

COLLAGE: Collaborative Human-Agent Interaction Generation using Hierarchical Latent Diffusion and Language Models

Divyanshu Daiya¹, Damon Conover², Aniket Bera¹

¹ Purdue University, ² DEVCOM Army Research Laboratory

Motivation

Challenges. Generating collaborative (Human–Object–Human) motions is difficult because of:

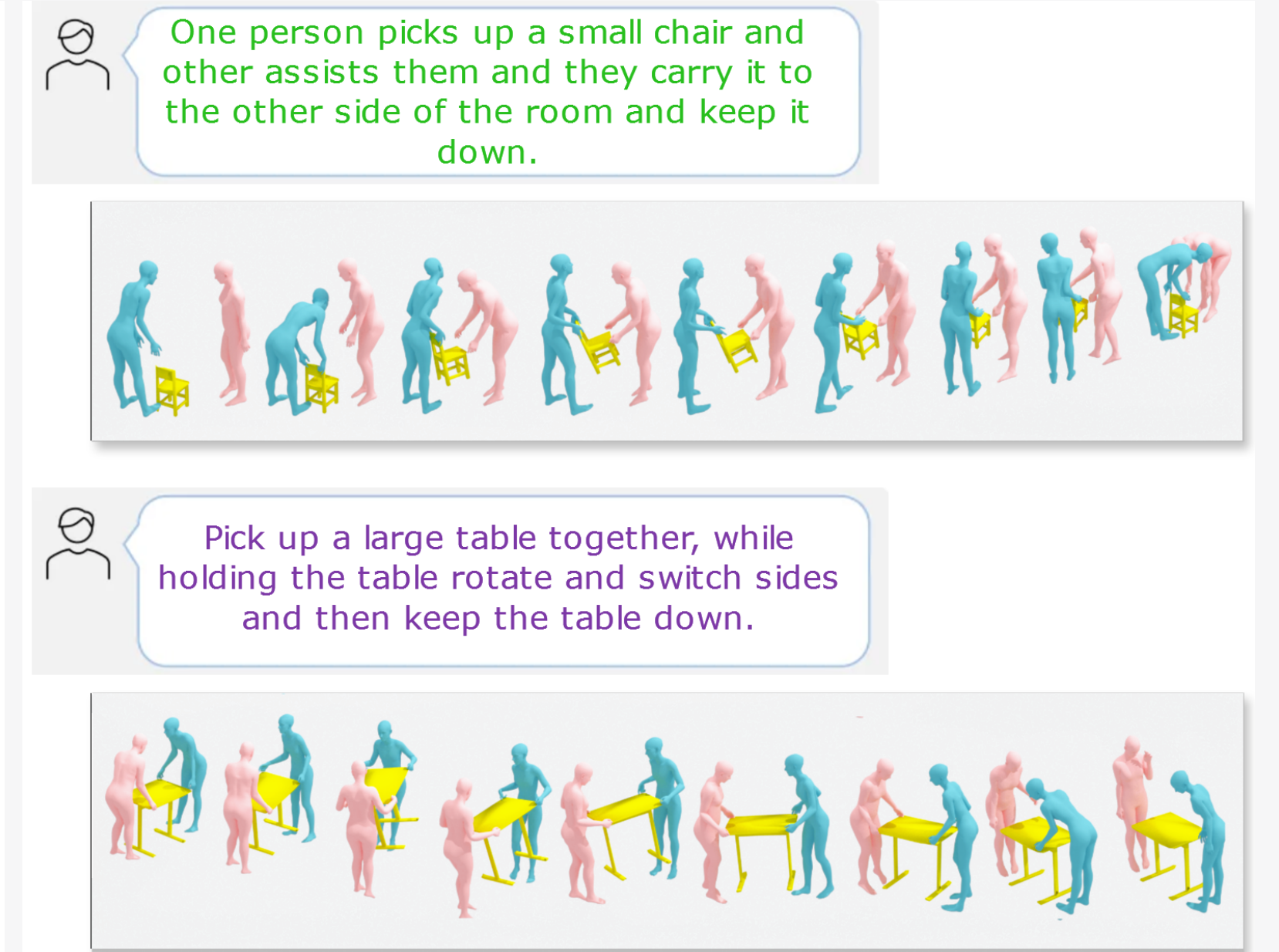
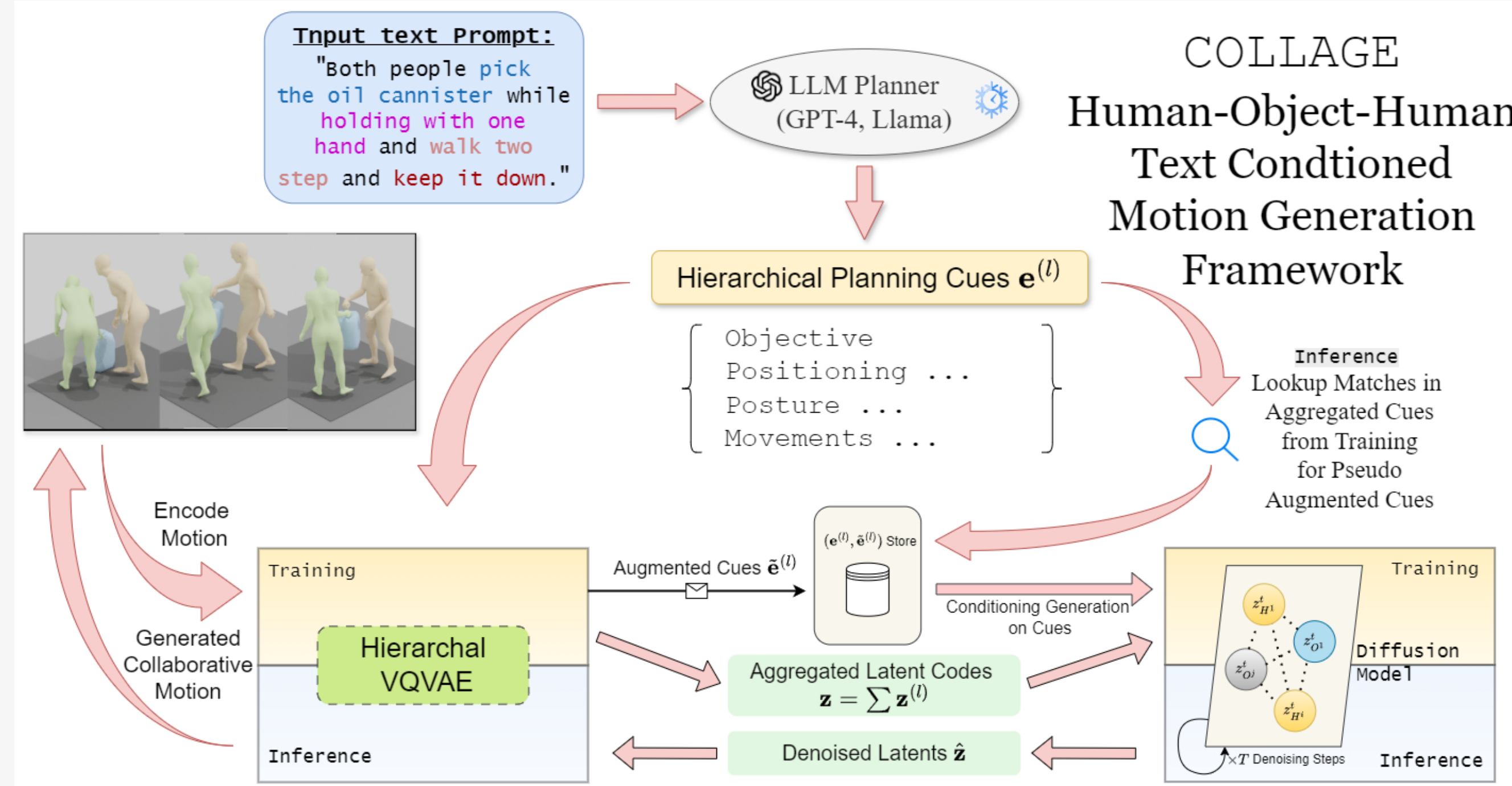
- Complex multi-agent planning,
- Limited **datasets that capture fine-grained synergy**,
- Need for fine control over interactions.

