

AI GOVERNANCE

LUNA™ AI-GRC + MCP Gateway

AI 거버넌스·실행 통제 플랫폼

AI 실행 통제, MCP 연결 검증, 보안감사 대응까지
하나의 거버넌스로 완성합니다.



AI를 쓰기 시작하면, 통제 밖 경로가 생깁니다

임직원과 AI Agent가 외부 모델과 사내 시스템에 직접 닿는 동안, 무엇이 오갔는지 확인할 지점이 마련되어 있지 않습니다.



● AI 도입이 만든 다섯 가지 공백

기존 보안 체계는 사람이 직접 쓰는 상황을 기준으로 만들어졌습니다. 이제는 AI가 대신 묻고 대신 처리하면서, **지금까지의 방식으로서는 확인하기 어려운 구간**이 생깁니다.

01

사내 자료가 외부 모델로 그대로 전송됩니다

문서를 붙여넣거나 파일을 올리면 그 내용이 외부 AI로 그대로 넘어가는데, 무엇을 보냈는지 회사에 남는 기록이 없습니다.

02

누가 어떤 AI를 쓰는지 알 수 없습니다

직원마다 개인 계정으로 따로 쓰다 보니, 어느 부서가 무슨 일에 AI를 쓰는지 회사가 파악하기 어렵습니다.

03

AI가 사내 시스템을 직접 건드립니다

요즘 AI는 사내 데이터를 조회하고 파일을 열고 다른 시스템과 연결되기도 하는데, 실행하기 직전에 권한과 요청 내용을 확인할 지점이 없습니다.

04

AI 관련 규제에 필요한 증빙을 남기기 어렵습니다

안전 조치, 사용 현황 점검, 기록 보관, 처리 이유 설명이 필요한데, 사용 이력이 여기저기 분산되어 있어 자료를 모으기 어렵습니다.

05

기록이 분산되어 사고 경위를 파악하기 어렵습니다

서비스마다 기록 방식과 보관 위치가 달라서, 문제가 생겼을 때 앞뒤 상황을 한 화면에서 이어 보기 어렵습니다.

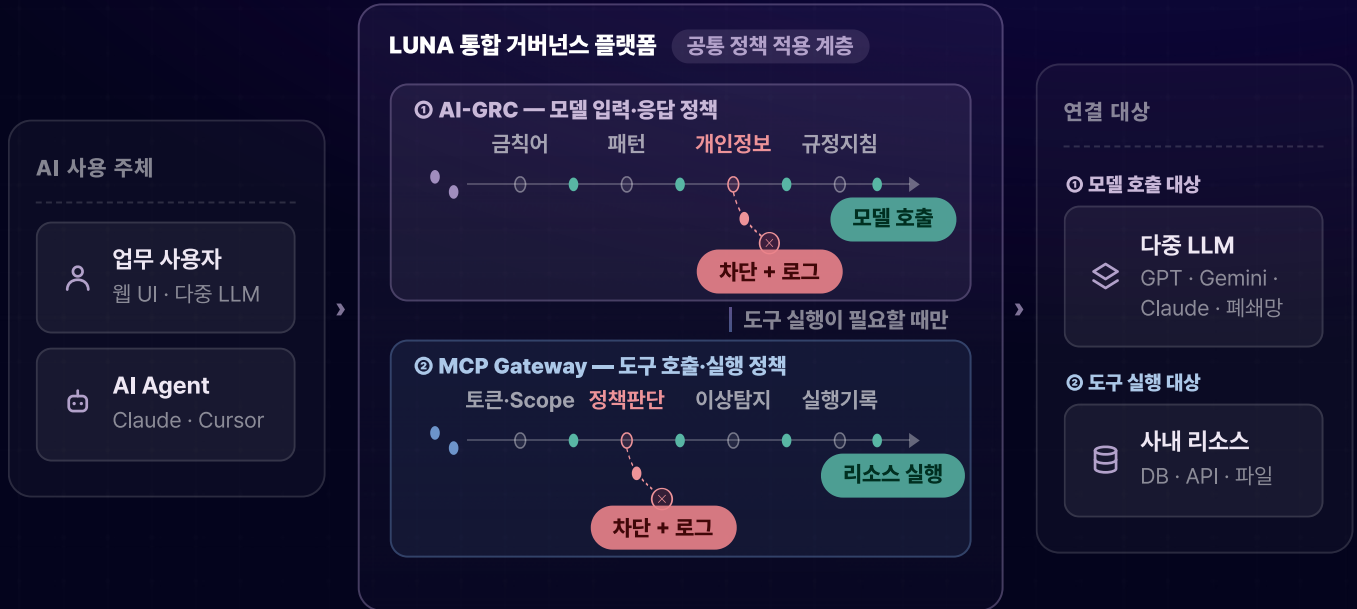
06

필요한 것은 사용 금지가 아니라, 오가는 내용을 확인할 지점입니다

AI가 주고받는 내용과 실제로 하는 일을, 하나의 기준과 기록으로 모읍니다.

