

A DETAILED PLATFORM DESCRIPTIONS

We collected data from nine Korean online platforms representing diverse user communities and domain expertise. Table 1 provides detailed information about each platform. These platforms were selected to ensure comprehensive coverage of different user demographics, expertise levels, and domain-specific knowledge, reflecting the diversity of real-world multimodal questions Korean users encounter online.

Table 1: Korean online platforms used for data collection

Platform	Category	Description
Naver KnowledgeIn	General Q&A	Korea’s largest general Q&A platform covering everyday queries, academic subjects, and technical issues
BRIC	Science Community	Specialized community for biological research and biotechnology with scientific discussions and professional knowledge sharing
Ruliweb	Gaming Community	Major gaming community covering video games, hardware reviews, game mechanics, and technical gaming issues
MonsterZym	Fitness Community	Fitness and bodybuilding community discussing workout routines, nutrition, supplements, and exercise techniques
Quasarzone	Hardware Community	Hardware enthusiast community focused on computer components, electronics, PC building, and technology reviews
i-Boss	Business Platform	Business and entrepreneurship platform for startup strategies, operations, marketing, and professional development
Inflearn	Coding Education	Online learning platform with community features for programming questions and coding experiences
Codeit	Coding Education	Coding education platform with forums for programming discussions and technical support
Okky	Developer Community	Developer community platform for programming discussions, career advice, and technical problem-solving

B STAGE 2 PROMPT EXCERPTS

We used three LLM-based filters in Stage 2: content safety, objectivity, and temporal dependency. Below we excerpt only the core exclusion criteria from the prompts (full wording omitted).

B.1 CONTENT SAFETY

Mark as inappropriate if the question–image pair includes:

- Political content (politicians, parties, elections, political opinions)
- Religious advocacy/criticism or conflicts
- Hate/discrimination
- Suicide or self-harm; sensitive mental-health topics
- Sexual/adult content, nudity, explicit innuendo

B.2 OBJECTIVITY

Mark as inappropriate if the pair is subjective or ambiguous, e.g.:

- Preference/aesthetic judgments (“pretty/ugly”, “which outfit is nicer?”)
- Suitability/personal advice without criteria
- Moral/intentionality speculation (“who is wrong?”, “good person?”)
- Multiple valid interpretations or unverifiable answers

Game (Stardew Valley) <i>“What is the circled item in the screenshot?”</i> • Identify circled item as a sap tap (수액채취기) • Mention install only on fully grown trees • Explain how to obtain/craft it • Note sap can be collected after time	Economics/Management <i>“Cost allocation: is S2 missing 100,000?”</i> • Provide correct S1/S2 values • Reset self-allocation entries to zero • Derive allocation ratios (0.5F, 0.4M)
Daily Life <i>“Is this ceiling tile asbestos?”</i> • Identify material as gypsum, not asbestos • Explain gypsum board contains no asbestos • Explicitly name “석고 텍스” • Assure user it is safe	Science <i>“Why does neutron mass ratio decrease?”</i> • Explain neutron beta decay • Clarify neutrons inside He nucleus • Relate x-axis to cosmic cooling • Interpret H:He ratio $\approx 3:1$

Figure 1: Examples of checklist decomposition across domains, generated in Stage 5. For brevity, the checklists shown here are abbreviated; full checklists typically contain 1–5 criteria per item.

B.3 TEMPORAL DEPENDENCY

Mark as inappropriate if the pair requires time-specific information, e.g.:

- “today/now” weather, traffic, store hours, last train
- Current events or status queries (“is it open now?”, “stock price today?”)
- Questions that become invalid/meaningless as time passes

STAGE 4 PROMPT EXCERPT (IMAGE DEPENDENCY RUBRIC)

Input: (Q), model answer with image, model answer without image, optional gold answer snippet.

Task: Compare the two answers and decide image dependency.

