

AI-FLOW

기획·실행 기록 · 4주, 혼자, 어떻게 만들었나 — 작업 방식의 기록

항목	내용
프로젝트	인천공항 AI-PORT 대국민 아이디어 공모전
기간	2026.03.09 — 04.03 (4주)
역할	개인 참가 · 기획 · 데이터 분석 · 시스템 설계 · 결과물 제작
작성자	김형준

웹 페이지가 "무엇을 만들었고 왜 그렇게 설계했는가"를 다룬다면, 이 문서는 그 뒤에 가려진 "어떤 순서로, 어떤 기준으로 일했는가"를 다룬다. 결과물이 아니라 4주의 작업 방식에 대한 기록이다. 데이터 진단(40% 쏠림, 26.3%)과 솔루션 설계(관제탑 모델), 결선작 분석은 웹 페이지에 정리되어 있으므로, 여기서는 그 결론에 이르기까지의 판단과, 결과물을 실제로 찍어낸 과정에 무게를 둔다.

본문은 시작 전에 정한 골격(01), 데이터를 다룬 방식(02), 제안서를 편집한 과정(03), 결과물을 찍어낸 방법(04) 순으로 이어지며, 마지막 05에서 이 작업이 기획·데이터·운영·도구 활용·품질 관리의 어떤 역량으로 정리되는지를 한 표로 모은다.

시작하기 전에 정한 네 가지

아이디어를 떠올리기 전에 작업의 골격부터 정했다. 네 가지를 먼저 고정하고 들어간 것이 4주를 흔들리지 않게 했다.

- 문제를 두 이해관계자로 나눠 정의했다.

이용객의 페인 포인트(어디로 가야 하는지 아무도 알려주지 않는다)와 운영 주체의 과제(특정 출국장 쏠림과 유휴 용량의 공존)를 따로 적고, 둘이 같은 뿌리에서 나온다는 것을 가설로 세웠다. 이 연결이 이후 모든 결정의 기준이 되었다.

- 방향의 우선순위를 정했다.

개인 맞춤 동선 안내(축 A)와 공항 전체 흐름 제어(축 B) 중, 축 A를 최우선으로 두고 축 B는 축 A가 누적되며 따라오는 효과로 배치했다. 개인이 체감하는 가치가 먼저 서야 시스템 효과도 성립한다고 봤기 때문이다. 브레인스토밍에서 모은 자료는 전부 쓰지 않고 이 축에 맞는 것만 남겼다.

- 생성과 비판을 분리했다.

한 AI로 브레인스토밍한 초기 기획을, 다른 AI에 "냉철한 비판가" 역할을 지정해 검증시키는 구조를 만들었다. 같은 도구 안에서 만들고 검증하면 자기 논리를 옹호하게 되기 때문이다. Google AI Essentials와 Google Prompting Essentials에서 배운 역할 지정·교차 검증을 실제 작업 구조로 옮긴 것이다.

- 제출물의 역할을 미리 갈랐다.

제안서는 문제와 해법을 압축해 전달하는 설득 문서로, 추가 결과물은 데이터와 시각 자료로 주장을 받치는 증거 문서로 정의했다. 같은 내용을 두 곳에 늘어놓는 대신 제안서엔 핵심만, 상세는 결과물로 보내 참조하게 했다.

데이터 — 수집보다 검증

데이터에서 한 일의 핵심은 "모으기"가 아니라 "의심하기"였다. (구체적 수치와 산출 과정은 웹 페이지의 데이터 섹션에 있다.)

■ 수집보다 계획이 먼저였다.

무엇을 모을지 정하기 전에 "이 정보는 어떤 주장의 근거로 쓸 것인가"를 항목별로 적었다. 출국장별 여객 데이터는 쏠림의 증거로, 공공 API 명세는 실현 가능성의 증거로 — 목적을 정한 뒤 수집했기 때문에 버려지는 데이터가 거의 없었다.

■ 수치를 만든 뒤 수치를 의심했다.

하루치 예고 데이터로 세운 논리가 과거 데이터에도 성립하는지 거꾸로 확인하고, 성수기에도 적용 가능한지 점검했다. 분석 도중 "지금 핵심이 아닌 주변 논리로 빠지고 있지 않은지"를 의식적으로 자문하며 경로를 교정했다.

