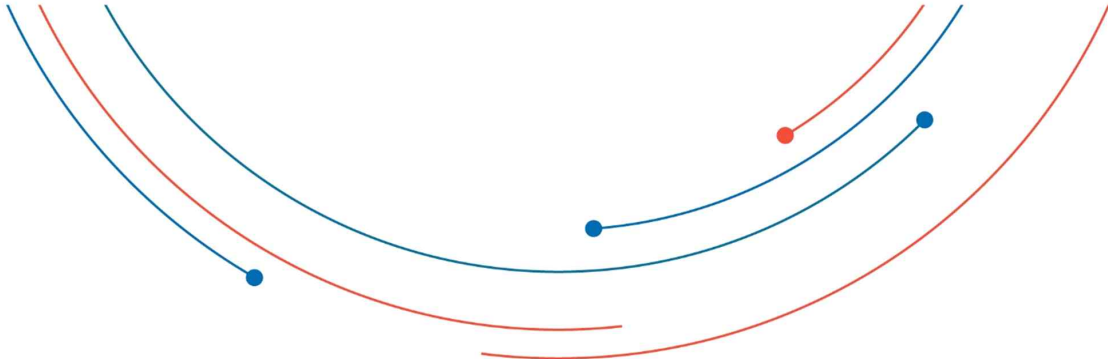




AI.GOV 이슈리포트 2026-1호

공공부문 AI 통제를 위한 사람 감독(Human Oversight)의 해외 사례 분석 및 시사점



NIA 한국지능정보원

일러두기 NOTE

「AI.GOV 이슈리포트」는 공공부문의 인공지능 전환 및 디지털정부와 관련된 다양한 이슈를 분석하고 향후 정책 방향을 모색하기 위해 한국지능정보사회진흥원에서 기획하여 발간하는 보고서입니다.

한국지능정보사회진흥원의 사전 승인 없이 본 보고서의 무단전재나 복제를 금하며, 가공·인용할 때는 반드시 출처를 명시하여 주시기 바랍니다.

본 보고서의 내용은 한국지능정보사회진흥원의 공식 견해와 다를 수 있으며, 문의 및 제안은 아래 연락처로 문의해주시기 바랍니다.

- 발행처: 한국지능정보사회진흥원
- 발행인: 김형철
- 작 성: 한국지능정보사회진흥원 인공지능정부본부 AI정부기획팀
박슬기 책임(psk64@nia.or.kr)
- 자 문: 한국IBM 백승현 전무
- 보고서 온라인 서비스: www.nia.or.kr

목 차 CONTENTS

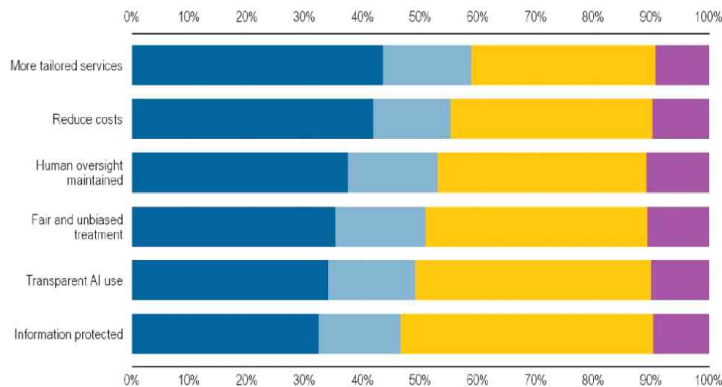
I. 배경 및 필요성	4
II. 공공 AI의 진화와 사람 감독 체계	6
1. 사람 감독의 유형과 AI 발전에 따른 변화	6
2. AI 활용과 사람 감독에 관한 국내외 법·제도	11
III. 공공부문 AI 사람 감독의 해외 사례	13
1. (영국) 내무부 난민심사 AI	13
2. (영국) Consult	14
3. (에스토니아) 뷰로크라트	15
4. (싱가포르) CLAIR	16
5. (미국) 보건부 AI 활용 심사·감독 체계	17
IV. 시사점	19
참고문헌	22

I 배경 및 필요성

□ 공공부문 인공지능(AI) 확산과 고위험 의사결정에서의 신뢰·책임성 요구 증대

- 민원, 복지, 조세, 인허가 등 공공행정 전 영역에서 AI 기반 의사결정이 빠르게 도입되면서, 기술적 성능 못지않게 신뢰성과 책임성 확보가 핵심 정책과제로 부상
 - AI는 속도 향상, 비용 절감 등 행정 효율성을 높이는 도구로 기대되나, 동시에 국민의 권리·기회에 직접 영향을 미치는 고위험 결정과 결부되어 엄격한 거버넌스 요구
- 네덜란드 SyRI, 미국의 COMPAS, 캐나다의 TRV, 호주의 Robodebt 사례는 공공부문 AI가 효율성 제고 수단으로 활용되더라도, 데이터 편향·검증 부재·설명 불충분·책임 회피가 결합될 경우 심각한 사회적 피해로 이어질 수 있음을 보여줌
 - ※ (네덜란드 SyRI) 복지 사기 적발용 자동화 시스템이 투명성·적법절차 문제로 위헌판결을 받음
 - (미국 COMPAS) 범죄 재범위험 예측 도구의 편향성과 설명책임 부족으로 공정성 논란
 - (캐나다 TRV) 비자 심사·판단 과정에서 오류 가능성 및 사람의 검토 부족 문제가 지적됨
 - (호주 Robodebt) 복지 부채 자동추징 시스템이 대규모 오청구와 심각한 사회적 피해 야기

| 공공부문 AI 활용에 대한 OECD 회원국 국민 신뢰 조사 결과(2026) |



※ 정부 AI의 서비스 개선·비용 절감에 대한 기대는 42~43%인 반면, 사람 감독의 유지·공정성·투명성·개인 정보 보호에 대한 신뢰는 32~37%에 그쳐 AI의 행정적 효용과 신뢰성 사이에 격차 존재 (OECD, Survey on Drivers of Trust in Public Institutions 2026 Results, 2026)

□ 공공 AI 거버넌스에서 사람 감독(Human Oversight)의 핵심 장치화

- 사람 감독(Human Oversight)은 AI의 설계·개발·도입·운영 전 과정에서 인간이 시스템의 작동과 결과를 검토하고, 필요할 경우 수정·중단·승인할 수 있도록 권한과

책임을 부여하는 구조를 의미하며, 공공부문에서는 기술적 보조수단이 아니라 민주적 통제와 행정 책임성을 확보하는 핵심 거버넌스 장치로 이해할 필요가 있음