Key Contributions.

- **LLM-guided hierarchical VQ-VAE** supplies a structured, coarse-to-fine action script, directly taming the complexity of multi-agent planning.
- **Text ↔ code semantic anchors** learned via contrastive training enrich sparse data with strong language priors, compensating for limited fine-grained synergy examples.
- **Anchor-conditioned diffusion** refines motions in codebook space within **5–55 steps**, delivering precise, high-fidelity control and achieving SOTA generative performance on CORE4D and InterHuman datasets.

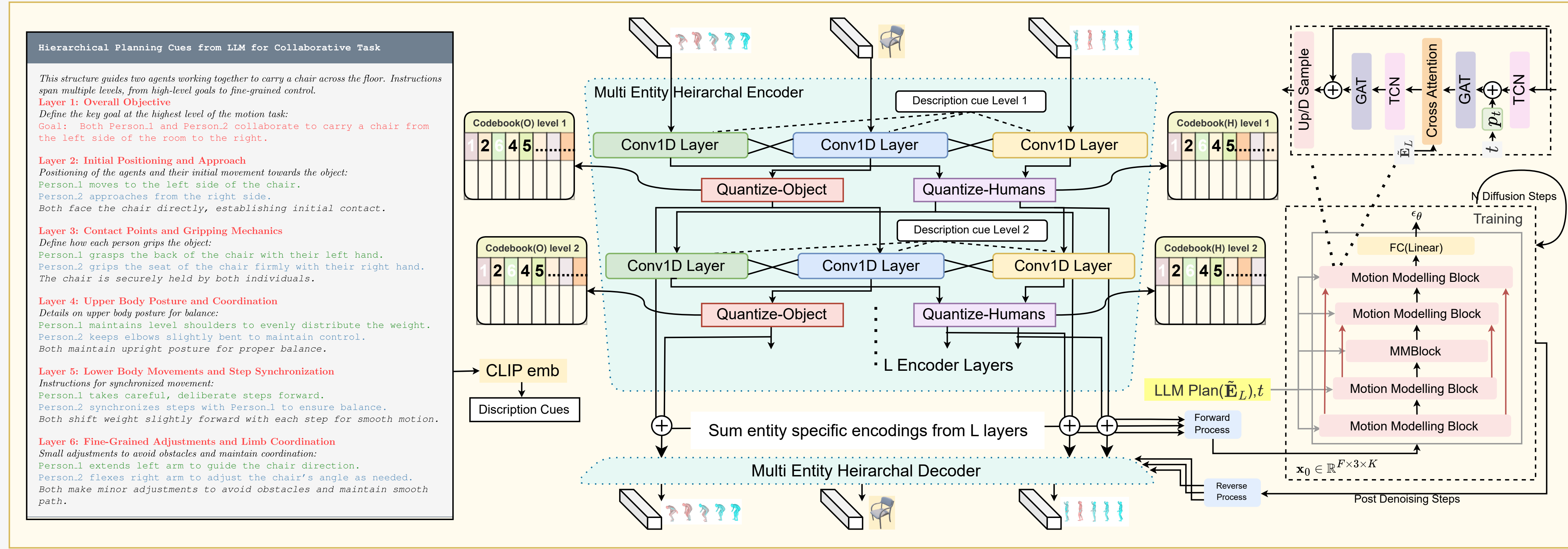
Overview

Goal: Generate realistic multi-actor + object collaboration from text or geometry.



High-Level Insight: LLM-based semantics + hierarchical latents yield coherent, controllable multi-agent interactions.

Methodology



Hierarchical VQ-VAE.

