



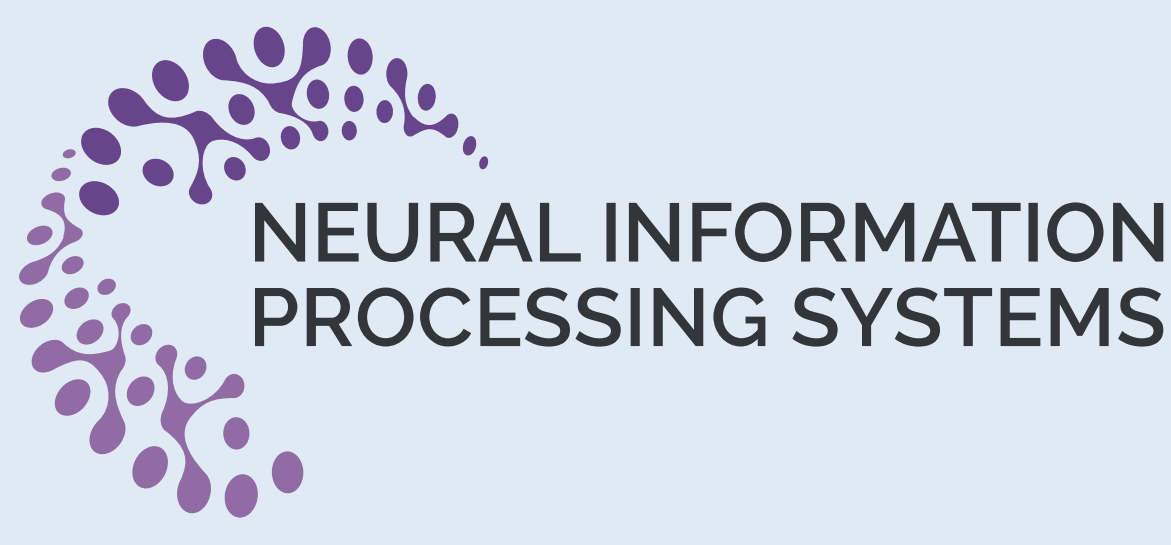
VidEmo: Affective-Tree Reasoning for Emotion-Centric Video Foundation Models

Zhicheng Zhang¹ Weicheng Wang¹ Yongjie Zhu³ Wenyu Qin³ Pengfei Wan³ Di Zhang³ Jufeng Yang^{1,2}

¹Nankai University

²Pengcheng Lab

³Kuaishou Technology

Di Zhang ³Jufeng Yang ^{1,2}

Introduction

- **High-Level emotion intelligence toward AGI requires explainable reasoning:** Human emotion in video is dynamic, open-set, and heavily context-dependent. Achieving AGI-level emotion understanding therefore requires models that not only perceive low-level cues but also construct *transparent, step-wise* affective rationales. However, even advanced VideoLLMs such as Gemini 2.0 reach only 26.3% accuracy on fine-grained sentiment analysis, underscoring the substantial gap in high-level emotional intelligence.
- **VidEmo: A Tree-Structured, Affective-Cue-Guided Reasoning Framework:** To close this gap, we propose VidEmo, a hierarchical reasoning model that decomposes emotion understanding into three stages: (1) *attribute perception*, (2) *expression analysis*, and (3) *emotion reasoning*. Inspired by recent progress in reasoning-style models (e.g., R1), VidEmo performs structured, cue-guided reasoning, integrating appearance, micro-expressions, temporal dynamics, and scene context into coherent emotional interpretations.

