

9 장

종합 실습 문제 답안 샘플

※ 다음 답안 샘플은 2026년 6월 25일부터 7월 7일 기간 동안 클로드 **Fable 5**로 테스트한 것입니다. 사용하는 모델과 환경에 따라 답은 달라질 수 있음을 참고하세요.

답안 활용 시 유의사항: 생성형 AI 결과의 변동성에 관하여

이 장의 세 분석은 변동성의 성격이 서로 다르다. VOC 분류(문항 1~2)는 생성형 AI의 언어 이해에 의존하는 정성 분석이므로 범주의 명칭과 건수가 AI마다 달라지는 것이 정상이다. 반면 SERVQUAL gap(문항 3)은 닫힌 계산이므로 어떤 AI를 쓰는 소수점까지 같아야 하고, Kano 분류(문항 4)는 표준 평가표를 따르는 준닫힌 계산이라 세부 규칙에 따라 소폭만 달라진다. 즉 문항이 뒤로 갈수록 허용되는 수치 변동의 폭이 좁아진다.

채점의 기준은 (1) VOC에서 반복 이슈의 실체를 포착하고 채널 간 동일 문제를 통합했는지, (2) SERVQUAL에서 gap의 크기와 재이용에 대한 영향력을 구분했는지, (3) Kano에서 당연·일원적·매력 품질의 구분을 우선순위 논리로 연결했는지, (4) 단계 2에서 세 분석이 서로를 교차 검증하는 구조를 자신의 판단으로 서술했는지이다. 각 문항의 주황색 상자에 허용되는 차이와 오답 기준을 명시하였다.

■단계 1. 실습 파일을 생성형 AI 에 업로드하여 질의 수행

[문제 1에 대한 답안 샘플]

120건을 전수 검토하면 표현은 다르지만 소수의 문제 유형이 반복된다. 네 축의 분류 체계와 대표 분포는 다음과 같다.

분류 축	범주 (건수)
고객 여정 단계	예약·일정 관리(약 40건) / 서비스 수행(약 42건) / 문제 해결·사후 처리(약 34건) / 전반·기타(약 4건)
이슈 유형	일정 변경·취소 불편(27) / 처리 진행상황 불투명(21) / 청소 품질 사각지대·편차(19) / 상담 이력 미연계(13) / 도착 지연·알림 부재(13) / 신뢰·안전(10) / 완료 확인 체계 부재(10) / 세제 안전 요청(2) / 칭찬(5)
감정 톤	불만(약 85건) / 우려·불안(약 16건) / 개선 제안(약 14건) / 칭찬(5건)
긴급도	높음 — "다음 예약 전에 개선됐으면 합니다"가 명시된 22건과 환불·보상 무응답 건 / 중간 — 일반 불만 / 낮음 — 제안·칭찬

〈표 1〉 VOC 4축 분류 체계와 분포 (120건)

[문제 2에 대한 답안 샘플]

순위	핵심 이슈	건수	대표 표현
1	일정 변경·취소의 셀프서비스 불편	27건	"변경 버튼을 찾기 어렵다", "상담원에게 다시 문의해야 한다", "대체 시간이 바로 보이지 않는다"
2	불만·보상 처리 진행상황 불투명	21건	"보상 접수 후 진행 상황이 보이지 않는다", "환불이 되는지 아무 안내가 없다", "재방문 가능 여부를 언제 알려주는지 모른다"
3	청소 품질의 사각지대와 편차	19건	"배수구 주변은 덜 됐다", "싱크대 안쪽·가스레인지는 그대로", "품질이 매번 다르다"
4(공동)	상담 이력 미연계·반복 설명	13건	"채팅 내용을 전화에서 처음부터 다시", "매번 다른 안내", "예약 번

			호를 반복 설명"
4 (공동)	도착 지연과 사전 알림 부재	13건	"25분 늦었는데 미리 알려주지 않았다", "수정 도착 시간이 안 보인다"

〈표 2〉 핵심 이슈 상위 5개 (그 외: 신뢰·안전 10건, 완료 확인 부재 10건, 세제 안전 2건, 칭찬 5건)

채널 간 통합의 예로, "배수구 주변 청소가 덜 됐다"는 동일한 문제가 취소·환불 사유(4건), 서비스 후기(4건), 앱 리뷰(2건), 콜센터 요약(1건)에 나뉘어 접수되었고, "예약 변경 버튼을 찾기 어렵다"는 5개 채널 전부에서 나타난다. 채널별로 따로 보면 각각 소수 의견처럼 보이지만, 문장의 의미 단위로 통합하면 상위 5개 이슈가 전체 VOC의 약 78%(93건)를 차지하는 집중 구조가 드러난다. 또한 표면적으로 다른 불만들, 즉 도착 알림 부재(운영 중), 진행상황 불투명(사후 처리), 완료 확인 부재(수행 직후)는 모두 "회사가 지금 무엇을 하고 있는지 고객이 볼 수 없다"는 정보 투명성이라는 하나의 상위 원인으로 수렴한다.