■ 실현 가능성을 데이터 단계에서 끝냈다.

아이디어를 다 만든 뒤 "데이터를 구할 수 없다"가 드러나면 전체가 무너진다. 그래서 승객예고 API가 자동 승인되는지, 호출 한도가 실서비스 수준인지(운영 계정 일 100만 건)를 먼저 확인했다.

■ 큰 근거부터 확정했다.

제안의 뼈대가 되는 큰 근거를 먼저 굳히고 초안으로 넘어갔다. 큰 근거가 흔들리면 문서 전체를 다시 써야 하지만 작은 근거는 유연하게 고칠 수 있기 때문이다. T2 관련 이슈는 "해결 불가·물리적 한계가 크면 후 순위, 해결책이 확실하고 시급하면 우선"으로 분류했다.

■ 모든 수치에 출처 장부를 만들었다.

수치마다 원문 인용·URL·산출 과정을 별도 문서로 기록하고 출처의 신뢰 등급까지 구분했다. 위키 기반 2차 출처에서 온 수치에는 "원 기사를 직접 찾거나 표현을 완화할 것"이라는 주의를 달았다. 분석 보고서 자체에도 "분석 한계" 섹션을 넣어 단일 날짜 데이터라는 점과 미반영 변수(항공사별 출국장 배정 규칙 등)를 명시했다. 한계를 밝히면 추정치는 약해 보이지만 유효 용량이 존재한다는 핵심은 흔들리지 않는다고 판단했다.

제안서 — 버리고 합성한 편집

제안서는 네 버전을 거쳤다. (버전별 변화의 의미는 웹 페이지에 정리되어 있다.)

버전	분량	성격
v1	84문단 · 15,410자	전체 논리를 빠짐없이 담은 초안
v2	37문단 · 11,390자	분량 규정에 맞춰 압축한 판본
v3	52문단	v1의 구조에 v2의 밀도를 입혀 합성한 판본
v4	12,382자	컨셉 전환과 비판 반영을 마친 최종본

작업 방식에서 남길 것은 세 가지다.

■ 양자택일 대신 합성했다.

v1과 v2를 "처음 읽는 사람이라면 어느 쪽에 설득되는가"로 비교했다. v1은 논리가 완결적이나 길고, v2는 읽기 쉬우나 뼈대가 약했다. 둘 중 하나를 고르는 대신 v1의 구조에 v2의 밀도를 입힌 v3을 만들었다.

■ 컨셉의 표현을 바꿨다.

초안에서는 자동차 내비게이션(TMAP) 비유가 첫머리에 있었으나, 정보를 주는 내비게이션과 흐름을 제어하는 관제탑은 본질이 다르다고 판단해 "내비게이션이 아니라 관제탑"을 전면에 세웠다. 다만 독자에게 친숙한 TMAP 비유를 완전히 버리지 않고 도입부의 디딤돌로 남기는 두 단계 구성을 택했다. 민간 기업 상호를 공모전 제안서에 쓰는 것이 문제가 없는지도 이 단계에서 별도로 확인했다.

■ 외부 비판을 전부 수용했다.

주변 사람과 AI에서 받은 세 비판 — 26.3%가 비현실적 가정 위의 수치라는 것, 풍선효과 해법 설명이 추상적이라는 것, 시나리오가 이상적이고 고령층·외국인이 빠졌다는 것 — 을 모두 받아들였다. 보수 추정치를 병기하고, 분배 원리를 구체화하고, 고령 페르소나와 키오스크를 추가했다. 비판 전과 후의 제안서는 다른 문서가 되었다.

결과물 — 4주 안에 혼자 짚어낸 방법

이 문서의 중심이다. 개인이 4주 안에 차트 5종·문서 4종·다이어그램·화면 시안·26장 통합 발표 자료를 만들 수 있었던 것은, 만드는 속도가 아니라 만드는 순서와 거르는 기준을 정해둔 덕분이다.

■ 제작 순서에 우선순위를 매겼다.