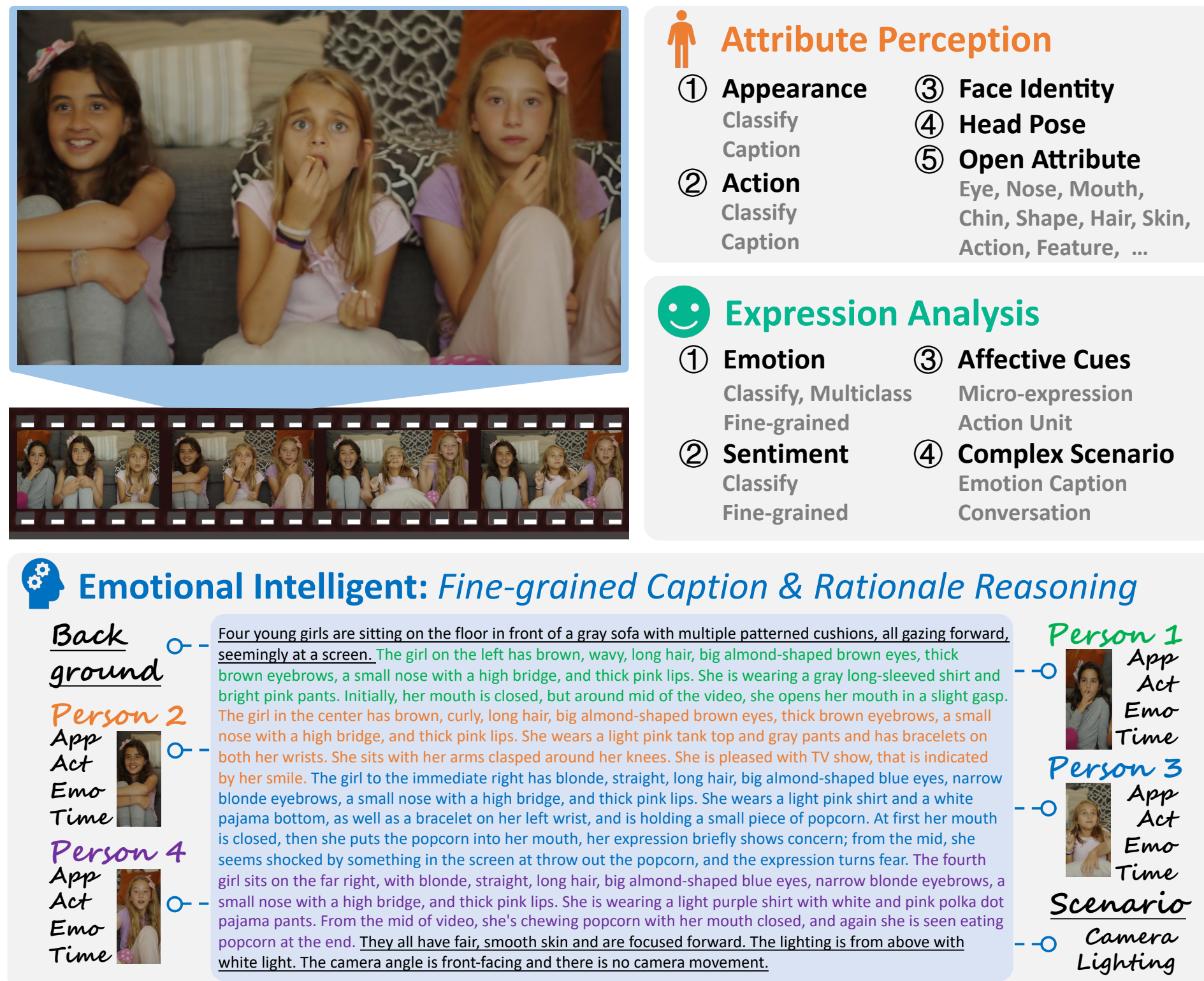


Figure 1. Selected examples of inputs and outputs obtained from . Apart from providing toolkits for basic attribute perception and expression analysis (top), extends the cognitive capacity and is able to generate fine-grained emotional captions with explainable rationale (bottom).

VidEmo: Video Emotion Foundation Models

Pipeline: To develop a family of emotion-centric video foundation models, we propose a comprehensive set of toolkits designed for the pre-training, post-training, and reasoning stages.