- 사람 감독은 공공부문 AI 도입의 부수적 요소가 아니라 정당성과 책임성 확보를 위한 핵심 거버넌스 장치로 부상
 - World bank(2021)는 공공 AI가 정부 신뢰와 국민의 건강·안전·복지에 영향을 미치는 만큼 인간의 지휘 아래 운영되어야 하며, 사람 감독은 데이터와 알고리즘에서 발생하는 편향을 발견하고 통제하는 추가적인 안전장치라고 강조
 - UNESCO(2021)는 AI 시스템이 인간의 최종적인 책임과 책무성을 대체해서는 안 된다고 규정하고, 감사·추적·영향평가와 감독 체계를 마련하도록 권고
- 공공부문 AI 활용에 사람 감독을 요구하는 정책은 널리 채택되었으나, 일각에서는 인간의 과도한 의존, 책임 하위 전가 등을 초래하여 결함 있는 알고리즘 사용을 정당화하는 효과를 낼 수 있다는 비판 제기
 - OECD(2025)는 공공부문 AI 도입에서 발생할 수 있는 대표적 위험으로 자동화 편향을 제시하며, 특히 공무원이 AI 결과와 다르게 판단했을 때 더 큰 책임을 질 것을 우려하거나, AI의 권고를 거부할 경우 별도의 사유를 작성해야 하는 조직에서는 AI 결과를 따르려는 유인이 강화될 수 있다고 분석
 - 즉, 사람 감독이 단순한 최종 확인이 아니라 설계, 의사결정, 운영 감시 및 중단·지휘에 이르는 AI 생애주기 전반의 통제 체계여야 함을 의미

□ 국민 수용성 강화를 위한 공공부문 AI 사람 감독의 중요성

- 정부가 AI의 효율성을 제시하는 것으로는 국민의 신뢰수준과 수용성 향상에 충분하지 않으며, 누가 AI를 감독하고 결과를 변경·중단하며 국민의 이의제기를 재검토할 것 인지를 명확히 해야 공공 AI의 정당성을 확보할 수 있음
- 공공 AI의 사람 감독은 AI의 활용을 제한하기 위한 장치가 아니라, AI의 이점을 활용하면서도 국민의 권리 및 공공서비스에 대한 신뢰를 유지하기 위한 핵심 설계 요건
 - 본 보고서는 해외 사례를 통해 공공부문 AI 활용 전 과정에서 사람 감독이 어떻게 구현되고 있는지를 살펴보고, 공공 AI의 신뢰성·책임성·안전성을 확보하기 위한 사람 감독의 필요성과 정책적 시사점을 도출하고자 함

II

공공 AI의 진화와 사람 감독 체계

1 사람 감독의 유형과 AI 발전에 따른 변화

□ 사람 감독 방식의 변화 : RPA에서 Agnetic AI로

- RPA(Robotic Process Automation)는 사전에 정해진 규칙과 절차에 따라 데이터 입력, 문서 처리, 정산 등 반복 업무를 수행하고, 규칙에 없는 사례나 시스템 오류가 발생했을 때 사람이 개입
 - 사람은 자동화가 처리하지 못한 예외를 해결하는 ‘예외 처리자’ 역할을 수행하며, 사람 감독의 목적은 프로세스 중단을 최소화하고 업무의 완결성을 높이는 데 있음
- Agnetic AI는 정해진 절차를 실행하는 데 그치지 않고 목표의 해석, 계획 수립, 데이터 조회, 도구 호출, 결과 평가 및 다음 행동을 선택하며 이에 따라 오류의 성격도 달라지게 됨
 - RPA의 오류의 데이터 오입력이나 단계 누락처럼 규칙을 잘못 실행하는 ‘실행 오류’ 중심인 반면, Agnetic AI의 오류는 잘못된 기준을 선택하거나 편향된 해석에 기초하여 부적절한 권고와 행동을 생성하는 ‘판단 오류’로 나타날 수 있음
 - 판단 오류는 한 단계의 업무 실패에 그치지 않고, Agent가 잘못된 전제를 채택한 채 후속 시스템을 연속 호출하면 오류가 의사결정과 집행 전반으로 확대될 위험
- 따라서 Agnetic AI 시대의 사람 감독은 AI 자율성의 범위를 설정하는 감독자, 중요한 의사결정을 검토·승인하는 통제자, 그리고 기관의 책임 체계 안에서 결과를 설명하고 통제하는 의사결정 주체로 재배치되어야 함

| RPA와 Agnetic AI에서 사람 감독 방식의 차이 |

구분	AI의 역할	사람의 역할	주요 오류	사람 개입 및 감독의 목적
RPA	정해진 규칙 실행	예외 처리	규칙 실행·예외 처리 오류	자동화의 완결성 확보
Agnetic AI	목표 기반 판단·계획·행동	감독·승인·책임	판단 오류, 실행 오류 및 연쇄 오류	자율성 통제 및 정당성 확보

□ 개입 시점과 통제 강도에 따른 사람 감독 유형

- 사람 감독을 구현하는 방식에는 AI의 결과가 실제 효과로 이어지기 전에 사람이 직접 검토·승인하는 Human-in-the-Loop, AI의 자율 처리를 사람이 지속적으로 감독하고 필요시 개입하는 Human-on-the-Loop, 그리고 사람이 최종 지휘권을 보유하는 Human-in-Command 등이 포함*

* EU AI 고위 전문가그룹(HLEG, High-Level Expert Group on Artificial Intelligence)이 「Ethics Guidelines for Trustworthy AI(2019)」에서 제시한 사람 감독 거버넌스 인용

- 모든 공공 AI에 동일한 수준의 사람 감독을 일괄 적용하는 것은 오히려 행정 효율성이나 책임성을 약화시킬 수 있으므로, 사람 감독의 수준은 국민의 권리·안전에 미치는 영향, 오류의 회복 가능성, AI 결과가 실제 의사결정에서 차지하는 비중 등을 함께 평가하여 결정되어야 함

| 위험 기반 사람 감독 유형 비교 |

영향도 / AI활용범위	보조·추천 중심	자동처리·행동 중심
낮음·회복가능	Human-in-the-Loop	Human-on-the-Loop
높음·권리영향	강화된 Human-in-the-Loop	Human-in-Command 원칙 AI 자동 실행 제한

- 또한 EU, OECD 등은 사람 감독을 AI 운영 단계에만 두지 않고 기획·설계 단계에서 윤리적 요구와 통제 장치를 반영하여 사전에 위험을 예방해야 하며, 이를 위해 다양한 이해관계자와 전문가가 참여해야 한다고 제시
- 즉, AI의 목적·허용범위·데이터·성능기준·금지행위·사람 개입 지점을 설계 단계에서 사전에 결정하고 시스템에 반영하는 통제방식이 필요