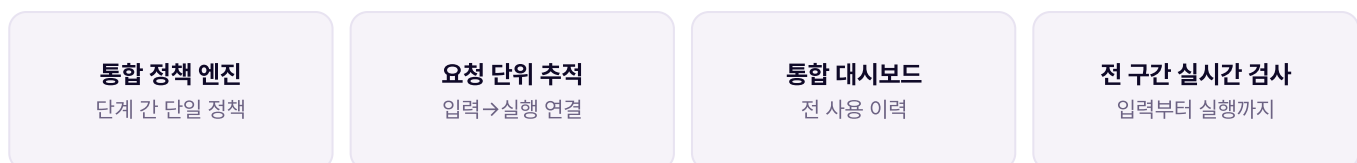
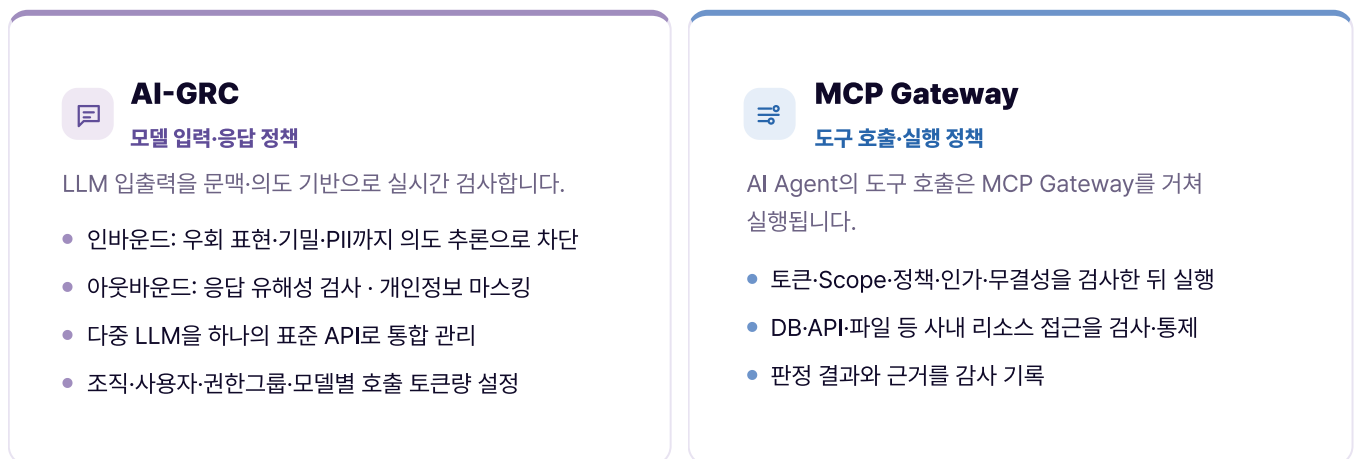
모델 호출부터 도구 실행까지, 하나의 정책으로 관리합니다

AI-GRC가 입력·응답을 판단하고, 도구 실행이 필요하면 MCP Gateway가 이어받아 권한과 정책을 검증합니다.



