는 것’이 아니라는 점을 말하고 있다<sup>24)</sup>. 영상 자동화에 있어 ‘AI가 만든 영상’임을 밝히고 우리가 구분할 수 있게 만들어 시청자의 거부감을 줄여야 한다. 또한 실제 운영에서는 문제가 발생하면 곧바로 멈출 수 있어야 하며, 필요에 따라 즉시 사람이 개입할 수 있도록 사용자 경험(UX) 관점에서 설계되어야 한다. 그리고 시청자의 감정과 맥락을 섬세하게 읽고 조정하는 일은 온전히 인간의 책임으로 남는다는 점이다.

물론 본 논문에서 제안한 구조가 **참조 아키텍처**인 만큼 조직의 데이터·규제·문화적 맥락에 맞춘 커스터마이징(Customizing)이 선행되어야 한다.

## 6. 결론

이 논문은 생성형 AI 에이전트를 활용한 영상 자동화 사례를 분석하여 각 사례에서 드러난 활용 가능성, 이점, 그리고 한계와 위험을 파악하고, 자동화

22) 최다혜, 구유리(2025). 「IT 서비스 협업 환경을 위한 AI 에이전트 서비스디자인 전략 연구: 경계확장자 역할 지원 중심」. 『Journal of Integrated Design Research』, 24권 3호, pp.113~130.

23) 최은진, 손민영(2025). 「AI 에이전트 시대의 영화 시나리오 저자성 재구성: 공동 창작과 산업적 함의」. 『멀티미디어학회논문지』, 28권 8호, pp.1136~1145.

24) 전해리, 윤명세, 김성희(2025). 「전문가형 AI 에이전트의 프로필 이미지 의인화 정도에 따른 사용자 경험 연구」. 한국HCI학회 학술대회. 강원.

구조의 실패 지점을 확인하여 보완할 수 있는 통합 아키텍처를 제안하였다. 이를 통해 미디어 콘텐츠 생성 시 사회·윤리적 문제를 해결할 수 있는 선례를 제시했으며 영상 콘텐츠 자동화 사례 분석을 통해서도 기존 작업에 AI를 어떻게 구조화하여 적용할 수 있는지 그 방안에 대해 고찰하였다.

결국 공통 파이프라인 구조를 영상 자동화의 관점에서 일반화하는 것을 시도하였고, n8n/Make/Opal 오케스트레이션, MCP/A2A 통신, HITL 거버넌스 계층을 아우르는 체계도를 완성했다. 이 아키텍처는 개별 사례를 넘어 뉴스나 광고 등 다양한 영상 제작의 자동화 시나리오를 통해 재사용 가능한 프레임워크를 제공하였다.

다만 기술 구현을 넘어 정확성·책임성, 투명성·신뢰, 편향·검열 등 인간과 AI 협업 거버넌스의 핵심 쟁점을 도출한 것 또한 연구의 의미가 있다고 하겠다.

앞으로의 연구에서는 제안한 아키텍처를 프로토타입으로 구현해 성능, 안정성, 사용자 경험에 대한 실험이나 연구를 통해 정량적으로 평가할 필요가 있으며 시청자와 이용자를 대상으로 한 AI 영상에 대한 신뢰 수용 여부나 거부감 등에 대한 실증적 조사를 진행하여 문화나 정치 환경이 다른 여러 나라에서 AI 영상 자동화가 어떤 거버넌스 구조를 낳는지 비교하는 연구도 가치가 있을 것이다.

결론적으로, 에이전트 기반 영상 자동화는 콘텐츠 산업의 효율성과 창의성을 동시에 끌어올릴 잠재력을 가지고 있다. 다만 구현 과정에서 기술적 설계와 사회적 책임을 함께 요구한다는 점에서 해결해야 할 문제가 있음은 연구의 한계로 남는다. 이에 본 논문이 AI 영상 생태계에 이론적·실무적으로 지침이나 구조적 도움이 되길 기대한다.

## 참고문헌

- Aldridge, N., Brooker, M., & Sivasubramanian, S. (2025). "Open Protocols for Agent Interoperability Part 1: Inter-Agent Communication on MCP". *AWS Open Source Blog* [On-line]. Available: <https://aws.amazon.com/blogs/opensource/>
- Arsanjani, A. (2025). "Complementary Protocols for Agentic Systems: Understanding Google's A2A & Anthropic's MCP". *Medium* [On-line]. Available: <https://medium.com/>
- Bain & Company. (2023). "Bain announces alliance with OpenAI; Coca-Cola first client". *Bain & Company* [On-line]. Available: <https://www.bain.com/>
- Coca-Cola Company. (2023). "Coca-Cola invites digital artists to 'Create Real Magic' using new AI platform". *The Coca-Cola Company* [On-line]. Available: <https://www.coca-colacompany.com/>
- Dmitrievna, Y., & Parsadanyan, E. (2024). "AI agentic workflows: a practical guide for n8n automation". *n8n Blog* [On-line]. Available: <https://n8n.io/blog/>
- "MBN introduces Korea's first AI news anchor". (2020). *JoongAng Daily* [On-line]. Available: <https://koreajoongangdaily.joins.com/>
- Paul, K., & agencies. (2023). "BuzzFeed to use AI to 'enhance' its content and quizzes - report". *The Guardian* [On-line]. Available: <https://www.theguardian.com/>
- Sato, M., & Roth, E. (2023). "CNET found errors in more than half of its AI-written stories". *The Verge* [On-line]. Available: <https://www.theverge.com/>
- Wood, C. (2024). "5 AI Case Studies in Customer Service and Support". *VKTR* [On-line]. Available: <https://vktr.com/>
- "China's Xinhua presents news using robot news anchor". (2019). *Reuters* [On-line]. Available: <https://www.reuters.com/>