- Each layer encodes a different motion scale, and **LLM text cues** are injected at every level to steer semantics in real time.
- **Cross-code / text alignment** binds discrete latent codes to sentence embeddings, giving robust semantic anchors.
- Those anchors guide **anchor-conditioned diffusion** at test-time, so motions stay text-consistent while requiring markedly fewer denoise steps.
- A **compact loss mix**—reconstruction, contact, penetration and alignment—jointly enforces physical plausibility and semantic fidelity.
- The resulting latents are **disentangled and stable**, enabling smooth multi-agent synergy.

Latent Generation & Planning.

- A **Graph U-Net** built from TCN + GAT blocks models multi-actor synergy.
- Training minimises the **noise-prediction loss** $\| \epsilon - \epsilon_\theta(\mathbf{x}_t, \mathbf{E}_L) \|^2$.
- **Time-modulated weights** emphasise broad constraints early and fine detail late.
- Converges in just **5–55 DDIM steps**.

Incorporating Language.

- **Contrastive links** tie textual prompts to the learned codebook, guaranteeing semantic relevance.
- **Cross-attention** fuses text and motion signals inside every TCN/GAT layer.
- **Hierarchical GPT cues** such as “Approach → Grasp → Lift” guide generation step-by-step along the diffusion timeline.

Iterative Refinement.

- **LLM sub-actions** are injected at each partial denoise step, acting as intermediate waypoints.
- Codes **progressively align** to those sub-actions while Gaussian noise shrinks, tightening the motion trajectory.
- The procedure preserves consistent agent – agent - object collaboration even on long, complex tasks.

Experimental Results

Datasets: CORE-4D & InterHuman; **Metrics:** R-Precision (Top-1/2/3), FID, Diversity, Multimodality, MM-Distance, Joint Position Error (RR.J_e), Vertex Position Error (RR.V_e), Contact Accuracy (C_{acc}).