Decision labels

- **IMAGE_REQUIRED:** with-image answer is substantially more accurate/specific; text-only answer is vague, incorrect, or explicitly requests the image.
- **NO_IMAGE_NEEDED:** both answers are comparable in correctness and specificity without relying on visual cues.
- **UNCERTAIN:** evidence is inconclusive (e.g., partial improvements or conflicting signals).

Scoring (1–5 quality gap)

- 1: negligible difference; 3: clear but moderate gain; 5: decisive gain (critical visual details).

Output (natural language)

- **Judgment:** IMAGE_REQUIRED / NO_IMAGE_NEEDED / UNCERTAIN
- **Reason:** brief comparison citing concrete differences
- **QualityGap:** integer in {1,2,3,4,5}

C STAGE 5 PROMPTS AND CHECKLIST EXAMPLES

This appendix provides the instruction prompt used for checklist generation along with illustrative examples of the resulting decompositions. We used GPT-4-mini to derive structured criteria directly from reference answers that users found satisfactory. These checklists therefore represent strict, human-aligned evaluation standards: a model must satisfy all listed criteria to be considered correct.

D PLATFORM-WISE FILTERING STATISTICS

Table 2 provides a detailed breakdown of data collection and filtering across all platforms.

Table 2: Detailed data collection and filtering statistics by platform (Stages 1–6). Coding platforms include Inflearn, Codeit, and Okky combined.

Platform	Raw Data	Appropri.	Difficulty	Image Dep.	Human Val.	Final	Survival
KnowledgeIn	31,484	10,495	1,404	648	441	441	1.4%
BRIC	291	291	163	60	42	42	14.4%
Ruliweb	305	240	54	42	32	32	10.5%
Coding	27,896	8,369	837	198	135	135	0.5%
MonsterZym	3,090	3,090	2,234	8	6	6	0.2%
Quasarzone	2,986	896	90	22	15	15	0.5%
i-Boss	20,000	20,000	578	62	42	42	0.2%
Total	86,052	43,381	5,360	1,040	713	653	0.76%

E INVESTIGATING FAILURE MODES

In Table 3, we observe that VARCO-VISION and HyperCLOVA X—two Korean-focused VLMs—underperform multilingual counterparts of similar scale. While the precise reasons remain unclear due to the closed nature of these models and limited information about their training, we propose two possible explanations:

- (A) **Training Data Coverage.** Current benchmarks that capture progress on culturally grounded, information-deficient queries are scarce. Model developers may not have explicitly emphasized such aspects in their training data, leading to weaker performance on this type of evaluation.
- (B) **Pretraining Scale and Robustness.** Robustness to imperfect or fragmented user queries may emerge from exposure to large-scale, diverse pretraining corpora. Larger multilingual models are more likely to encounter noisy, colloquial, or partially specified inputs, thereby preparing them better for benchmarks of this kind.

Table 3: Complete performance across all 13 categories for all evaluated models (scores in %). All scores are reported as mean_{SE}, where SE is the standard error over 3 independent runs (n=3).