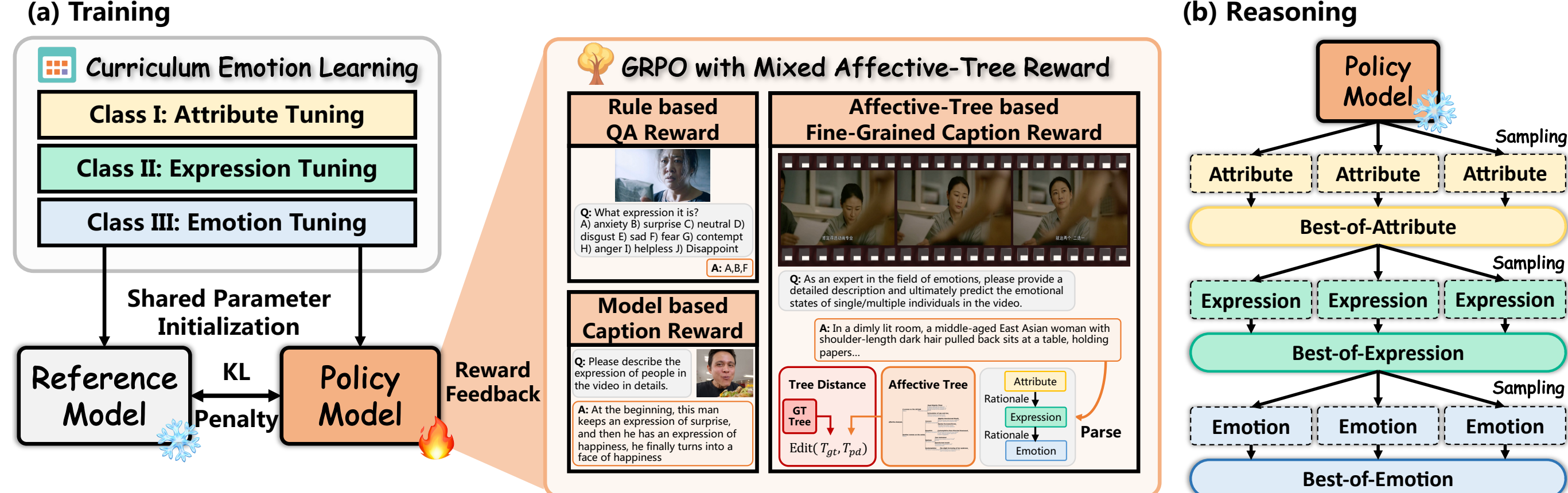


Figure 2. Pipeline of VidEmo. (a) Training: The model is trained using curriculum emotion learning, divided into three stages: attribute, expression, and emotion tuning. A reference model provides initial parameters, and a policy model is trained with reward feedback. (b) Reasoning: The policy model performs hierarchical reasoning by sampling from the best attributes, expressions, and emotions to generate the final emotional output.

Pre-training: Curriculum Emotion Learning

- **Curriculum Emotion Learning:** A progressive 3-stage curriculum that injects knowledge into base model: (1) attribute perception -> (2) expression analysis -> (3) emotion understanding. This reduces perplexity and stabilizes learning across difficulty levels.

Post-training: GRPO via Mixed Affective-Tree Reward

- **Starting from GRPO:** Formally, let q be a query, GRPO samples a group of outputs $\{o_i\}_{i=1}^G$ with the number of G from the old policy model $\pi_{\theta_{\text{old}}}$, and train a policy model by maximizing the following objective:
$$\mathcal{J}_{\text{GRPO}}(\theta) = \mathbb{E}_{q, \{o_i\} \sim \pi_{\theta_{\text{old}}}} \left[\frac{1}{G} \sum_{i=1}^G \frac{1}{|o_i|} \sum_{t=1}^{|o_i|} \min \left(r_t(\theta) \hat{A}_{i,t}, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon) \hat{A}_{i,t} \right) \right] - \beta D_{\text{KL}}(\pi_{\theta} \parallel \pi_{\text{ref}}),$$
- **Rule-based QA reward:** The model is evaluated on its ability to respond to emotion-related queries using predefined rules of Acc and F1 score.
- **Model-based Caption reward:** For the short caption of action, appearance, and emotion, we use a generative reward model to score the quality of captions generated by the model.
- **Affective Tree-based Fine-Grained Caption Reward:** Given a generated caption $\hat{o}$, we first parse it into a set of aspect-item pairs at three semantic levels: attribute ($\mathcal{A}$), expression ($\mathcal{E}$), and emotion ($\mathcal{M}$). These elements are organized into a three-level affective tree T_{pred} , where each node represents an extracted item and directed edges encode rationale:

$$\mathcal{A} \xrightarrow{\text{rationale}} \mathcal{E} \xrightarrow{\text{rationale}} \mathcal{M}. \quad (1)$$