Text-Conditioned Generation on CORE4D and InterHuman

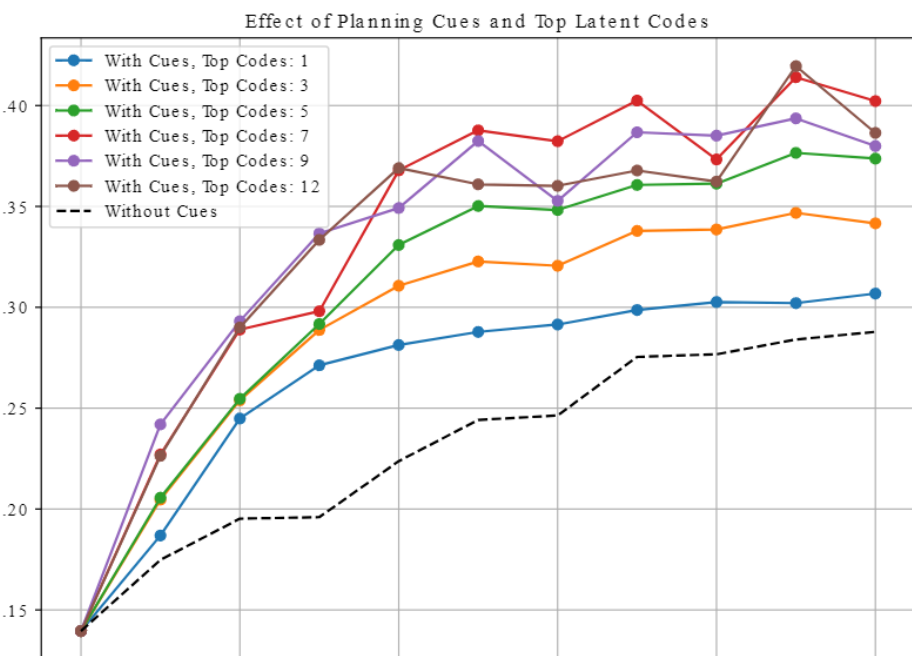
Methods(CORE-4D)	R Precision [↑]			FID [↓]	MM Dist [↓]	Diversity [→]	MModality [↑]	Methods (InterHuman)	R Precision [↑]			FID [↓]	MM Dist [↓]	Diversity [→]	MModality [↑]
	Top 1	Top 2	Top 3						Top 1	Top 2	Top 3				
Real	0.312±0.007	0.587±0.006	0.673±0.006	0.005±0.0005	4.124±0.019	8.151±0.091	-	Real	0.452±0.008	0.610±0.009	0.701±0.008	0.273±0.007	3.755±0.008	7.948±0.064	-
TEMOS	0.065±0.006	0.179±0.006	0.211±0.005	9.214±0.0758	8.536±0.019	4.671±0.091	0.510±0.052	TEMOS	0.224±0.010	0.316±0.013	0.450±0.018	17.375±0.043	6.342±0.015	6.939±0.071	0.535±0.014
T2M	0.195±0.003	0.141±0.003	0.267±0.002	11.258±0.0694	5.867±0.013	2.738±0.078	1.672±0.041	T2M	0.238±0.012	0.325±0.012	0.464±0.014	13.769±0.072	5.731±0.018	7.046±0.022	1.387±0.076
MDM	0.163±0.013	0.257±0.010	0.348±0.008	9.671±0.0629	10.219±0.020	7.395±0.090	3.526±0.074	MDM	0.153±0.009	0.260±0.011	0.339±0.012	9.167±0.056	7.125±0.018	7.602±0.045	2.355±0.080
MDM(GRU)	0.168±0.009	0.279±0.008	0.361±0.010	9.587±0.1382	10.228±0.025	6.951±0.151	3.170±0.046	MDM(GRU)	0.179±0.006	0.299±0.005	0.387±0.007	32.617±0.1221	9.557±0.019	7.003±0.134	3.430±0.035
ComMDM	0.187±0.005	0.256±0.007	0.301±0.007	9.217±0.0727	7.541±0.023	5.367±0.080	0.721±0.065	ComMDM	0.223±0.010	0.334±0.008	0.466±0.012	7.069±0.054	6.212±0.021	7.244±0.038	1.822±0.052
InterGen	0.206±0.007	0.312±0.008	0.401±0.008	7.217±0.2321	10.251±0.017	6.162±0.225	3.402±0.063	InterGen	0.371±0.010	0.515±0.012	0.624±0.010	5.918±0.079	5.108±0.014	7.387±0.029	2.141±0.063
COLLAGE	0.229±0.008	0.332±0.009	0.435±0.009	6.890±0.2198	5.526±0.016	7.373±0.237	3.589±0.066	COLLAGE	0.383±0.005	0.547±0.006	0.657±0.006	4.987±0.2061	4.992±0.012	7.515±0.214	2.872±0.057
w/o Hierarchy	0.201±0.007	0.309±0.008	0.411±0.008	7.452±0.2381	5.582±0.018	6.995±0.224	3.209±0.058	w/o Hierarchy	0.355±0.009	0.521±0.010	0.632±0.009	5.543±0.2154	5.048±0.015	7.137±0.201	2.492±0.061
w/o LLM	0.208±0.007	0.315±0.008	0.419±0.008	7.235±0.2305	5.561±0.017	6.549±0.230	3.152±0.063	w/o LLM	0.362±0.008	0.528±0.009	0.639±0.008	5.326±0.2087	5.027±0.014	6.691±0.207	2.435±0.059
w/o Time Modulation	0.218±0.008	0.317±0.009	0.420±0.009	7.071±0.2251	5.556±0.017	7.263±0.234	3.474±0.065	w/o Time Modulation	0.372±0.007	0.536±0.008	0.647±0.007	5.162±0.2129	5.021±0.013	7.405±0.211	2.767±0.062