결과의 변동 가능성 (SI에 따라 다르게 나올 수 있는 부분)

VOC 분류는 생성형 AI의 언어 이해에 의존하므로 범주의 명칭, 개수, 건수가 SI마다 달라지는 것이 정상이다. 특히 경계 사례의 귀속에 따라 건수가 수 건씩 이동한다: "완료 후 사진·체크리스트를 원한다"를 품질 이슈로 볼지 확인 체계 요청으로 볼지, "재청소 일정 확정이 오래 걸린다"를 처리 불투명으로 볼지 응답 속도로 볼지 등이다. 채점은 건수의 일치가 아니라 (1) 일정 변경 불편, 처리 불투명, 품질 사각지대, 도착 알림, 상담 이력이라는 실체가 상위권에 포착됐는지, (2) 채널 간 동일 문제의 통합 사례를 제시했는지, (3) 칭찬(5건)을 불만과 분리했는지를 본다.

오답 기준: 채널별 건수 분포(서비스 후기 34건, 채팅상담 34건 등)를 이슈 분석으로 대체하거나, 칭찬을 불만 건수에 합산하거나, 분류 없이 개별 VOC를 나열만 한 답안은 '분류를 통한 구조화'라는 문항의 요구를 충족하지 못한 것으로 처리한다.

[문제 3에 대한 답안 샘플]

gap을 [경험 - 기대]로 정의하면(음수일수록 기대에 미달) 8개 차원 모두 음의 gap을 보이며, 크기 순서는 다음과 같다.

품질 차원	기대	경험	gap	gap과 재이용 의도의 상관
문제 처리 투명성	6.22	4.62	-1.60	0.18
일정 신뢰성	6.40	4.99	-1.41	0.42
청소 품질	6.30	5.12	-1.18	0.45
응답 속도	6.02	4.90	-1.12	0.03
예약 편의성	6.04	4.96	-1.08	0.13
공감성	5.89	5.33	-0.56	0.01
확신성	6.06	5.51	-0.55	-0.12
물적 요소	5.82	5.29	-0.53	0.05

〈표 3〉 SERVQUAL 기대-경험 gap ($n = 420$, 전반 만족 평균 4.79, 재이용 의도 평균 4.85)

gap 상위 3개는 문제 처리 투명성(-1.60), 일정 신뢰성(-1.41), 청소 품질(-1.18)이다. 이는 VOC 상위 이슈(처리 불투명 21건, 도착 지연 13건, 품질 사각지대 19건)와 정확히 대응하여, 정형·비정형 데이터가 같은 실패 지점을 가리킴을 교차 확인해 준다. 특히 문제를 접수한 경험이 있는 고객(issue_reported=1)의 투명성 gap은 -1.91로 무경험 고객(-1.37)보다 훨씬 커, 서비스 실패 후 회복 과정에서 투명성 실패가 이중으로 누적됨을 보여준다.

다만 gap의 크기와 재이용에 대한 영향력은 구분해야 한다. 8개 gap으로 재이용 의도를 회귀하면($R^2 = 0.38$) 표준화 계수는 청소 품질(0.40)과 일정 신뢰성(0.36)이 압도적이고, 예약 편의성(0.13)과 투명성(0.13)이 그다음이며 나머지는 사실상 0이다. 즉 투명성은 기대 미달의 폭이 가장 크지만, 재예약을 직접 좌우하는 것은 청소 품질과 약속 시간 준수다. 따라서 재예약을 회복의 관점에서는 "품질·일정을 핵심으로 개선하되, 실패가 발생했을 때의 투명성으로 이탈을 방어한다"는 이원 구조로 읽는 것이 타당하다.