① AI 기획·설계 단계의 사전 통제

- 기획·설계·학습 단계에서 사람이 AI의 목적, 데이터 사용 범위, 허용되지 않는 행동과 운영 조건을 정하는 것으로, 핵심은 일반적인 윤리 원칙을 선언하는 것이 아니라 설계 단계부터 사람이 AI가 사용할 기준과 자율성의 경계를 승인할 통제권을 갖는 데 있음
- 데이터·모델의 편향과 위험, 고영향 판단의 금지 범위, 운영 전 승인 기준을 명시함으로써 잘못된 설계가 실제 서비스에서 반복되는 것을 차단

- 설계 단계의 통제가 없으면 Human-in-the-Loop은 잘못된 설계 판단을 반복 승인하는 절차가 될 수 있고, Human-on-the-Loop은 문제가 발생한 이후에야 위험을 발견하는 체계가 될 위험이 있음
 - 다만 사전 검토만으로는 법령 변화, 데이터 변화, 예상하지 못한 사용 행태를 모두 통제할 수 없으므로 운영 단계의 승인 및 모니터링과 결합해야 함

② Human-in-the-Loop: 중요한 의사결정 직전의 검토·승인

- Human-in-the-Loop은 AI가 분석이나 권고를 생성하되 법적·행정적 효과가 발생하기 전에 사람이 결과를 검토하고 최종 결정을 내리는 방식
 - 복지, 채용, 인허가, 세무, 행정처분처럼 국민의 권리나 경제적 이익을 미치는 업무에서는 AI의 평균 정확도가 높아도 사람의 독립적 승인 단계가 필요
- 이를 위해 AI 결과가 곧바로 실행되지 않도록 사람이 승인하는 관문을 두고, 일정 위험도 이상이거나 정책상 예외 조건에 해당하는 건은 담당자에게 자동으로 전달되어야 함
 - 담당자는 AI가 사용한 핵심 데이터, 판단 근거, 신뢰 수준과 예외 사항을 확인할 수 있어야 하며, AI의 승인 또는 반려를 수정·취소할 수 있는 AI 결과의 수정·반복·무효화 권한을 가져야 함
- 실질적인 Human-in-the-Loop은 승인 버튼 하나를 추가하는 것으로 성립하지 않으며, 정보·권한·책임이 함께 설계될 때 실질적 통제로 기능
 - 사람이 AI 결과를 그대로 수용하는 자동화 편향을 줄이기 위해 반대 근거와 예외 조건을 함께 제시하고, 승인·반복 사유를 기록하며, 검토에 필요한 시간과 책임 범위가 보장되어야 함

③ Human-on-the-Loop: 자율 처리에 상응하는 상시 감독

- Human-on-the-Loop은 AI가 사전에 정한 범위에서 업무를 자율 처리하고 사람은 전체 운영을 감시하면서 이상 상황에 개입하는 방식
 - 민원 안내, 문서 요약, 내부 지식 검색, 단순 분류처럼 오류의 회복이 가능하고 대량 처리가 필요한 업무에 적합
 - 모든 건을 사전에 검토하지 않는 대신 감사 추적 기록, 이상 탐지, 정기 표본검사와

이의제기 경로를 통해 자율 처리의 품질과 위험을 지속적으로 확인

- 이는 사람 통제가 악화되는 것이 아니라 개별 승인에서 시스템 감독으로 이동하는 것임
 - AI Agent가 여러 도구와 시스템을 호출하는 경우 최종 출력만 저장해서는 충분하지 않으며, 어떤 데이터와 도구를 사용하고 어떤 중간 판단을 거쳐 다음 행동을 선택했는지 추적할 수 있어야 함
 - 위험 신호가 발생하면 자동 실행 권한을 즉시 회수하고 수동 처리로 전환하는 긴급 중단 장치도 필수적
- Human-on-the-Loop은 낮은 위험의 업무 또는 사후 회복이 가능한 업무에는 효과적이지만, 국민의 자격 박탈이나 제재처럼 영향이 크고 회복이 어려운 업무에 단독으로 적용할 경우 위험성이 커짐

④ Human-in-Command: 고영향 업무의 최종 지휘권 보유

- Human-in-Command는 AI의 사용 범위와 권한을 사람이 결정하고, 고영향 의사 결정과 실행에 대한 최종 지휘·통제권을 사람과 기관이 보유하는 원칙
 - AI를 분석 및 비의사결정적 지원에 한정하고, 최종 판단과 실행 명령 권한은 사람에게만 부여하는 통제 유형
 - 의료, 사법, 안보, 제재 처분처럼 오류가 생명·신체·기본권 또는 법적 지위에 중대한 영향을 미치고 사후 회복이 어려운 영역에 적용될 수 있음
- AI에게 독자적인 실행 권한을 부여하지 않고 사람이 AI의 권고를 받되 독립적으로 판단할 수 있도록 보장하는 데 있음
 - 담당자는 AI 권고를 따르지 않을 수 있어야 하고, 그러한 결정이 조직적으로 보호되어야 함. 이를 위해 전문성 교육, 권한과 책임의 명시, 충분한 검토 시간, 독립적 판단 근거의 기록이 기술적 통제와 함께 마련되어야 함

□ 사람 감독과 AI 감사의 연계

- 사람 감독이 AI의 설계·운영 과정에서 사람의 통제권을 행사하는 구조라면, AI 감사는 그 통제가 실제로 작동했는지 검증하는 체계
 - AI가 어떤 판단을 했고 사람이 무엇을 근거로 승인·수정했으며 최종적으로 어떤

행정 효과가 발생했는지를 기록하지 못하면 사람 개입은 책임성을 입증할 수 없는 형식적 절차로 전락

- 무슨 일이 있었는지 남기고(**감사 추적 기록**), 왜 그런 행동이 발생했는지 재구성하며 (**의사결정 추적 및 재현**), 현재의 위험을 어떻게 탐지하고 중단할지(**긴급중단장치 및 대시보드**) 관리할 수 있어야 사람 감독이 형식적 승인 절차가 아닌 검증 가능한 공공 AI 통제 체계로 기능할 수 있음

○ (**감사 추적 기록**) AI의 입력, 모델과 버전, 추론 결과, 신뢰 수준, 사람의 승인·반려·수정, 최종 결과를 시간 순서대로 기록하는 기능

- 기술 장애 분석을 넘어 민원 대응, 이의제기, 감사와 책임소재 규명에 필요한 최소 기반
- 로그를 단순 보유하는 데 그치지 않고, 해당 기록의 변경·삭제 여부가 검증되며 필요한 기간동안 검색·보존될 수 있도록 관리해야 함

○ (**의사결정 추적 및 재현**) AI가 특정 결과에 도달하기까지 사용한 데이터, 중간 판단, 정책·규칙, 도구 호출과 행동 순서를 추적하는 기능