Model	IT	Health	Game	Econ	Sci	Mach	Daily	Shop	Math	Ent	Trans	Nature	Code	Avg
<i>Proprietary Models</i>														
Gemini 2.5 Pro	50.73 _{1.63}	62.17 _{2.33}	36.67 _{1.75}	51.09 _{1.72}	39.93 _{1.43}	56.22 _{4.32}	51.32 _{2.57}	46.00 _{5.15}	60.94 _{1.36}	44.37 _{1.18}	57.69 _{2.81}	53.45 _{0.57}	50.91 _{0.92}	48.54 _{0.18}
Gemini 2.5 Flash	42.98 _{0.34}	56.10 _{3.95}	26.70 _{1.04}	48.05 _{6.64}	39.62 _{3.14}	45.86 _{1.70}	44.59 _{2.10}	46.13 _{5.19}	51.14 _{3.87}	31.92 _{2.63}	48.12 _{2.37}	44.37 _{0.99}	39.26 _{2.25}	41.05 _{1.38}
Gemini 2.5 Flash Lite	25.92 _{1.82}	43.54 _{4.48}	17.97 _{1.92}	38.84 _{3.82}	41.73 _{1.47}	38.30 _{3.10}	30.67 _{0.88}	28.82 _{2.38}	45.62 _{7.49}	18.82 _{0.68}	34.10 _{3.78}	27.16 _{0.32}	32.63 _{2.66}	30.29 _{0.42}
GPT 5	59.95 _{2.01}	62.61 _{2.59}	32.34 _{2.08}	58.41 _{1.63}	36.31 _{1.60}	52.85 _{4.72}	46.93 _{2.83}	55.96 _{3.14}	54.70 _{4.54}	33.80 _{2.20}	54.97 _{2.43}	53.42 _{1.23}	55.07 _{1.24}	48.01 _{0.32}
GPT 5 Mini	49.59 _{3.74}	60.45 _{2.40}	29.22 _{1.71}	50.19 _{5.53}	52.49 _{0.44}	51.68 _{1.47}	50.28 _{4.75}	44.33 _{4.96}	58.19 _{3.94}	25.54 _{2.20}	49.22 _{3.11}	41.17 _{2.73}	57.02 _{0.53}	45.21 _{1.21}
GPT 5 Nano	22.99 _{2.64}	45.98 _{1.02}	10.46 _{1.71}	24.81 _{0.65}	11.47 _{1.21}	26.59 _{7.45}	21.49 _{1.41}	26.42 _{3.26}	23.56 _{6.12}	12.81 _{0.67}	26.17 _{1.63}	25.27 _{1.60}	32.84 _{4.80}	21.22 _{0.46}
Grok 4	39.64 _{1.89}	36.96 _{1.16}	19.00 _{0.49}	44.44 _{2.79}	40.70 _{1.13}	47.63 _{1.86}	40.57 _{1.60}	36.73 _{1.65}	22.09 _{2.86}	24.77 _{1.78}	50.43 _{4.90}	30.29 _{1.28}	39.02 _{0.20}	36.08 _{0.53}
<i>Open-source Models</i>														
<i>Mistral/Pixtral Family</i>														
Mistral Medium 3.1	24.77 _{2.76}	37.01 _{5.96}	16.01 _{1.49}	28.48 _{2.48}	33.70 _{1.23}	34.14 _{1.20}	24.41 _{1.06}	25.22 _{2.51}	38.99 _{3.41}	11.46 _{2.32}	25.46 _{2.65}	19.62 _{2.62}	31.09 _{2.35}	24.86 _{0.98}
Pixtral Large	19.09 _{1.82}	35.09 _{2.33}	11.33 _{1.33}	24.32 _{3.36}	27.40 _{2.29}	23.41 _{1.21}	24.16 _{1.09}	19.01 _{4.66}	19.89 _{1.10}	11.54 _{2.50}	21.93 _{1.34}	18.08 _{0.62}	22.64 _{0.67}	20.10 _{0.41}
Mistral Small 24B	15.38 _{1.93}	25.07 _{4.25}	7.00 _{1.35}	22.29 _{1.84}	20.47 _{1.46}	21.53 _{2.01}	13.07 _{2.67}	15.34 _{3.08}	18.57 _{1.81}	7.76 _{1.09}	13.84 _{2.48}	10.61 _{1.77}	16.36 _{2.46}	14.43 _{0.41}
Pixtral 12B	8.76 _{0.77}	24.19 _{3.58}	6.74 _{0.94}	17.49 _{0.53}	14.12 _{0.11}	16.46 _{2.65}	11.66 _{2.40}	11.44 _{1.34}	6.83 _{2.27}	6.17 _{0.35}	15.06 _{2.39}	9.60 _{0.46}	12.94 _{2.80}	11.20 _{0.02}