We compare the predicted tree T_{pred} with gt tree T_{gt} using the tree edit distance $\text{Edit}(T_{\text{gt}}, T_{\text{pred}})$, which quantifies the minimal number of edit operations (insertions, deletions, substitutions) required to transform one tree into the other as reward R :

$$R = \exp \left(-\lambda \cdot \text{Edit}(T_{\text{gt}}, T_{\text{pred}}) \right), \quad (2)$$

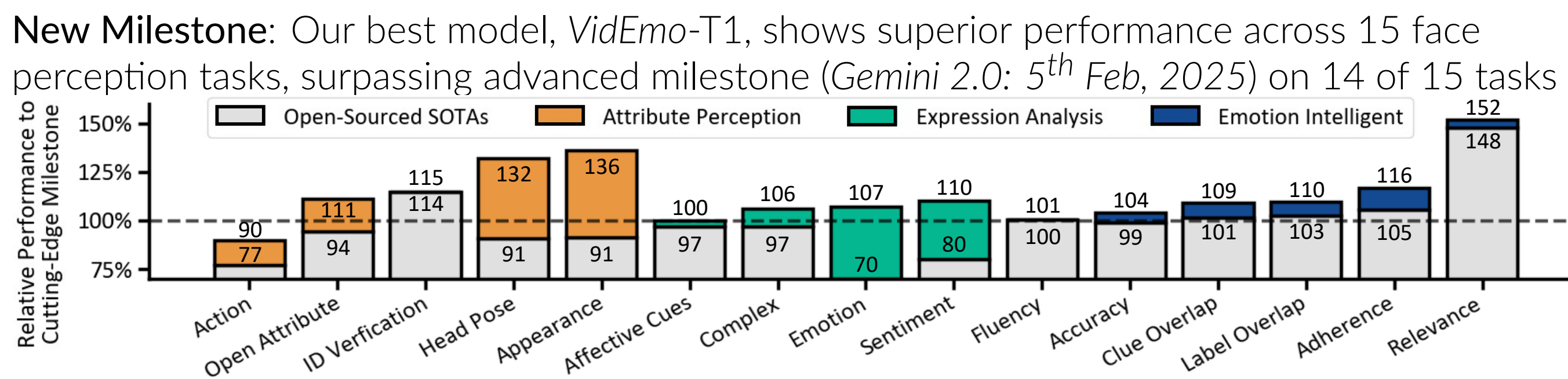


Figure 3. Results Overview.

SOTA comparison: We compare VidEmo with 18 leading VideoLLMs on 14 face attribute perception, 11 expression analysis and 6 fine-grained emotion understanding tasks. VidEmo achieves over a 16.3% and 14.2% improvement compared to existing open-source VideoLLMs on 1-3B and 7-8B scales.