- Agentic AI는 멀티턴(Multi-turn) 방식으로 계획을 수정하고 외부 도구를 연속 호출하므로 최종 출력만으로는 판단 경로 설명이 어려움
- 따라서 공공 Agentic AI에서는 단순 로그를 넘어 과거 실행의 주요 판단 경로를 세션 단위로 재구성·검증할 수 있는 의사결정 재현 환경이 필요
- 이를 통해 감사자는 결과가 틀렸는지 뿐만 아니라 어느 판단 지점에서 오류가 시작되었고 사람 통제는 왜 작동하지 않았는지 분석할 수 있게 됨

○ (**긴급중단장치**) AI가 비정상적으로 작동하거나 중대한 위험을 초래할 가능성이 있을 때 작동과 실행 권한을 즉시 회수하는 장치

- 단순한 시스템 종료 버튼이 아닌 신규 자동처리 중지, 외부 도구 접근 차단, 진행 중 업무의 격리, 수동 처리 모드 전환, 영향 범위 확인 및 복구절차를 포함
- 대량 자동처리가 가능한 Agentic AI에서는 작은 판단 오류가 반복·확산되기 전 사람이 개입할 수 있는 최후의 통제선으로 기능

○ (**AI 감독 대시보드**) 오류율, 사람 개입률, 승인·반복 비율, 신뢰도 분포, 예외 처리, 사용자 불만, 모델 성능과 데이터 분포 변화 등을 통합하여 보여주는 감독 도구

- 지표를 단순 시각화하는데 그치지 않고, 임계치 초과 시 자동 경보 및 이관하여 담당자가 승인 수준 강화, 서비스 중단과 같은 조치를 취할 수 있도록 연결

2 AI 활용과 사람 감독에 관한 국내외 법·제도

□ AI 활용시 사람 감독에 대한 법령 및 지침

- 「인공지능 발전과 신뢰 기반 조성 등에 관한 기본법」(26.1.22 시행)은 고영향 인공지능의 안전성과 신뢰성 확보를 위한 조치 중 하나로 사람의 관리·감독과 관련 문서의 작성·보관을 포함하고 있음
 - 관련 의무는 인공지능 사업자를 중심으로 규정되어 있으나, 발주처인 부처 및 공공기관은 이를 공공 AI 사업의 발주·도입·운영 요건으로 구체화함으로써 고영향 AI에 필요한 사람 감독을 실질적으로 구현할 필요가 있음

인공지능 발전과 신뢰 기반 조성 등에 관한 기본법

제34조(고영향 인공지능과 관련한 사업자의 책무)

제1항 4호. 고영향 인공지능에 대한 사람의 관리·감독

제1항 5호. 안전성·신뢰성 확보를 위한 조치의 내용을 확인할 수 있는 문서의 작성과 보관

- EU 「인공지능법(AI Act)」은 고위험 인공지능 시스템이 사용되는 동안 사람에 의해 효과적으로 감독될 수 있도록 설계·개발할 것을 의무화하고, 사람 감독의 목적·수준·방법과 감독자에게 보장되어야 할 구체적인 권한을 규정하고 있음
 - 사람 감독을 단순한 운영 절차가 아니라 고위험 AI가 갖추어야 할 필수적인 설계·개발 요건으로 규정하고, AI의 위험도·자율성 수준·사용 맥락에 비례하여 감독 조치를 마련하도록 함
 - 특히 감독자가 AI의 성능과 한계를 이해하고, 이상·오작동을 감시하며, AI 결과를 해석·거부·변경할 수 있을 뿐만 아니라 필요한 경우 시스템 작동에 개입하거나 안전하게 중단할 수 있어야 한다고 구체적으로 규정함
 - 이는 사람이 형식적으로 최종 승인하는 수준을 넘어, AI의 판단을 실질적으로 검토하고 수정하거나 시스템을 정지시킬 수 있는 권한과 기술적 수단을 갖추어야 한다는 의미임

EU 인공지능법(AI Act)

제14조(사람 감독, Human oversight)

제1항. 고위험 인공지능 시스템은 사용되는 기간 동안 적절한 인간-기계 인터페이스 도구 등을 통하여 자연인이 효과적으로 감독할 수 있도록 설계·개발되어야 한다.

제4항. 고위험 인공지능 시스템은 사람 감독을 담당하는 자연인이 적절하고 비례적인 범위에서 다음의 사항을 수행할 수 있도록 배포자·운영자에게 제공되어야 한다.

- 그 외 많은 국가는 법률상 사람 감독의 의무를 부과하기보다는 정부 지침·위험관리 프레임워크 등을 통해 AI의 위험 수준에 상응하는 사람의 개입과 감독 체계를 마련하도록 요구하고 있음
- 각국의 지침은 사람이 AI의 결과를 형식적으로 최종 승인하는 것만으로는 충분하지 않으며, AI의 한계와 오류 가능성을 이해한 담당자가 결과를 실질적으로 검토하고, 필요한 경우 판단을 변경하거나 AI 사용을 중단할 수 있는 ‘의미 있는 사람 감독 (meaningful human oversight)’이 필요하다는 방향을 공통적으로 제시함
- 이에 따라 AI의 최종 결과를 담당 공무원이 단순 확인하는 수준을 넘어, 업무의 위험도와 국민의 권리·안전에 미치는 영향에 따라 감독자 지정, 검토 대상과 개입 시점 설정, AI 결과의 수정·거부 권한, 이의제기 및 비상 중단 절차 등을 구체적으로 설계할 필요가 있음

| 해외 주요국의 AI 활용시 사람 감독에 관한 지침 등 |

국가	주요내용
미국	「NIST 인공지능 위험관리 프레임워크(NIST AI RMF)」를 통해 AI의 설계·개발·도입·운영 전 과정에서 인간의 역할과 책임을 명확히 하고, AI 위험을 최소화하기 위해 필요한 경우 사람의 개입이 가능하도록 관리할 것을 제시
영국	「혁신 친화적 AI 규제 접근법(A pro-innovation approach to AI regulation)」과 「영국 정부를 위한 AI 플레이북(AI Playbook for the UK Government)」을 통해 AI 활용기관이 안전성·책임성·투명성 원칙에 따라 사람의 감독과 책임 체계를 마련하도록 하고 있음
캐나다	자동화된 행정결정 시스템을 도입하기 전 의무적으로 알고리즘 영향평가 (Algorithmic Impact Assessment)를 실시하며, 영향도가 높을수록 사람 개입과 외부 검토 요건을 강화
싱가포르	「모델 AI 거버넌스 프레임워크(Model AI Governance Framework)」에서 AI 의사결정의 위험도에 따라 사람의 개입 수준을 구분하고, 조직이 적절한 사람 감독 방식을 사전에 결정하도록 권고하고 있음
일본	총무성과 경제산업성이 공동으로 마련한 「AI 사업자 가이드라인(AI Guidelines for Business)」을 통해 인간 중심의 AI 활용과 AI 활용 주체의 책임·감독 체계 구축을 요구하고 있음
호주	「정부의 책임 있는 AI 사용 정책(Policy for the Responsible Use of AI in Government)」을 통해 AI 영향평가·고위험 AI의 위원회 또는 고위 경영진 승인·운영 중 재검토 등 조직 차원의 다단계 감독 체계를 마련하도록 의무화