<i>Google Gemma Family</i>														
Gemma 3 27B	20.31 _{1.18}	40.90 _{1.49}	13.75 _{1.55}	31.71 _{3.21}	34.93 _{1.52}	27.36 _{6.28}	26.72 _{1.12}	24.07 _{2.01}	23.85 _{2.74}	9.43 _{1.30}	20.66 _{2.62}	18.61 _{0.40}	20.81 _{2.15}	22.53 _{0.28}
Gemma 3 12B	15.15 _{0.69}	36.60 _{1.32}	10.52 _{1.44}	27.91 _{1.30}	28.79 _{1.39}	27.44 _{3.60}	19.20 _{1.27}	22.40 _{1.47}	17.25 _{2.89}	7.23 _{1.12}	21.01 _{1.65}	13.43 _{1.47}	23.41 _{0.13}	18.76 _{0.63}
Gemma 3 4B	12.43 _{1.63}	34.23 _{1.08}	8.91 _{0.96}	19.67 _{4.37}	22.50 _{0.12}	21.25 _{1.33}	15.59 _{0.87}	18.21 _{1.21}	13.54 _{2.63}	6.84 _{1.10}	19.56 _{2.12}	14.68 _{1.08}	13.45 _{0.88}	15.47 _{0.78}
<i>AIDC-AI Ovis2 Family</i>														
Ovis2-34B	15.90 _{1.35}	40.15 _{2.16}	9.87 _{0.77}	19.44 _{0.45}	23.97 _{0.56}	29.46 _{1.47}	19.43 _{0.58}	20.27 _{3.31}	22.91 _{2.46}	9.18 _{1.41}	21.86 _{2.89}	18.77 _{1.37}	16.78 _{0.26}	18.50 _{0.03}
Ovis2-16B	11.20 _{1.67}	38.98 _{0.75}	8.08 _{0.18}	21.58 _{1.27}	24.68 _{0.80}	23.94 _{3.50}	21.20 _{3.52}	14.83 _{3.00}	24.32 _{1.31}	8.72 _{1.57}	20.21 _{0.84}	16.47 _{0.63}	16.12 _{1.92}	17.18 _{0.50}
Ovis2-8B	9.80 _{1.30}	33.62 _{1.54}	6.07 _{0.30}	19.18 _{3.28}	19.45 _{1.85}	21.02 _{1.98}	18.37 _{1.83}	13.51 _{1.33}	19.81 _{5.29}	8.04 _{0.53}	17.42 _{3.08}	13.17 _{0.35}	14.77 _{1.80}	14.46 _{0.37}
Ovis2-4B	6.76 _{1.75}	23.66 _{3.93}	6.00 _{0.27}	15.89 _{2.76}	16.16 _{1.17}	17.05 _{3.15}	16.43 _{1.51}	10.68 _{2.89}	13.16 _{0.84}	7.17 _{0.50}	17.65 _{3.01}	14.26 _{0.58}	8.31 _{1.00}	12.18 _{0.11}
Ovis2-2B	6.14 _{0.22}	16.10 _{1.01}	5.30 _{0.83}	13.74 _{2.34}	12.24 _{1.70}	13.64 _{4.43}	11.99 _{1.14}	11.27 _{2.01}	6.57 _{1.32}	7.28 _{0.64}	11.33 _{2.19}	9.73 _{0.56}	8.98 _{3.88}	9.54 _{0.22}
Ovis2-1B	4.83 _{0.91}	12.62 _{2.58}	4.74 _{0.31}	8.07 _{1.07}	7.52 _{0.71}	5.95 _{1.12}	8.03 _{0.98}	8.11 _{1.97}	6.57 _{2.40}	5.05 _{0.98}	8.10 _{2.55}	6.80 _{1.38}	4.43 _{1.13}	6.52 _{0.25}
<i>OpenGVLab InternVL3.5 Family</i>														
InternVL3.5 38B	14.94 _{0.63}	30.95 _{4.82}	9.09 _{1.57}	24.85 _{1.52}	28.79 _{0.27}	20.90 _{4.44}	19.25 _{1.90}	18.40 _{0.19}	24.54 _{2.47}	8.53 _{0.17}	21.10 _{0.84}	16.41 _{1.98}	14.76 _{2.12}	18.01 _{0.39}
InternVL3.5 14B	15.50 _{2.05}	26.81 _{4.46}	8.26 _{1.12}	20.72 _{0.96}	24.64 _{1.18}	17.41 _{3.95}	14.67 _{1.98}	17.70 _{3.09}	26.45 _{1.63}	7.74 _{0.53}	15.76 _{1.58}	12.05 _{1.07}	19.72 _{3.13}	16.04 _{0.37}
