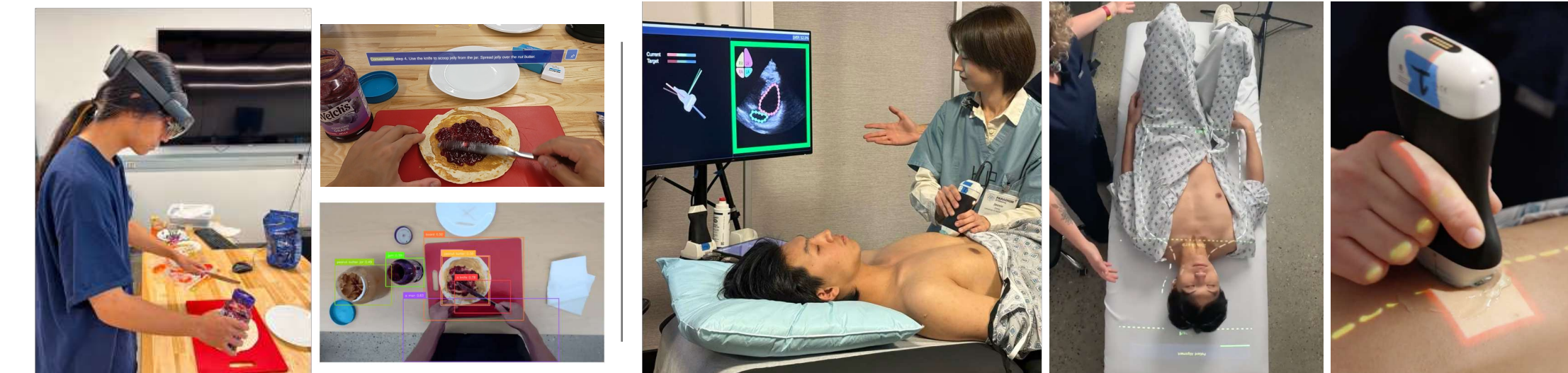


Mistake Understanding — What? & Why?

- Envisioning Instructional AI:

Upskill novice to complete professional tasks



Daily Cooking

Professional Radiology

- Human mistake is inevitable:
Beyond Recognition and Prediction

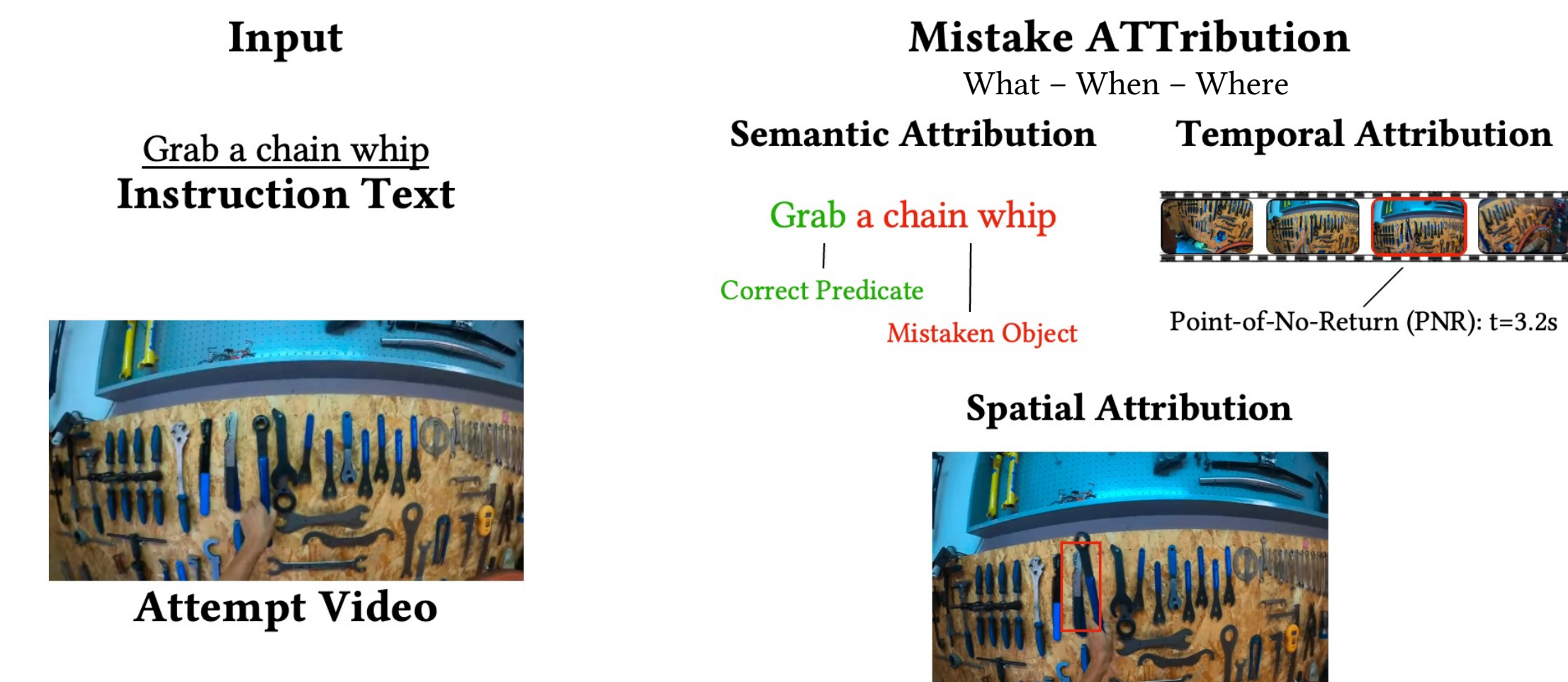
Existing Mistake Understanding

- Lack of diagnostic details



Instruction Text: "Grab a chain whip"

Mistake ATtribution — Problem Formulation



Technical Challenges

- Benchmark: Hard to manually collect large scale mistake dataset and attribution annotation
- Method: Lack of a unified model to produce semantic (what), temporal (when), spatial (where) attribution