III 해외사례

1 (영국) 내무부 난민심사 AI

□ 개요

- 영국 난민 신청 증가와 심사 적체로 인해 심사관이 수십 페이지에 달하는 인터뷰 기록과 국가별 정책자료를 검토하는 데 많은 시간이 소요되는 문제 발생
 - 심사 품질은 유지하면서 반복 업무를 줄이기 위해 AI 기반 업무지원 도구 개발
- 핵심 도구는 Asylum Case Summarisation(ACS)와 Asylum Policy Search(APS)이며, 두 도구의 결과는 심사관에게 참고자료로 제공되고 난민 인정 여부를 자동으로 결정하지 않음

| 영국 난민심사 AI의 주요 도구 및 기능 |

도구	주요기능
Asylum Case Summarisation (ACS)	• 난민 인터뷰 녹취록을 LLM이 자동 요약 • 핵심 사실관계, 신청 사유 등을 정리하여 심사관에게 제공 • 심사관은 요약본과 원문을 함께 검토
Asylum Policy Search (APS)	• 국가별 정책정보를 생성형 AI가 검색·요약* * 단순한 웹검색을 통해 공개된 국가정보를 제공하는 것이 아닌, 내무부 내부에 축적된 국가정보 DB를 생성형 AI로 검색 및 요약 • 심사관 질문에 맞는 정책 근거를 대화형 방식으로 제공

- 난민심사 AI 도입으로 인터뷰 기록 검토 시간이 건당 23분, 국가별 정책정보 검색 시간이 평균 37분 감소하여 난민 심사의 효율성이 제고됨

※ '24년 시범운영 후 '26년 현재 APS는 전면 운영으로 확대, ACS는 단계적 확장 운영 중

□ 사람 감독 운영 방식

- 내무부는 이 사업을 '난민심사 자동화'가 아닌 심사관의 업무 지원으로 명확히 규정하고, AI 심사 결정을 하지 못하도록 제도적으로 설계

- 심사관은 모든 원본 인터뷰 자료와 증거를 반드시 직접 검토하고, AI 요약과 비교
 - AI 결과는 참고자료이며 법적 판단 및 결정문 작성의 책임은 심사관이 보유
 - AI 활용 이후에도 기존 품질관리, 증거 검토, 오류 검증 절차를 유지하여 AI가 행정 책임을 대신하지 않도록 설계
- AI는 문서요약·자료검색 등 반복 업무를 수행하고, 국민의 권리·의무에 영향을 미치는 최종 행정 판단은 사람이 수행
- 단순히 사람이 최종 결재하는 수준이 아니라, AI 결과를 검증·수정·배제할 수 있도록 사람 감독을 시스템 설계 단계에서 구현

2 (영국) Consult

□ 개요

- 영국 정부는 매년 약 500건의 정책·법령·규제 관련 공공협의를 실시하며, 수천 건 이상의 자유서술형 의견을 공무원이 직접 분류·분석하는 데 많은 시간과 비용이 소요
- AI를 활용한 분석 자동화를 추진하여 정책전문가가 정책 설계에 집중하도록 지원
- ※ 영국 과학혁신기술부(DSIT)와 정부디지털서비스(GDS)가 개발한 공무원 업무 지원용 생성형 AI 플랫폼 (도구 모음)인 Humphrey에 포함된 의견분석 AI 서비스
- 대규모 정책 협의 응답을 AI가 주제별로 분류·요약하고 핵심 의견을 도출하되, 전문가가 결과를 수정·재분류·검증하여 최종 분석을 확정하는 정책 협의 지원 도구

Consult 적용 workflow

- ① AI가 자유서술형 의견 응답의 성향을 분석하고 쟁점을 추출하여 초기 주제를 생성·통합
- ② 정책 전문가가 AI 응답을 검토하여 유사 주제 통합, 누락 주제 추가 등 보완
- ③ AI가 최종 주제 체계를 기준으로 개별 응답 분류하고 분류 근거 제시
- ④ 전문가가 주제와 분류 결과를 검증하고, 누락·오분류 및 주요·특이 응답 상세 검토
- ⑤ 사람이 전체 정책 맥락을 종합하여 최종 해석·권고 및 의사결정 수행

- (의견 자동분석) 국민이 제출한 자유서술형 의견을 생성형 AI가 분석하여 주요 주제, 쟁점 및 의견 유형을 자동 추출
- (주제별 분류) 유사한 의견을 자동 군집화하고 질문별 찬성·반대 의견 및 주요 이슈를 시각화하여 정책 담당자가 신속하게 검토할 수 있도록 지원

- (정책보고서 작성 지원) AI 분석 결과를 기반으로 정책 담당자가 협의 결과보고서 초안을 작성할 수 있도록 지원
 - (대시보드 제공) 질문별 주요 의견, 응답 비율, 핵심 쟁점 등을 대시보드 형태로 제공하여 정책 검토의 효율성 제고
 - 생성형 AI를 활용하여 대규모 공공협의 의견 분석 시간을 대폭 단축하고 정책 담당자가 심층적인 정책 검토에 집중할 수 있는 환경 조성
 - AI 결과는 전문가가 수행한 분석과 거의 동일한 수준의 주요 주제를 도출
 - 연간 약 7만 5천일의 행정업무와 약 2천만 파운드의 분석 비용 절감 효과 기대
- ※ Independent Water Commission 분석 사례에서는 5만건 이상의 응답을 약 2시간 만에 분류했고 AI 처리비용은 약 240파운드, 전문가 검증은 약 22시간 소요

□ 사람 감독 운영 방식

- Consult는 정책 결정을 수행하는 시스템이 아니라 정책분석을 지원하는 분석도로 활용
 - AI는 다단계 분석 워크플로우를 수행하고 사람의 수정사항을 반영해 다시 작업함
 - AI가 생성한 주제는 사용자가 직접 수정하거나 재분류할 수 있도록 설계하여 사람의 판단과 전문성을 유지하고, 분석의 신뢰성과 정책 책임성을 확보
- AI가 정책을 결정하지 않고, 사람은 분류기준·예외·품질·정책적 의미를 지속 통제

3 (에스토니아) бюрократ(Bürokratt)

□ 개요

- 에스토니아는 국민이 여러 정부기관 홈페이지를 일일이 찾아다니지 않고 자연어로 정부와 소통할 수 있도록 범정부 AI 기반 가상비서 бюрократ을 개발
 - 단일 챗봇이 아니라 각 정부기관의 AI 챗봇을 연결하는 범정부 AI 상담 네트워크를 구축하여 24시간 공공서비스를 제공하는 것을 목표로 추진
- 에스토니아의 공식 로드맵은 2026년부터 각 기관·분야가 개인화된 AI agent를 통합 네트워크의 일부로 사용하고, agent들이 독립적으로 소통하도록 발전시키는 것을 목표로 함
 - 장기적으로는 사용자의 요청에 따라 필요한 정보를 찾고 문의·신청·예약 등의 행동을