- 강정현, 박희영, 유희주, 정태양, 황동욱(2025). 「VR 기반 분리배출 교육 효과 연구 제안: AI 에이전트와 가이드북 비교연구」. 『디지털콘텐츠학회논문지』, 26권 7호, pp.1843~1852.
- 백창원(2025). 「AI 에이전트와 MCP를 활용한 데이터 보안 취약 대응 전략」. 『한국지능시스템학회 논문지』, 35권 5호, pp.465~474.
- 유호현, 유재연, 노가빈, 황혜민(2025). 「인간-AI 결합 에이전트의 시너지틱 인텔리전스 거버넌스 모델 제안」. 한국HCI학회 학술대회. 강원.
- 이선재(2025). 「AI 에이전트의 발전과 연구 동향」. 『한국통신학회지(정보와통신)』, 42권 11호, pp.43~52.
- 임가영, 김애경(2025). 「지각된 신뢰성의 역할: 인간과 AI 에이전트 선택에서 서비스 영역의 상호작용 효과」. 『지식경영연구』, 26권 3호, pp. 151~174.
- 전혜리, 윤명세, 김성희(2025). 「전문가형 AI 에이전트의 프로필 이미지 의인화 정도에 따른 사용자 경험 연구」. 한국HCI학회 학술대회. 강원.
- 정동균, 이종화(2025). 「설명 가능한 GPT 주식투자분석 시스템 연구: RAG 구조와 멀티 AI 에이전트 통합 설계」. 『경영정보학연구』, 27권 3호, pp.243~266.
- 정유림(2025. 5. 20). 「“AI로 만든 콘텐츠입니다”...네이버 ‘AI 활용 콘텐츠’ 표시 도입」. 『아이뉴스24』 [On-line]. Available: <https://www.inews24.com/>
- 최다혜, 구유리(2025). 「IT 서비스 협업 환경을 위한 AI 에이전트 서비스디자인 전략 연구: 경계확장자 역할 지원 중심」. 『Journal of Integrated Design Research』, 24권 3호, pp.113~130.
- 최승혁, 송민(2025). 「데이터 거버넌스 프레임워크 및 다중 생성형 AI 에이전트 오케스트레이션 프레임워크를 통한 빅데이터 자동 분석 시스템」. 한국경영정보학회 춘계통합학술대회, pp.185~191.
- 최은진, 손민영(2025). 「AI 에이전트 시대의 영화 시나리오 저자성 재구성: 공동 창작과 산업적 함의」. 『멀티미디어학회논문지』, 28권 8호, pp. 1136~1145.
- 최종근(2025). 「교육 현장에서 LLM 기반 AI 에이전트의 활용 가능성: 과학

글쓰기 평가를 중심으로』. 『대한지구과학교육학회지』, 18권 2호, pp. 124~139.

## Case Studies of AI-Agent-Based Automation: Focusing on Video Content Automation\*

Jeon, Joon Hyun / Hansung University\*\*

The rapid advancement of generative AI is accelerating the automation of end-to-end video production, and AI agents are increasingly used to create and distribute news, advertising, and tutorial videos. This study investigates automation cases based on generative AI agents, with a primary focus on video automation, and analyzes their orchestration structures and human-AI collaboration governance.

First, we conceptualize low-code workflow tools such as n8n and Make as the orchestration layer of multi-agent systems, and theoretically examine how Anthropic’s Model Context Protocol(MCP) and Google’s Agent-to-Agent(A2A) protocol standardize internal and external communication between agents. We then analyze text/news automation cases at BuzzFeed, CNET, and Naver, AI news anchor systems at MBN and Xinhua and others. From these cases we derive a common agentic workflow consisting of script generation, fact-checking, video synthesis, editing, and distribution, together with Human-in-the-Loop(HITL) checkpoints. On this basis, we propose an integrated reference architecture for agent-based video automation that combines n8n/Make-based orchestration, MCP/A2A-based agent communication, and human verification modules for fact-checking, ethical review, and editorial approval.

Finally, we discuss key governance issues—including accuracy and responsibility, transparency and bias, labor restructuring, and regulation and standardization—and argue that technical design and social accountability must be

---

\* This research was financially supported by Hansung University.

\*\* Hansung University College of Creativity and Convergence, Associate Professor, naturaljeon@hansung.ac.kr

addressed jointly at the architecture level.

**Key words:** generative AI, AI agents automation, human-AI collaboration

투고일자: 2025년 11월 16일

심사일자: 2025년 11월 21일

게재확정일자: 2025년 11월 30일