주장의 토대인 데이터를 가장 앞에, 그 토대 위의 구조물을 그다음에 두는 순서다. 무엇보다 만들지가 정해져 있어 막히는 구간이 없었다.

순위	결과물	우선한 이유
1	데이터 분석 보고서	모든 주장의 토대
2	시스템 아키텍처	토대 위에 세우는 구조
3	UI · 사용자 시나리오	구조를 사용자 경험으로 번역
4	부록(FAQ · 로드맵)	본체를 보강하는 자료

■ 도구 활용은 두 축이었다 — 변환과 검수.

하나는 머릿속 기획을 도구가 한 번에 알아듣는 지시로 바꾸는 프롬프팅이다. 01에서 교차 검증 구조의 토대가 된 Google AI Essentials·Google Prompting Essentials의 프롬프트 설계 방식을, 여기서는 산출물을 빠르게 뽑아내는 데 적용했다. 이 변환이 빨라서 같은 시간에 더 많은 산출물이 나왔다. 다른 하나는 산출물의 오류를 잡아내는 검수다. 사용자 시나리오 초안에 "면세품을 수령했다가 다시 수령한다"는 현실에서 불가능한 동선이 들어 있었는데, 이런 논리 결함은 도구가 스스로 못 잡는다. 산출물을 그대로 쓰지 않고 한 장씩 읽으며 걸러냈고, 빈약한 부분은 검토를 지시해 보강했다. 생성 속도와 검증이 결합될 때 시너지가 나며, 결과물의 품질을 결정하는 것은 도구가 아니라 활용하는 사람이라는 것이 이 4주의 결론이다.

■ 시각 자료는 될 때까지 반복했다.

차트 5종은 Software&AI 융합전공에서 익힌 Python + 데이터 시각화로 직접 그렸고, 같은 데이터라도 어떻게 보여주느냐에 따라 설득력이 달라진다는 것을 기준으로 여러 번 고쳤다. 아키텍처 다이어그램은 글자 겹침과 가독성 문제로 세 번째 버전에서 확정했고, UI 화면은 영문 시안을 한국어로 다시 만들고 배치를 2×2로 키운 두 번째 버전을 채택했다. "글자가 작다", "흐리다" 수준의 문제도 넘기지 않았다. 읽는 사람이 보는 것은 만든 사람의 의도가 아니라 화면이기 때문이다.

- **분량은 늘리지 않고 줄였다.**

결과물이 많으면 오히려 인상이 흐려진다고 보고, "이 자료가 이용객의 문제와 운영 과제를 푸는 증거가 되는가"를 기준으로 자체 검열했다. 26장 통합 발표 자료를 구성할 때도 같은 기준으로 페이지를 배치했다 — 화면 시안과 설명을 한 페이지에 묶고, 아키텍처는 계층별로 잘라 설명을 붙이는 식이다.

이 작업 방식이 증명하는 것

본문의 작업 과정을 역량 단위로 정리하면 다음과 같다.

역량	본문의 증거
기획	이용객·운영 과제의 동시 정의(01) · 축 A/B 우선순위(01) · 관제탑 컨셉 전환(03)
데이터 분석·활용	가설 주도 수집 계획(02) · 역검증과 한계 명시(02) · 출처 등급 장부(02) · Python 데이터 시각화 차트 5종(04)
프로젝트 운영	제출물 역할 분리(01) · 큰 근거 우선 확정(02) · 결과물 우선순위(04) · 분량 자체 검열(04)
AI·도구 활용	생성과 비판을 분리한 교차 검증(01) · 변환 프롬프팅과 오류 검수의 결합(04)
검증·품질 관리	독자 관점 버전 합성(03) · 외부 비판 수용(03) · 산출물 논리 검수(04) · 시각 자료 반복 개선(04)

마케팅 실무의 단위 업무도 구조가 같다고 본다. 캠페인은 고객의 페인 포인트를 정의하는 데서 시작해, 데이터로 근거를 만들고, 제작물을 만들어 검수하며 집행된다. 이 프로젝트는 그 전체 사이클을 4주 동안 혼자 한 바퀴 돌린 기록이며, 각 단계에서 어떤 기준으로 판단하는지를 보여주는 자료다.