수행하는 구조를 지향

- 일상 언어 기반의 공공서비스 검색, 기관별 FAQ와 고객상담 자동화, 담당기관의 전문 챗봇 자동 연결 등 정부 데이터와 서비스를 연결해 단순히 관련 홈페이지 링크만 주는 것이 아니라 이용자가 대화 과정에서 서비스를 완료하도록 지원
 - 복잡한 민원, 개별 상황에 대한 판단, 예외적인 요청 등은 사람 상담원에게 이관되어, 사람 상담원은 앞선 대화 맥락을 확인 후 이어서 처리할 수 있음

□ 사람 감독 운영 방식

- 뷰로크라트는 개별 답변을 공무원이 일일이 승인하는 구조가 아니라, AI 상담 서비스 전반을 사람이 설계·관리하는 방식
 - AI가 국민에게 직접 답변하지만 답변 가능한 범위, 활용 데이터, 연결 기관 등은 모두 사람이 설계하여 AI는 정책을 전달하는 역할에 머무르도록 통제
 - 잘못 실행된 거래를 취소·복구할 수 있는 절차와 전체 행동 로그를 유지

4 (싱가포르) CLAIR

□ 개요

- 싱가포르 기업회계규정청(ACRA)는 기업의 재무제표를 대규모로 분석하여 회계 공시상 잠재적인 문제를 찾아내는 생성형 AI 규제 지원 도구 CLAIR를 개발
 - 단순히 재무제표를 요약하는 것이 아니라, 재무제표의 내용과 공시사항을 분석하여 규제기관이 추가로 확인해야 할 잠재적 공시 문제를 식별·선별
 - ※ 정식 출시 솔루션이 아닌 ACRA 내부 재무보고 감독 업무에 사용하는 생성형 AI 도구로, PoC로 시작하여 가치 검증을 거쳐 실제 규제 업무에 확장 적용 중
- CLAIR 도입으로 ACRA의 감독 가능 범위가 두 배 이상 확대되어 PoC를 넘어 실제 확장 가능한 규제 업무 도구로 발전
 - ACRA는 결과의 일관성과 규제 기준 준수 여부를 CLAIR의 핵심 성공 지표로 설정하며, 고영향 AI에서 단순히 모델의 정확도 뿐만 아니라 행정 기준에 따라 안정적으로 작동하는지 여부를 중요하게 관리

□ 사람 감독 운영 방식

- AI의 역할은 기업의 위반 여부를 최종 확정하거나 제재하는 것이 아니라 감독기관이 검토해야 할 대상을 발굴하고 쟁점을 제시하는 것
 - 사람은 AI가 표시한 문제를 회계기준, 기업의 구체적 사정 및 제출 자료와 대조하여 실제 규정 위반 여부와 후속 조치를 판단
 - AI가 작성한 문장을 단순 교정하는 것이 아니라, AI가 제시한 규제 위험을 전문적으로 검증하고 행정권한 행사 여부를 결정하는 방식으로 이루어짐

5 (미국) 보훈부 AI 활용 심사·감독 체계

□ 개요

- 미국 보훈부 사례는 단일 AI 서비스가 아니라, 보훈의료·급여·행정업무에 도입되는 모든 AI 사용사례를 등록·분류·심사·승인·모니터링하는 부처 차원의 AI 거버넌스 체계
 - 보훈부는 의료정보·장애정보·급여 등 민감한 데이터를 다루고, AI 결과가 재향군인의 권익에 영향을 줄 수 있어 보다 높은 수준의 안전·개인정보·책임성 관리 필요, 이에 따라 부처 차원의 AI 활용 거버넌스를 구축
- ※ 이는 미국 관리예산실(OMB)이 연방기관에 최고AI책임자 지정, AI 사용사례 목록관리, 고영향 AI 영향평가 등을 요구함에 따라 보훈부 차원에서 범정부 정책 이행 체계를 마련하는 일환으로 시행됨
- 사람 감독을 개별 감독 공무원의 최종 점검에 한정하지 않고, 단계별로 책임과 권한을 분산하는 조직 차원의 상호 견제 체계 구현
 - AI 사업 책임자, 독립 심사자, 전문조직, 최고AI책임자와 차관급 위원회가 단계별로 책임을 나누며, 심사 요건 불충족시 시정·중단을 요구할 수 있는 시스템을 갖추

| 미국 보훈부 AI 활용 심사·감독체계 |

단계	수행내용
사업부서 신청 및 자체검토	• 부서의 AI 사업 책임자가 사전 접수서를 제출 • AI 해당 여부, 고영향 여부 및 시스템 생애주기 단계 확인 • 고영향 AI는 AI 영향평가와 위험 완화 계획 추가 제출
독립 검토	• 고영향 AI의 영향평가와 위험 완화 계획은 사업 책임자가 작성하되, 별도의 심사자가 적합성을 검증 • 심사자에는 최고AI책임자 조직뿐 아니라 각 본부·참모조직이 지정한 인력이 참여하여 중앙의 기준과 현장 전문성을 결합

다분야 검토	법무, 개인정보, 재무 등 전문조직이 참여하여 위험성에 대해 다분야 검토 ※ 개인정보 침해, 사이버보안, 법률·행정절차 위반, 조달·계약상 책임, 재정적 위험, 재향군인 경험과 소수집단 경험 등
지휘·조정	차관급 위원회가 중요한 AI의 배포 전 시험과 성능평가에 의견을 제시하거나 직접 참여, 운영 후 실제 편익과 성과가 예상대로 나타나는지 검토
시정·중단	- 요건을 충족하지 못한 AI는 업무부서에 시정 요구 - 시정이 불가능할 경우 예외 승인을 검토하거나 운영에서 제거 - 필요시 보안 조직은 정보시스템 운영 승인 거부 또는 네트워크 연결 차단을 통해 사용을 중단하도록 설계

□ 사람 감독 운영 방식

- AI 결과를 공무원 한 명이 마지막에 확인할 경우 발생할 수 있는 한계*를 보완하고 책임 분산과 책임 명확화를 동시에 구현

* AI의 기술적 한계를 모두 이해하기 어려움, 개인정보·보안·법률 등 전문 영역의 영향을 동시에 판단하기 어려움, 자동화 편향 발생 가능성, 조직 전체에서 운영되는 AI의 연계 파악 곤란 등