Model	Size	Appearance				Action	Intent	Verbal	Eye	On-Shop				Attr				Fut	Acc	Gend	Skin	Avg	Model	Size	Emotion				Sentiment	Cues	Complex	Avg	Emotion Understanding				
		Cue	Cap	Chp	Cup					Eye	Most	Nose	Hand	Chin	Shp	Face	Acc								Gend	Skin	Size	Min					Emot	Fin	Sent	Em	Ave
Closed MLM3 API																																					
Gemini 2.0 [62]	422	564	414	622	865	258	524	721	729	873	641	619	721	800	787	641	619	721	800	787	641	619	721	800	787	641	619	721	800	787	641	619	721	800	787	641	619
Gemini 1.5 [62]	301	388	357	611	616	220	420	544	773	661	619	619	721	800	787	641	619	721	800	787	641	619	721	800	787	641	619	721	800	787	641	619	721	800	787	641	619
Gemini-VLM-MAX [63]	411	543	540	683	835	320	519	696	794	844	640	715	644	714	789	933	640	715	644	714	789	933	640	715	644	714	789	933	640	715	644	714	789	933	640	715	644
Gemini-VLM [63]	301	388	357	611	616	220	420	544	773	661	619	619	721	800	787	641	619	721	800	787	641	619	721	800	787	641	619	721	800	787	641	619	721	800	787	641	619
GPT-4o-mini [64]	205	258	401	560	693	278	342	450	520	326	416	401	450	450	450	450	450	450	450	450	450	450	450	450	450	450	450	450	450	450	450	450	450	450	450	450	
Open-source 1.3B Video MLM																																					
InternV2.5 [23]	177	402	307	476	503	144	256	410	600	509	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518
InternV2.5 [23]	177	402	307	476	503	144	256	410	600	509	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518
VidMLA-X [86]	208	363	445	482	892	206	334	551	502	717	452	367	417	552	854	593	502	717	452	367	417	552	854	593	502	717	452	367	417	552	854	593	502	717	452	367	417
VidLLaVA-N [65]	177	402	307	476	503	144	256	410	600	509	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518
Qwen2.5-VLM [66]	38	436	41	302	499	957	304	600	642	543	728	201	408	521	589	706	367	621	626	605																	
VidFiMo-base	38	57.0	67.9	37.7	47.9	100	70.7	66.9	84.7	82.9	85.2	75.9	80.8	78.0	83.4	84.0	90.9	88.8	61	82.8																	
Open-source 78s Video MLM																																					
InternV2.5 [23]	177	402	307	476	503	144	256	410	600	509	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518
InternV2.5 [23]	177	402	307	476	503	144	256	410	600	509	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518
VidLLaVA-N [65]	177	402	307	476	503	144	256	410	600	509	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518	467	408	481	298	876	201	619	452	518
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	122	46.8	61.2	155	36.6	56.3	59.0	55.0	72.2	52.4	36.2	52.7	61.5	60.9	57.6	59.3																		
VidLLaVA-O [25]	38	36.9	37	12																																	

Emo-CFG: Emotion-Centric Fine-Grained Video Dataset

Data Curation: We curate Emo-CFG from 17 diverse video datasets, preserving rich meta information. We further obtain rationale annotations linking attribute, expression, and emotion. All data is rigorously filtered through a committee, ensuring high-quality emotion annotations.

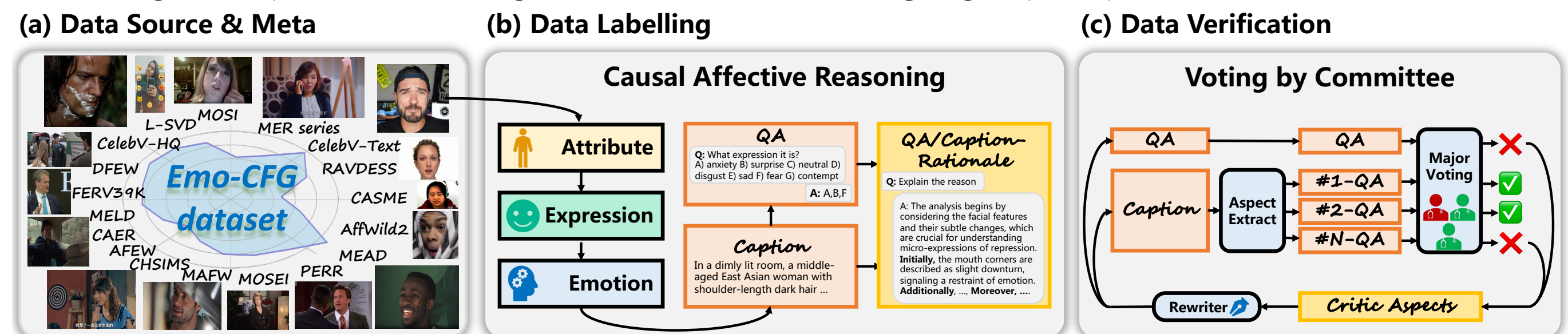


Figure 4. **Data Curation Pipeline of the dataset.** (a) The source of data from 17 datasets. (b) The illustration of data labeling steps. (c) The illustration of data verification loop.

Data Statistics: Emo-CFG covers three major task types with diverse facial views, video lengths, and affective annotations. Its large scale, rich label variety, and inclusion of fine-grained emotions and rationales make it substantially more comprehensive than existing emotion or video datasets.