결과의 변동 가능성 (AI에 따라 다르게 나올 수 있는 부분)

gap의 수치는 닫힌 계산이므로 어떤 AI를 쓰든 <표 3> 과 일치해야 하며, 다르면 데이터 처리 오류다. 허용되는 차이는 정의의 방향뿐이다: gap을 [기대 - 경험]으로 정의하면 부호가 반대(+1.60 등)가 되며, "양수일수록 기대 미달"로 일관되게 해석했다면 정답이다. 만족도·재이용과의 연결 방법도 상관계수, 회귀, 단순 순위 비교 등 다양하며 모두 인정한다.

이 문항에서 답안의 수준을 가르는 것은 'gap이 큰 차원'과 '재이용에 영향이 큰 차원'이 다르다는 사실의 발견이다(투명성은 gap 1위지만 영향력은 4위, 청소 품질은 gap 3위지만 영향력 1위). 상위 3개 gap만 나열한 답안도 문항의 요구는 충족하지만, 이 불일치를 지적하고 우선순위 함의를 끌어낸 답안은 우수 답안으로 평가한다.

오답 기준: gap의 부호 방향과 해석이 어긋난 답안(예: 경험-기대가 음수인데 "기대를 초과했다"로 해석)은 SERVQUAL의 개념을 오해한 것이므로 틀린 것으로 처리한다.

[문제 4에 대한 답안 샘플]

각 응답자의 긍정형(기능 제공 시)·부정형(미제공 시) 응답 쌍을 표준 Kano 평가표에 대입해 개선안별로 집계하면 다음과 같다. Better 계수 = $(A+O)/(전체-R-Q)$ 는 제공 시 만족 증가의 정도, Worse 계수 = $-(O+M)/(전체-R-Q)$ 는 미제공 시 불만의 정도다.

개선안	M	O	A	I	최빈 유형 (비율)	Better	Worse
K01 방문 전 도착 알림 + 지연 자동 재안내	156	43	14	23	당연 품질 M (65.0%)	0.24	-0.84
K02 완료 체크리스트·주요 공간 사진	46	126	44	15	일원적 품질 O (52.5%)	0.74	-0.74
K03 원클릭 일정 변경/취소 + 대체 시간 추천	50	133	37	13	일원적 품질 O (55.4%)	0.73	-0.79
K04 불만 접수 후 진행상황 표시	45	124	37	24	일원적 품질 O (51.7%)	0.70	-0.73

K05 선호 매니저 재예약 기능	11	38	137	45	매력 품질 A (57.1%)	0.76	-0.21
K06 친환경 세제·반려동물 안전 옵션	21	33	122	51	매력 품질 A (50.8%)	0.68	-0.24

〈표 4〉 Kano 분류 결과 ($n = 240$, R·Q 응답 개선안별 4~13건은 표에서 생략)

해석.

K01(도착 알림)은 당연 품질이다. 제공돼도 만족은 크게 오르지 않지만(Better 0.24) 없으면 강한 불만(Worse -0.84)을 낳는, '있어야 기본'인 기능이다. VOC의 도착 지연 불만 13건이 이를 뒷받침한다. K02·K03·K04는 일원적 품질로, 잘할수록 만족이 오르고 못하면 불만이 커지는 성과 비례형이다. 셋 모두 Better 0.70대·Worse -0.7대로 개선 투자의 만족 환수가 크다. K05(선호 매니저)·K06(친환경 세제)은 매력 품질로, 없어도 불만은 작지만(Worse -0.2대) 제공하면 만족과 차별화를 만드는 기능이다.

따라서 자원 투입의 우선순위는 [당연 품질의 결핍 해소 → 일원적 품질의 강화 → 매력 품질의 선택적 도입] 순서, 즉 K01을 최우선으로 하고 K03·K04·K02를 그다음에, K05·K06을 차별화 단계에 배치하는 것이 Kano 논리에 부합한다.

결과의 변동 가능성 (AI에 따라 다르게 나올 수 있는 부분)

Kano 분류는 준단한 계산이다. 표준 평가표를 쓰면 응답자별 분류가 결정되므로 〈표 4〉의 빈도는 재현되어야 하지만, 다음 세부 규칙에 따라 소폭 달라질 수 있다: (1) 평가표의 판정 셀 일부(특히 역(reverse)·의문(questionable) 처리)는 문헌에 따라 다르고, (2) 개선안의 대표 유형을 최빈값으로 정할지, $M > O > A > I$ 우선순위 규칙이나 Better-Worse 좌표로 정할지에 따라 근소한 차이가 난다. 이 데이터는 최빈 비율이 50~65%로 명확해 어느 규칙을 써도 K01=당연, K02~K04=일원적, K05~K06=매력이라는 결론이 안정적이므로, 이 정성적 분류에 도달했다면 빈도의 소수점 차이와 무관하게 정답이다. Better-Worse 계수나 산점도를 함께 제시한 답안은 우수 답안이다.