- 책임을 분산하면서도 '누가 무엇을 책임지는지' 명확히 하여 책임 공백 방지
- AI 생애주기 전체에 사람 감독을 배치하여 AI의 도입부터 종료까지의 지속적 관리 구현

| 미국 보건부 AI 활용 심사·감독 체계의 주체별 역할 |

수행주체	역할
AI 사업 책임자	AI 활용목적과 업무성과 관리, AI 영향평가 수행
독립 심사자	위험관리 적합성 검증
전문조직	개인정보, 보안, 법률, 재무 등 전문분야 위험성 검토
AI거버넌스위원회	정책 조정, 중대 위험 관리

- 저위험 사업 등에는 상대적으로 간소한 절차를 적용하고, 고영향 AI에는 강화된 검토·승인·중단 체계를 적용하여 감독 역량 집중

- 사람 감독을 과도한 규제로 만들지 않으면서 고위험 분야에 조직의 전문성과 책임을 집중하는 위험비례형 감독의 사례
- 또한 기존 행정통제(보안 운영 승인, 개인정보 평가 등)와 AI 거버넌스를 결합하여, 형식적 심사가 아닌 기관의 일상적인 사업·예산·정보화 관리 체계에 내재화

IV 시사점

□ Human-in-Governance로의 확장

- 공공 AI의 사람 감독은 개별 업무 단계에 사람을 배치하는 수준을 넘어 기술적 통제와 조직적 책임, 행정적 절차를 통합하는 거버넌스 체계로 확장되어야 함
 - 기술적 측면에서는 AI의 판단 근거와 처리 이력을 확인하고, 결과를 수정·거부하거나 위험 발생 시 실행을 중단·권한 회수할 수 있는 기능 확보
 - 조직적 측면에서는 업무담당자, AI 책임자, 개인정보·보안·법무·감사 부서 및 고위 의사결정기구 간 역할과 책임, 보고·승인 경로를 명확히 설정
 - 행정적 측면에서는 AI 활용 사실과 책임 주체의 공개, 국민의 이의제기와 사람에 의한 재검토 절차 마련, 사고보고·감사·운영중단 절차를 기존 행정통제와 연계
 - 개별 AI 결과의 정확성뿐 아니라 누가 도입을 승인하고, 어떤 위험을 수용하며, 문제가 발생했을 때 누가 개입·책임질 것인지 관리하는 기관 차원의 감독 체계 구축 필요
- 이를 위해 AI의 설계·도입·운영·사후감사 전 과정에 사람의 권한과 책임을 배치하고, 조직 전체가 AI의 사용 범위와 위험을 지속적으로 통제하여 의사결정 구조에 사람의 책임을 내재화하는 Human-in-Governance 관점의 공공 AI 거버넌스 정립 필요

□ 행정적 영향에 비례한 감독 차등화

- 모든 AI에 동일한 승인·검토 절차를 적용하기보다, AI가 국민의 권리·안전 등에 미치는 영향에 따라 사람 감독의 수준을 달리하는 위험비례형 접근을 채택
 - 위험이 낮은 분야에는 간소화된 감독을 적용함으로써 통제의 실효성과 행정 효율성을 함께 확보하고, 고위험 분야는 오류의 발생 가능성뿐 아니라 피해의 중대성·비가역성·사회적 파급효과를 고려하여 더 강한 사람 감독을 적용
- 공공 AI의 사람 감독은 특정 의사결정 시점에 사람을 배치하는 것이 아닌 AI의 권한·위험·성과를 지속적으로 관리하는 체계로 확대하되, AI가 수행하는 기능과 그 결과가 국민에게 미치는 영향을 기준으로 사람 감독 수준을 설정할 필요

□ 사람 감독의 실효성은 인간의 존재보다 실질적인 권한 보장에 의해 좌우

- AI 처리 과정에 사람이 참여하더라도 AI 결과를 변경·거부·중단할 수 있는 권한과 충분한 정보가 없다면 형식적 감독에 그칠 가능성이 있으며, 실효성 있는 사람 감독을 위해서는 다음과 같은 권한과 조건을 갖출 필요가 있음
 - AI 결과와 근거자료를 확인할 수 있는 접근 권한
 - AI 결과를 수정하거나 채택하지 않을 수 있는 판단 권한
 - 중대한 오류·위험 발견 시 운영을 일시 중단하거나 상급자에게 이관할 수 있는 권한
 - AI와 다른 결정을 내렸을 때 그 사유를 기록할 수 있는 절차
 - 국민에게 최종 결정의 이유와 AI 활용 사실을 설명할 수 있는 정보
 - 자동화된 결과에 대한 국민의 이의제기를 사람이 재검토하는 구체 절차

□ 개인 담당자 중심의 감독에서 조직적·다층적 거버넌스로 전환할 필요

- 복잡한 공공 AI의 기술·법률·개인정보·보안·업무 위험을 최종 담당자 한 사람이 모두 파악하고 책임지는 것은 현실적으로 곤란
 - 미국 보훈부는 AI 사용부서, 사용사례 검토자, 개인정보·정보보안 조직 및 차관급 AI 거버넌스위원회가 역할을 분담하는 조직적 감독 체계를 구축
- 특히 공공부문 고위험 AI 의사결정은 오류 발생시 행정적 파급력 및 사회적 피해규모가 크므로 개인 차원이 아니라 조직 차원의 검증·책임 분산·재검증 문제로 접근 필요
 - AI가 권고·평가를 통해 행정 판단을 실질적으로 좌우할수록 실시간 사람 승인과 강한 번복권이 요구되며 이를 개인의 판단 및 책임으로 남겨둘 경우 더 큰 위험을 초래
- 고영향 AI의 사람 감독은 기술적·조직적·절차적 차원에서 동시에 구현되어야 함
 - (기술적) 승인 관문, AI 결과의 수정·반복·무효화, 감사 기록과 긴급중단의 시스템 내장
 - (조직적) 검토자·승인자·운영 책임자의 역할과 책임 구분, 사람이 AI의 권고를 거부할 수 있는 권한 보장
 - (절차적) 승인·반려·수정 사유와 이의제기·재검토 과정 기록, 중요한 모델 업데이트나 업무 범위 변경시 위험등급 및 통제 수준 재평가