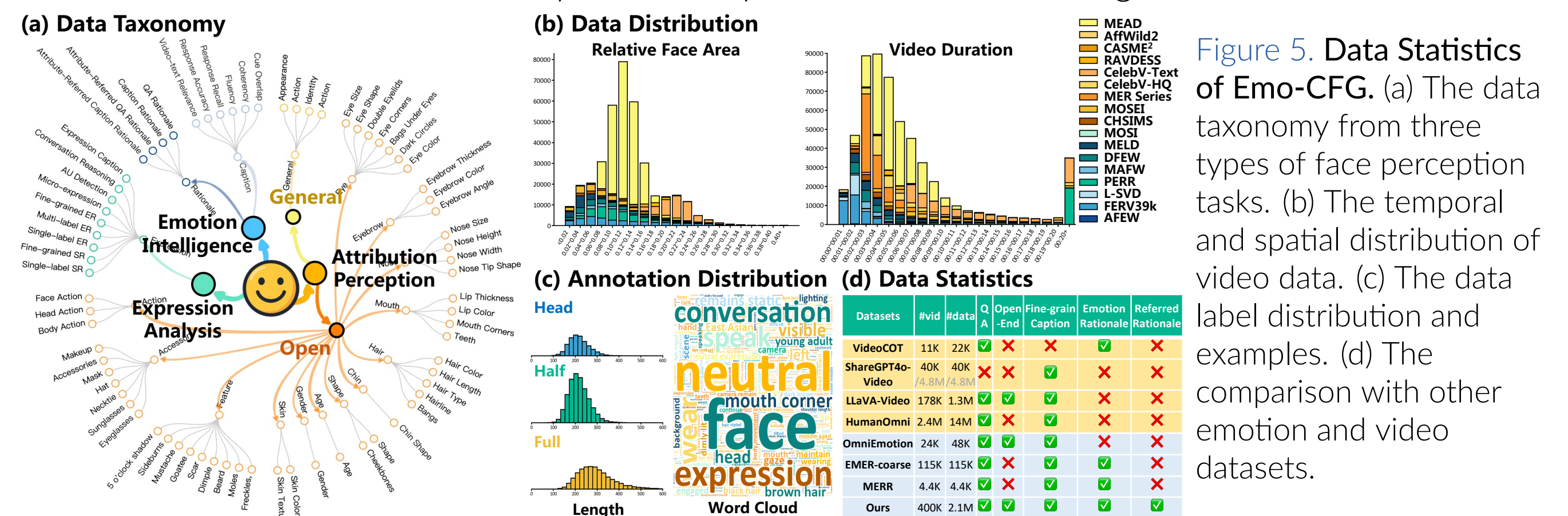


Figure 5. Data Statistics of Emo-CFG. (a) The data taxonomy from three types of face perception tasks. (b) The temporal and spatial distribution of video data. (c) The data label distribution and examples. (d) The comparison with other emotion and video datasets.

User Study: Emo-CFG is built to preserve high data quality even beyond 50K samples, a scale at which consistency becomes difficult for human-labeled datasets such as CelebV-Text. As a fully automated and systematically verified pipeline, Emo-CFG delivers reliable, large-scale emotional annotations across diverse video sources.

Dimension	Pairwise			#Vid	#Usr	Prefer	p-value
	Win	Tie	Loss				
Precision	964	204	82	50	25	95.5%	0.00021
Rationality	1082	87	81	50	25	92.1%	0.00015
Complementary	1172	23	55	50	25	93.0%	0.00008

Table 2. User study between VidEmo and CelebV-Text. We test the label quality on precision, rationality, and complementary through pairwise comparison with 50 videos and 25 users. All three dimensions show significant preference for VidEmo.

Conclusion

- We introduce VidEmo, a unified video emotion foundation model that integrates attribute perception, expression analysis, and high-level affect reasoning.
- We build Emo-CFG, a 2.1M-sample emotion-centric dataset with rich fine-grained annotations, establishing a comprehensive data infrastructure for community.

Contact
MODA
*gloryzzc6@
sina.com*

Website Project Code