요청부터 실행까지, 하나의 정책 흐름

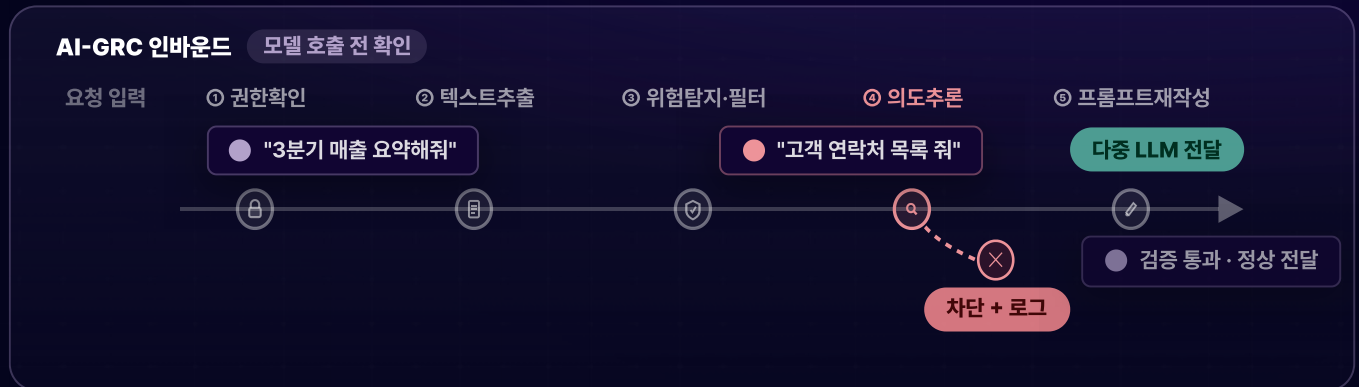
AI-GRC가 **모델 입력·응답**을 판단하고, 도구 실행이 필요하면 MCP Gateway가 이어받아 **권한·정책**을 검증합니다. 두 단계 모두 실시간으로 검사하고 감사 기록을 남깁니다.



모델 호출도 도구 실행도, 판단과 기록은 하나의 플랫폼에서 이루어집니다.

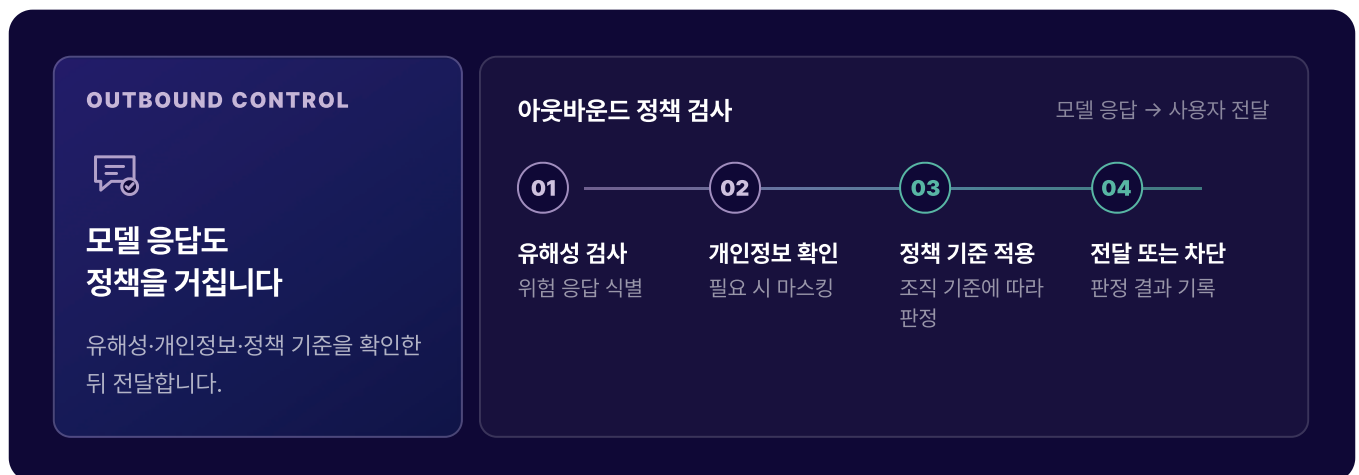
프롬프트가 모델에 닿기 전, 다섯 단계를 거칩니다

권한 확인부터 의도 추론까지 위험 요소를 확인한 뒤에만 모델로 전달하고, 처리 결과는 감사 기록으로 남습니다.



● AI-GRC — 입력과 응답을 연결하는 정책

입력 단계에서는 위험을 확인하고, 모델 응답은 다시 검사합니다. 정책 판단과 처리 결과는 **감사 기록**으로 남습니다.



 **다중 LLM 연동**

ChatGPT·Gemini·Claude 등 SaaS 모델
OpenAI 호환 표준 API로 연동

 **배포 환경**

온프레미스·폐쇄망 구축 지원
외산 대비 국내 규정 대응에 유리

 **사용량 관리**

조직·사용자·권한그룹·모델별
호출 토큰 쿼터 설정

1,000+ 문서 형식

이미지·오피스 문서 추출

우회 표현 탐지

의도 추론 기반

감사 기록 보존

규제 대응 기반

이상행위 탐지

비정상 패턴

입력부터 응답까지, 판단과 기록은 **AI-GRC** 하나로 관리됩니다.