InternVL3.5 8B	10.22 _{1.08}	23.11 _{3.12}	7.14 _{0.57}	20.44 _{1.87}	20.14 _{1.89}	16.16 _{3.35}	11.27 _{1.56}	11.99 _{2.92}	22.96 _{2.08}	5.29 _{0.79}	12.68 _{1.73}	12.57 _{0.25}	13.01 _{1.22}	13.16 _{0.82}
InternVL3.5 4B	7.70 _{1.15}	23.33 _{0.30}	7.72 _{0.52}	19.71 _{1.60}	23.20 _{1.24}	18.84 _{0.42}	15.11 _{1.52}	14.98 _{2.03}	25.72 _{2.64}	6.48 _{0.96}	13.83 _{1.36}	11.78 _{1.37}	14.90 _{2.25}	14.09 _{0.28}
InternVL3.5 2B	5.32 _{0.25}	20.86 _{4.13}	5.24 _{0.34}	15.50 _{1.86}	16.05 _{0.87}	12.69 _{2.87}	8.94 _{1.54}	7.69 _{1.60}	14.18 _{1.71}	5.63 _{1.26}	10.14 _{3.14}	7.03 _{0.51}	9.07 _{2.28}	9.48 _{0.49}
InternVL3.5 1B	3.21 _{0.43}	7.94 _{2.99}	3.39 _{0.09}	10.32 _{0.07}	9.12 _{0.32}	5.74 _{0.58}	3.29 _{1.02}	7.79 _{1.30}	10.22 _{1.53}	3.24 _{0.57}	7.31 _{1.05}	2.93 _{0.74}	5.64 _{0.43}	5.43 _{0.13}
<i>Qwen Family</i>														
Qwen2.5 VL 72B	16.53 _{1.36}	31.30 _{1.38}	11.80 _{2.24}	25.55 _{1.06}	28.46 _{2.62}	23.55 _{1.42}	19.72 _{0.38}	25.86 _{3.14}	32.32 _{7.22}	9.97 _{0.45}	21.02 _{2.62}	19.36 _{0.79}	25.31 _{1.59}	20.58 _{0.80}
Qwen2.5 VL 7B	10.33 _{0.70}	21.04 _{4.51}	5.95 _{1.26}	18.96 _{1.05}	20.49 _{3.89}	18.50 _{3.79}	13.70 _{0.92}	17.00 _{4.00}	13.26 _{4.07}	6.71 _{0.28}	14.06 _{1.66}	12.35 _{0.74}	13.28 _{2.86}	13.15 _{0.86}
Qwen2.5 VL 3B	6.08 _{2.15}	18.49 _{3.90}	2.82 _{0.44}	12.76 _{1.17}	11.54 _{1.70}	13.76 _{2.51}	9.22 _{0.16}	6.89 _{1.47}	10.14 _{0.98}	4.88 _{0.18}	10.31 _{3.46}	7.85 _{0.38}	6.54 _{0.84}	8.20 _{0.36}
<i>Other Open-source</i>														
Skywork-R1V3-38B	27.12 _{0.74}	47.94 _{2.92}	15.30 _{1.63}	32.37 _{2.44}	36.84 _{0.69}	37.25 _{1.80}	26.43 _{2.63}	28.27 _{1.95}	41.71 _{4.53}	14.76 _{1.96}	30.10 _{2.73}	27.38 _{0.26}	26.42 _{0.26}	27.76 _{0.58}
<i>Korean-specialized Models</i>														
VARCO-VISION-2.0-14B	11.90 _{0.79}	34.76 _{4.78}	7.94 _{0.85}	17.83 _{2.30}	22.03 _{2.71}	23.46 _{3.16}	21.89 _{1.09}	14.05 _{2.80}	12.68 _{1.90}	7.80 _{2.64}	18.84 _{1.75}	14.97 _{0.57}	13.31 _{1.31}	15.55 _{0.50}
HyperCLOVA-3B	8.42 _{0.98}	29.74 _{2.98}	6.33 _{0.49}	15.17 _{1.40}	18.54 _{0.41}	15.80 _{2.14}	13.38 _{0.67}	13.43 _{3.83}	9.86 _{2.44}	6.16 _{0.70}	14.20 _{1.75}	16.21 _{0.93}	9.53 _{1.86}	12.66 _{0.18}
VARCO-VISION-2.0-1.7B	8.09 _{1.21}	21.34 _{1.50}	5.95 _{2.38}	16.07 _{0.96}	17.79 _{0.63}	16.22 _{1.16}	12.70 _{0.32}	12.88 _{1.08}	12.54 _{5.35}	8.11 _{0.68}	12.81 _{1.01}	12.13 _{0.72}	10.46 _{3.57}	11.87 _{0.46}