오답 기준: 긍정형 응답 하나만으로(예: '마음에 든다' 비율 순위) 유형을 정한 답안은 긍정·부정 문항 쌍의 조합이라는 Kano 모형의 핵심 구조를 위반한 것이므로 틀린 것으로 처리한다. 이 방식으로는 당연 품질(K01)과 매력 품질(K05)을 구분할 수 없기 때문이다.

■ 단계 2. A4 한 페이지 이내로 제출물 작성

[답안 샘플]

진단:

재예약률 하락은 '한 번의 큰 실패'가 아니라 '보이지 않는 서비스'의 누적이다. 세 데이터는 독립적으로 수집되었지만 같은 지점을 가리킨다. VOC 120건의 상위 이슈(일정 변경 불편 27건, 처리 불투명 21건, 품질 사각지대 19건, 상담 이력 단절·도착 지연 각 13건)는 SERVQUAL의 gap 상위 3개 차원(투명성 -1.60, 일정 신뢰성 -1.41, 청소 품질 -1.18)과 일대일로 대응하고, Kano 설문에서 고객이 기본으로 요구하는 기능(K01 도착 알림)과 성과 비례로 반응하는 기능(K03 일정 변경, K04 진행상황, K02 완료 확인)도 정확히 같은 문제의 해법이다. 이 수렴은 진단의 신뢰성을 교차 검증한다.

판단:

우선순위는 세 근거의 교집합에서 나온다. 첫째, K01(도착 알림)은 Kano상 당연 품질(Worse -0.84 최대)이면서 SERVQUAL 재이용 영향력 2위 차원(일정 신뢰성)의 직접 해법이므로 최우선이다. 없애야 할 불만이지 더할 매력 이 아니다. 둘째, K03(원클릭 일정 변경)은 VOC 최다 이슈(27건)의 해법이자 일원적 품질(Better 0.73)이므로 두 번째다. 셋째, K04(진행상황 표시)는 gap이 가장 큰 투명성의 해법이며, 문제를 겪은 고객에서 gap이 -1.91로 깊어지는 이탈 방어 장치다. 넷째, K02(완료 체크리스트·사진)는 재이용 영향력 1위인 청소 품질 인식을 보완한다. 품질 자체의 편차 관리(매니저 교육·표준 절차)와 병행해야 효과가 난다. 반면 K05·K06(선호 매니저, 친환경 세제)은 매력 품질이므로, 위 네 과제로 기본기를 회복한 뒤 차별화 단계에서 도입한다.

한 문장 결론:

"고객은 청소가 잘 됐는지, 매니저가 언제 오는지, 접수한 불만이 어떻게 처리되는지 아무것도 볼 수 없다고 말하고 있다. 도착 알림(K01)과 일정 변경(K03), 진행상황 표시(K04), 완료 확인(K02)으로 서비스의 전 과정을 '보이게' 만드는 것이 재예약률 회복의 최우선 과제이며, 선호 매니저·친환경 옵션은 그다음 차별화 단계의 과제다."

결과의 변동 가능성 (SI에 따라 다르게 나올 수 있는 부분)

단계 2는 본인의 판단을 요구하는 정성 문항이므로 정답이 하나가 아니다. 예컨대 재이용 영향력 1위인 청소 품질을 근거로 K02와 품질 표준화를 최우선에 둔 답안, 또는 VOC 최다 건수를 근거로 K03을 최우선에 둔 답안도 각자의 근거가

데이터와 일치하면 정답이다. 채점은 (1) 세 분석의 결과가 상호 참조되는지(하나라도 누락 시 감점), (2) 우선순위의 근거가 자신의 분석 수치와 일치하는지, (3) AI 응답의 단순 요약이 아니라 트레이드오프(예: gap 크기 vs 재이용 영향력, 당연 품질 vs 매력 품질)에 대한 본인의 판단이 드러나는지를 본다. 분량(A4 1페이지 이내) 준수도 평가 대상이다.