도구 호출이 실행되기 전, 다섯 단계를 거칩니다

토큰 검증부터 톨 무결성까지 위험 요소를 확인한 뒤에만 실행하고, 처리 결과는 감사 기록으로 남습니다.



● MCP Gateway — 도구 호출·실행 거버넌스

AI Agent의 도구 호출은 5단계로 위험을 확인한 뒤 실행되고, 실행 결과와 응답도 다시 감사합니다. **Agent별 권한 분리**와 **실행 대상 관리**도 하나의 화면에서 이루어집니다.

실행 대상 도구 실행 범위

검증을 통과한 호출만 실제 도구로 전달됩니다.

- DB MCP · External API(Jira·Calendar 등) · File MCP
- 실행 결과와 응답도 감사 기록으로 보관
- 이상행위 탐지 대상에 도구 실행 이력 포함

Agent 연동·권한 분리 Scope 기반 권한 관리

Codex·Claude·Cursor 등 다양한 Agent를 Scope 단위로 구분해 관리합니다.

- Global·Group·Single 단위로 Scope 분리
- 톨 스키마 해시로 변조·포이즈닝 탐지
- 성공·차단 모두 감사 기록

● 도구를 더 쉽게 연결합니다 — 묶음 제공 · 자동 추천

도구 묶음 제공 · Tool Group

자주 쓰는 도구를 미리 묶어 신청하면, 토큰 하나로 그 묶음 전체를 한 번에 불러옵니다.

필요한 도구 자동 추천

내장된 AI가 요청 내용과 사용 목적을 살펴, 지금 필요한 도구를 골라 추천합니다.

도구 호출부터 실행까지, 판단과 기록은 **MCP Gateway** 하나로 관리됩니다.

국내외 AI 규제에, 하나의 플랫폼으로 대응합니다

AI기본법 4대 의무부터 N2SF·개인정보보호법·EU AI Act까지, 별도 솔루션 없이 하나의 플랫폼에서 대응합니다.

● 산업별 규제 대응 매핑

AI기본법은 이미 전면 시행 중이고, 금융·공공은 소관 부처의 별도 기준이 추가로 적용됩니다. 필요한 것은 선언이 아니라 **통제 근거와 기록**입니다.

산업·대상	적용 기준·요구 사항	LUNA 대응
전 산업 공통 AI기본법	2026.01.22 전면 시행 · 시행령 2026.07.21 시행 — 고영향 AI 위험관리, 이용 기록 5년 보관	공통 프롬프트·응답·툴 호출을 사람 단위로 전수 기록
금융·보험·증권 금융위·금융보안원	2026.06.22 시행 — 거버넌스·보안성 등 7대 원칙, 폐쇄망 완화 시 대체 통제 수단	공통 정책 엔진 기반 실시간 차단·통합 관제·폐쇄망 대체 통제
금융 AI 에이전트 금융보안원 위협 분석	도구 오염, 간접 프롬프트 인젝션, 연동 도구 확대에 따른 공격 표면 증가	MCP Gateway 호출 인자 검사·스키마 무결성 검증·권한 범위 제한
공공·공기업 행안부·국정원	행정 활성화법 개정 2026.08.28 시행 — 공공부문 AI 도입 가이드, N2SF 등급별 보안 통제	공통 폐쇄망 On-Premise 구축·등급별 정책 분리 운영
개인정보 처리 개인정보보호위원회	생성형 AI 개인정보 처리 안내서, 처리방침 작성지침 개정 — 학습 활용 여부 고지·거부 절차	AI-GRC 입력 탐지·차단·응답 마스킹·이력 보관

폐쇄망 On-Premise

외부 연결 없이 내부망 구축

단계적 도입

필요한 축부터 확장

기존 시스템 연동

운영 중인 보안체계와 연동

AI 거버넌스, 하나의 플랫폼으로 시작하세요

프롬프트 통제부터 도구 실행 통제까지, LUNA 하나로 시작할 수 있습니다.

Company

(주)오픈소프트랩

Tel

02-6941-0501

Mail

osl@opensoftlab.kr



주소 서울 금천구 가산디지털1로 196, 907호 · 제품 상담 및 도입 문의는 QR 또는 osl@opensoftlab.kr로 연락해 주세요.