Object-Conditioned

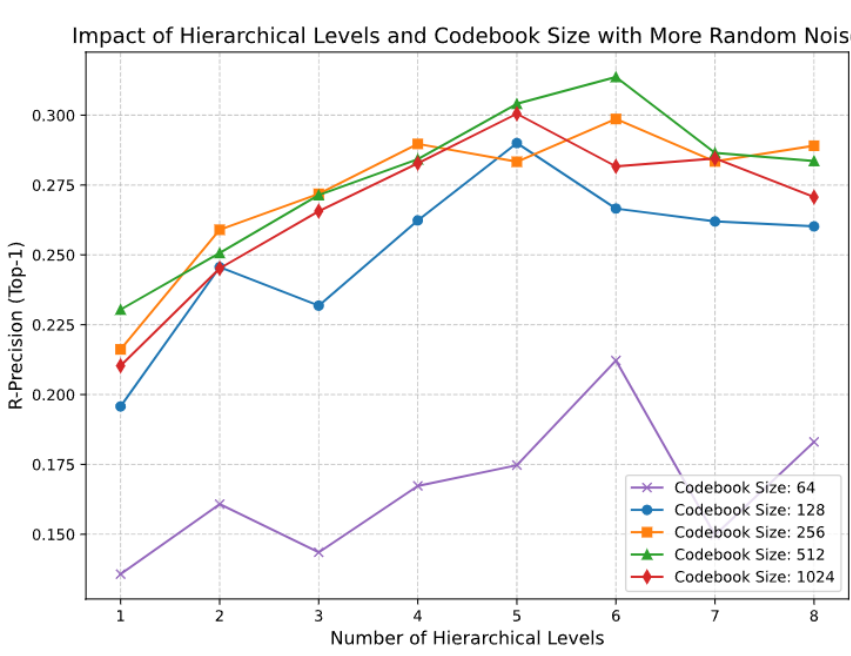
Test Set	Method	RR.J _e (mm, ↓)	RR.V _e (mm, ↓)	C _{acc} (% ↑)	FID (↓)
S1	MDM	138.0 (± 0.3)	194.6 (± 0.2)	76.9 (± 0.5)	7.7 (± 0.2)
	OMOMO	137.8 (± 0.2)	196.7 (± 0.3)	78.2 (± 0.5)	8.3 (± 0.6)
	COLLAGE	131.2 (± 0.2)	185.1 (± 0.2)	80.5 (± 0.4)	7.2 (± 0.2)
S2	MDM	145.9 (± 0.2)	208.2 (± 0.2)	76.7 (± 0.1)	7.7 (± 0.2)
	OMOMO	145.2 (± 0.6)	209.9 (± 1.0)	77.8 (± 0.3)	8.3 (± 1.0)
	COLLAGE	138.5 (± 0.5)	198.7 (± 0.8)	79.9 (± 0.2)	7.3 (± 0.8)

Observations

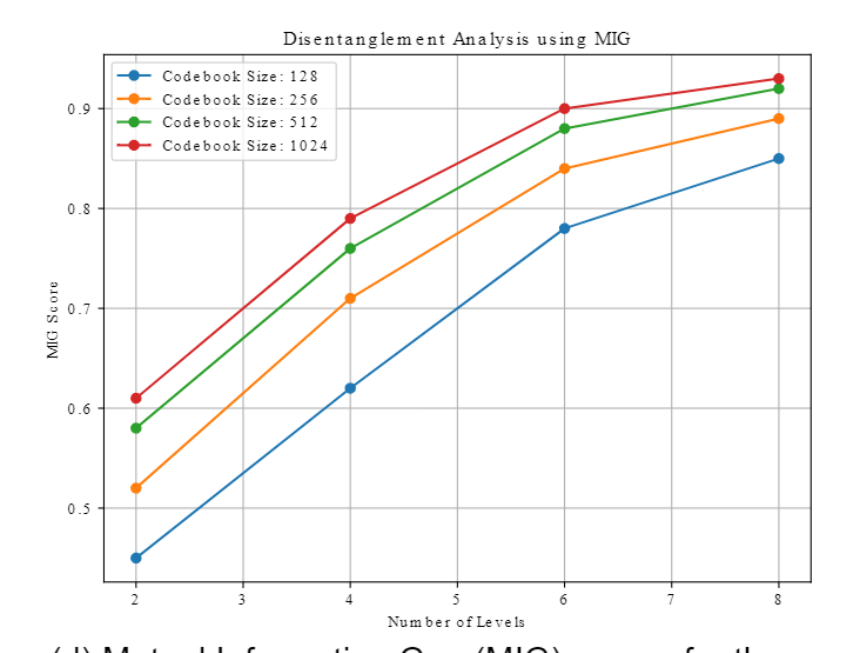
- **Hierarchy Matters:** Removing levels harms synergy and quality, drops R-Precision by **12 %** and worsens FID by **8 %**
- **Language Guidance is crucial** Excluding LLM cues lowers R-Precision by **9 %** and raises FID by **5 %**.
- With only 15 steps, FID is **4 %** better than InterGen while **65 %** faster; 55 steps yields peak quality.
- **Diversity & Task Conformity:** Hierarchical latent diffusion explores a **broad codebook space** for varied interactions, while LLM-guided anchors ensure **prompt consistency**.
- **SOTA Performance:** Against InterGen, R-Precision improves by **+11 %** (CORE-4D) and **+3 %** (InterHuman) and FID drops by **−5 %** and **−16 %**.