□ 국민 접점에서는 설명·이의제기·사람 연결이 사람 감독의 핵심 요소

- 국민이 체감하는 사람 감독은 내부 위원회나 승인 절차보다, AI가 관여한 사실을 알 수 있고 결정의 이유를 이해하며 사람에게 문제를 제기할 수 있는지에 의해 좌우
 - 단순한 정보 제공을 넘어, 국민이 자신의 상황에 맞는 판단을 이해하고 대응할 수 있도록 돕는 ‘설명 가능성’과 ‘접근 가능성’을 동시에 확보하는 것이 중요
 - AI가 활용된 경우 해당 사실을 명확히 고지하고, AI가 수행한 역할과 인간이 수행한 역할을 구분하여 안내할 필요
 - 또한 복잡한 알고리즘이나 기술적 설명이 아닌, 행정적 판단의 흐름과 결과에 영향을 미친 요소 중심으로 설명이 제공되어야 공공 AI에 대한 국민의 신뢰 확보 가능
- 공공 AI의 신뢰 확보를 위해서는 내부 감독뿐 아니라 국민이 접근할 수 있는 설명·정정·이의제기·사람 재검토 절차를 함께 설계해야 함
 - 나아가 이러한 절차가 실제로 작동하는지를 지속적으로 점검하고 개선하는 것이 필요
 - ※ 국민의 이의제기 건수·인용률 등을 분석하여 제도의 실효성 평가, 설명의 이해도와 만족도에 대한 이용자 피드백 수집, 반복적으로 발생하는 오류나 불만 유형을 분석하여 AI 모델과 행정절차 개선에 반영
- 결국 국민 접점에서의 사람 감독은 기술적 통제 장치가 아니라, 국민이 행정 결정 과정에 참여하고 자신의 권리를 보호받을 수 있도록 하는 신뢰 기반의 제도적 장치로 설계되어야 함

□ ‘통제 가능성’ 중심의 공공 AI 관리체계 확립 필요

- 공공 AI의 핵심은 자동화 범위의 확대가 아니라, AI의 설계·운영·사후관리 전 과정에서 사람이 통제권을 유지해야 할 의사결정 지점을 명확히 설정하는 데 있음
- 이를 위해 공공 AI 사업 평가를 정확도·처리속도 중심에서 통제 가능성 중심으로 확대할 필요
 - 사람이 판단 근거를 이해할 수 있는지, 결과를 수정·반복할 수 있는지, 위험 발생 시 실행권한을 회수·중단할 수 있는지, 사후에 책임 경로를 재구성할 수 있는지 등
- 설계 단계의 AI 권한·경계 설정, 운영 단계의 사람 승인·개입·반복, 시스템 수준의 모니터링·중단 및 사후 감사가 하나의 체계로 연결될 때 공공 AI의 책임성과 신뢰성 확보 가능

참고문헌

- 1) OECD(2026), OECD Survey on Drivers of Trust in Public Institutions
https://www.oecd.org/en/publications/oecd-survey-on-drivers-of-trust-in-public-institutions-2026-results_9eb63fec-en/full-report/trustworthy-artificial-intelligence-in-the-public-sector_6f98c91a.html
- 2) World Bank(2021), Artificial Intelligence in the Public Sector
<https://openknowledge.worldbank.org/entities/publication/2c62df59-a40b-5d54-8d7d-3ee9d6c61715>
- 3) UNESCO(2021), Recommendation on the Ethics of Artificial Intelligence
<https://www.unesco.org/en/artificial-intelligence/recommendation-ethics>
- 4) NIST(2023), Artificial Intelligence Risk Management Framework
<https://airc.nist.gov/airmf-resources/airmf/appendices/app-c-ai-risk-management-and-human-ai-interaction/>
- 5) OECD(2025), Governing with Artificial Intelligence
https://www.oecd.org/en/publications/governing-with-artificial-intelligence_795de142-en/full-report/how-artificial-intelligence-is-accelerating-the-digital-government-journey_d9552dc7.html
- 6) European Data Protection Supervisor. (2025). TechDispatch #2/2025: Human Oversight of Automated Decision-Making
https://www.edps.europa.eu/data-protection/our-work/publications/techdispatch/2025-09-23-techdispatch-22025-human-oversight-automated-making_en
- 7) European Union. (2024). Regulation (EU) 2024/1689 Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act)
<https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- 8) High-Level Expert Group on Artificial Intelligence. (2019). Ethics Guidelines for Trustworthy AI. European Commission.
<https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>

- 9) OECD. (2023). Advancing Accountability in AI: Governing and Managing Risks throughout the Lifecycle for Trustworthy AI. OECD Digital Economy Papers, No. 349. OECD Publishing.
https://www.oecd.org/en/publications/advancing-accountability-in-ai_2448f04b-en.html
- 10) Haitsma, L., & Brink, B. (2025). From human intervention to human involvement: A critical examination of the role of humans in (semi-)automated administrative decision-making. Digital Government: Research and Practice, 6(3), Article 33. doi: 10.1145/3716173
<https://research.rug.nl/en/publications/from-human-intervention-to-human-involvement-a-critical-examination/>
- 11) European Commission(2021), Ethics by Design and Ethics of Use Approaches for Artificial Intelligence
https://ec.europa.eu/info/funding-tenders/opportunities/docs/2021-2027/horizon/guidance/ethics-by-design-and-ethics-of-use-approaches-for-artificial-intelligence_he_en.pdf
- 12) UK Home Office(2025), Evaluation of AI trials in the asylum decision making process
<https://www.gov.uk/government/publications/evaluation-of-ai-trials-in-the-asylum-decision-making-process/evaluation-of-ai-trials-in-the-asylum-decision-making-process>
- 13) UK DSIT (2025), Government-built “Humphrey” AI tool reviews responses to consultation for first time
<https://www.gov.uk/government/news/government-built-humphrey-ai-tool-reviews-responses-to-consultation-for-first-time-in-bid-to-save-millions>
- 14) Scottish Government (2025), Regulation and licensing of non-surgical cosmetic procedures: consultation analysis and response – Methodology
<https://www.gov.scot/publications/regulation-licensing-non-surgical-cosmetic-procedures-consultation-analysis-response/pages/3/>
- 15) UK Incubator for AI(2025), Consult Evaluation: Scottish Government’s

Non-surgical cosmetic procedures consultation

<https://ai.gov.uk/evaluations/consult-evaluation-scottish-government-non-surgical-cosmetic-procedures-consultation/>

- 16) UK DSIT/DEFRA (2025), Ground-breaking use of AI saves taxpayers' money and delivers greater government efficiency

<https://www.gov.uk/government/news/ground-breaking-use-of-ai-saves-taxpayers-money-and-delivers-greater-government-efficiency>

- 17) Estonian AI and Data portal, Bürokratt – Virtual assistant and roadmap

<https://www.kratid.ee/en/burokratt>

- 18) Smart Nation Singapore (2026), Update to NAIS

https://isomer-user-content.by.gov.sg/39/cb52d9a0-0d6c-484d-96fe-571031b37eb7/NAIS_update.pdf

- 19) VA Artificial Intelligence(2024),Department of Veterans Affairs Compliance Plan for OMB Memorandum M-24-10

<https://department.va.gov/ai/department-of-veterans-affairs-compliance-plan-for-omb-memorandum-m-24-10/>

- 20) VA Artificial Intelligence(2025), Department of Veterans Affairs Compliance Plan for OMB Memorandum M-25-21

<https://department.va.gov/ai/department-of-veterans-affairs-compliance-plan-for-omb-memorandum-m-25-21/>