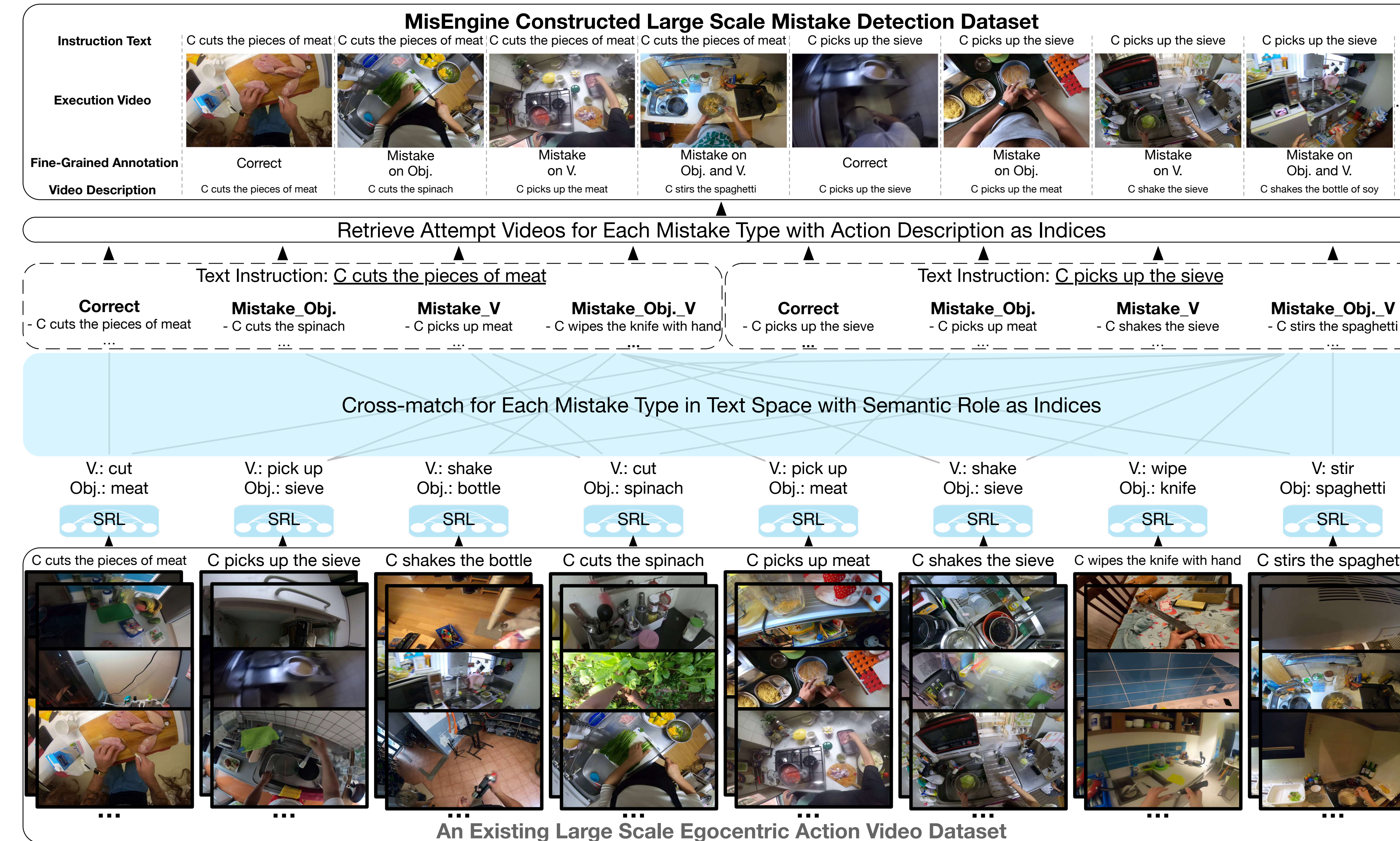


Yayuan Li¹, Aadit Jain¹, Filippos Bellos¹, Jason J. Corso^{1,2}
¹University of Michigan, ²Voxel 51



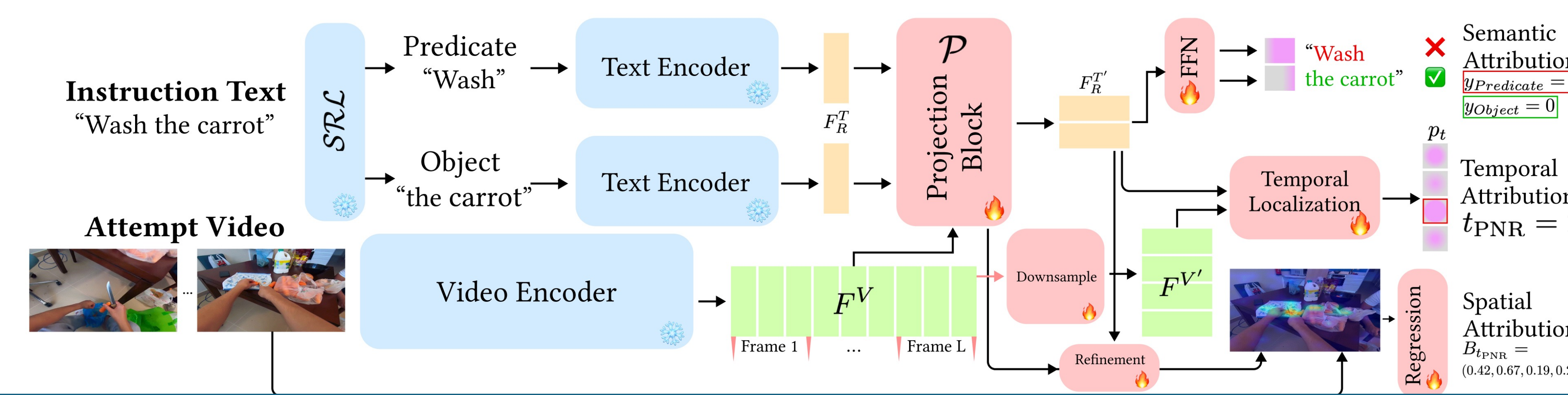
MisEngine – Automatic Dataset Construction

- The essence of mistakes is the divergence of video from text



MisFormer – Unified Structure for MATT

- Pretrained VLM for feature extraction



- Tailored attribution heads

Quantitative Results

Dataset	Method	# Samples		Semantic Var.		Visual Var.		Annotation			
		Total	By activity	Activ.	Dm.	Env.	Part.	Det.	Sem.	Temp.	Spa.
EgoPER [23]	LLaVA-Vid [49]	599	9.7	62	1	2	11	✓	▲	×	▲
	Assembly101 [4, 40]	707	2.0	358	1	1	53	✓	▲	×	×
	HoloAssist [46]	7,562	0.91	8,285	2	-	222	✓	▲	×	×
	CaptionCook4D [37]	1,964	5.6	352	1	10	8	✓	▲	×	×
	Epic-Tent [18]	626	52.2	12	1	-	24	✓	▲	×	×
EPIC-KITCHENS	Ego4D	89,975	5.0	18,003	1	45	37	✓	×	×	▲
	Ego4D	153,367	2.1	75,423	14+	100+	931	✓	×	×	▲
EPIC-KITCHENS-M	Ego4D	221,094	18.0	12,283	1	45	37	✓	✓	×	▲
	Ego4D-M	257,584	16.0	16,099	14+	100+	248	✓	✓	✓	✓