F ANNOTATION GUIDELINES

Seven Korean-speaking annotators conducted human validation using a custom web-based tool (Figure 2). Annotators were provided with comprehensive guidelines covering five evaluation dimensions:

F.1 CORE EVALUATION CRITERIA

Image-Question Relevance: Assess whether images provide essential visual information required to answer the question. Images should contain specific visual elements that cannot be determined from text alone.

Question-Answer Quality: Evaluate question clarity, answerability, and reference answer accuracy. Questions should be unambiguous with verifiable answers, while reference answers should be comprehensive and correct.

Checklist Validation: Review each checklist item for necessity, clarity, and completeness. Items should capture essential answer components using unambiguous, measurable criteria that together represent the full scope of required understanding.

Category Appropriateness: Verify correct classification into one of 13 domain categories based on primary subject matter and required expertise.

Overall Assessment: Flag items with fundamental issues such as inappropriate content, cultural insensitivity, or unsolvable questions.

F.2 QUALITY STANDARDS

Annotators applied conservative filtering criteria, removing any item flagged as problematic to maintain high dataset quality. Specific attention was given to:

- Cultural specificity and Korean context requirements
- Image dependency verification through visual inspection
- Checklist comprehensiveness and granularity
- Answer completeness relative to community expectations

The annotation process employed automatic progress saving and session management to ensure data integrity and allow work resumption across multiple sessions.

G HUMAN EVALUATION FEEDBACK ANALYSIS

Table 4 presents representative examples of human annotator feedback for inappropriate judge evaluations, revealing systematic failure patterns.

Rating	Human Reasoning (translated)
Very Inappropriate	"Judge awarded points based on superficial word matching rather than actual checklist compliance"
Inappropriate	"Judge gave 1 point despite response not addressing checklist criteria, incorrectly interpreting explicit mention as meeting requirements"
Inappropriate	"Checklists 1,2,4 satisfied. Item 3 not clearly inappropriate but ambiguous and open to interpretation"
Inappropriate	"Even if intent aligns with checklist, response lacks clarity and remains ambiguous"
Inappropriate	"Judge overlooked insufficient explanations that clearly failed checklist requirements"

Table 4: Representative human feedback explaining inappropriate judge ratings.

Analysis reveals judge failures primarily stem from: (1) superficial keyword matching without semantic understanding, (2) excessive leniency toward incomplete responses, and (3) difficulty distinguishing between implicit intent and explicit satisfaction of requirements.