(a) Effect of LLM-guided planning cues and the number of top latent codes on the diffusion speed and denoising process.

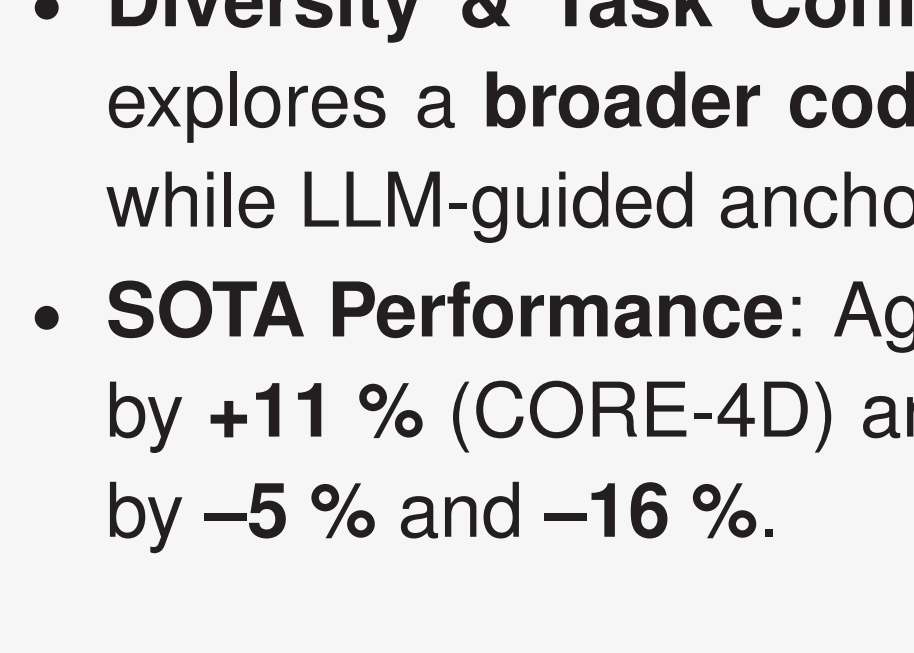


(b) Impact of the number of hierarchical levels and codebook sizes in the VQ-VAE on the R-Precision (top-1) scores on the CORE-4D dataset.



Method	Inference Time (s)
TEMOS	1.24 (± 0.08)
T2M	0.96 (± 0.05)
MDM	0.85 (± 0.06)
ComMDM	1.12 (± 0.09)
InterGen	0.79 (± 0.04)
OurModel	0.58 (± 0.02)

(c) Average inference time per sample on the CORE-4D dataset.



(d) Mutual Information Gap (MIG) scores for the hierarchical VQ-VAE model across different numbers of levels and codebook sizes.

For more information, and visualizations, Please visit / Scan →