Dataset Profile

Method	MAE (frames) ↓	MAE (seconds) ↓
EgoMotion-COMPASS [24]	48.96	1.632
EgoT2 [48]	24.48	0.816
MisFormer (Ours)	19.14	0.638

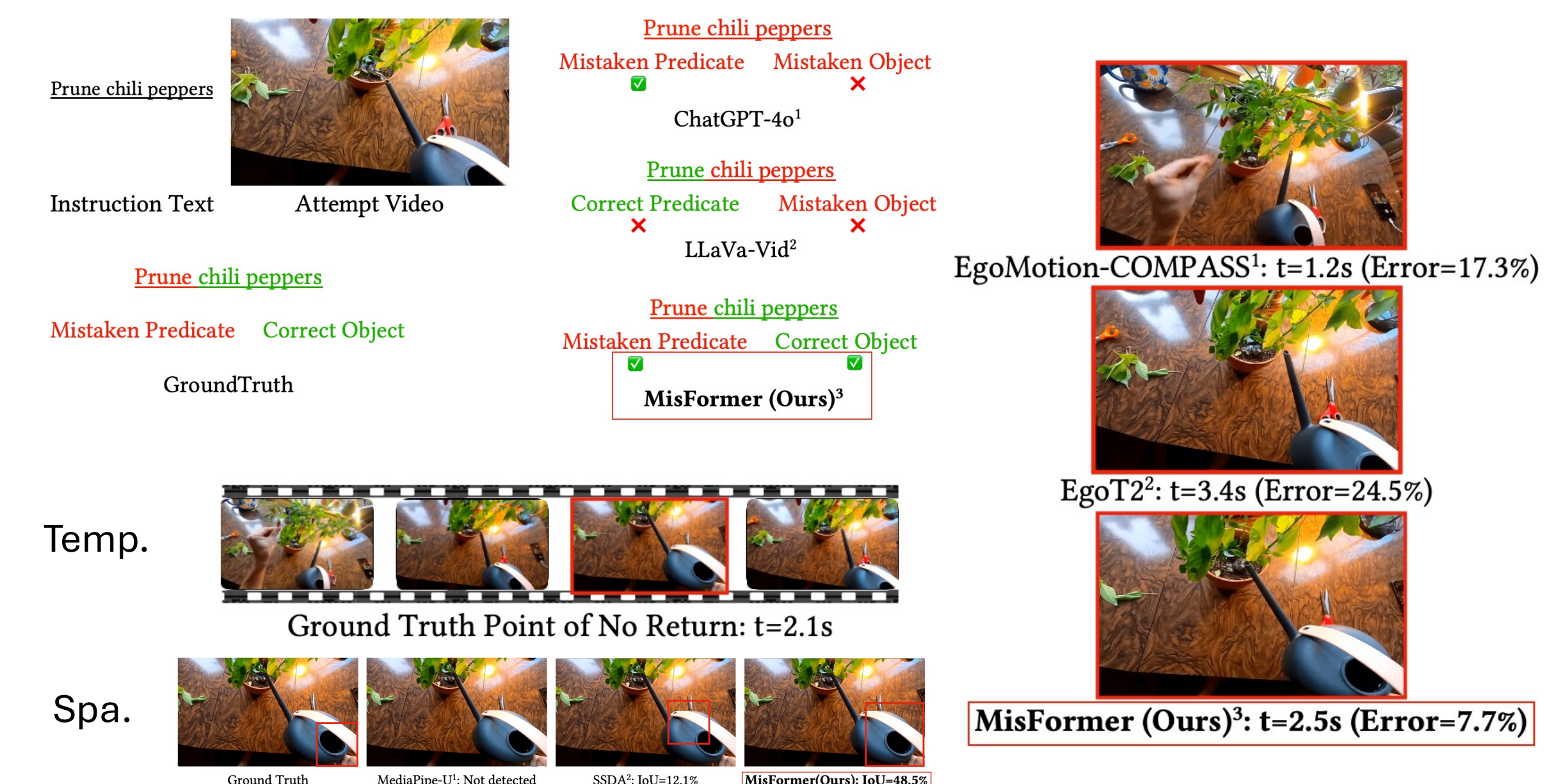
Temporal Attribution

Semantic Attribution

Method	mIoU (%) ↑	CD (%) ↓	BSE (%) ↓
MediaPipe-U [28]	49.88	13.47	20.98
SSDA [25]	64.54	7.21	12.34
MisFormer (Ours)	59.21	10.36	16.27

Spatial Attribution

Qualitative Results



Instruction Text: **Empty** the kettle

