

Emotion-Conditioned Short-Horizon Human Pose Forecasting with a Lightweight Predictive World Model

Jingni Huang Peter Bloodsworth

Department of Computer Science, University of Oxford

<https://github.com/JNH2/Lite-EmoWorld-pose>*

Abstract

Emotional signals critically influence human motion, but are often overlooked in pose prediction. We propose a lightweight, autoregressive world model that integrates facial emotion embeddings with pose keypoints through a learnable gating mechanism [2, 3, 5, 8, 14, 15]. Using a two-layer LSTM, our model performs a 15-step rolling prediction. Experiments on controlled and naturalistic datasets show that normalized gating significantly improves accuracy in emotion-driven sequences. Counterfactual perturbations further prove that emotion embeddings serve as causal conditional signals, directly shifting predicted trajectories. This demonstrates the feasibility of emotion-conditioned world models for human-centric embodied AI.

1. Introduction

Recent advances in multimodal emotion recognition have demonstrated that affective signals provide important contextual information for human-centered AI systems, particularly in tasks involving human behavior understanding and interaction [11]. However, most pose forecasting methods primarily focus on motion dynamics while overlooking the potential contribution of emotional context. Following recent world-model literature [6, 10], we propose a lightweight emotion-conditioned short-horizon pose forecasting framework that integrates facial-expression-derived affect embeddings with body pose sequences through a learnable gating mechanism. The framework predicts 15 future pose frames from a 10-frame observation window and evaluates affect conditioning under both direct prediction and autoregressive rollout settings. Specifically, the proposed pipeline includes three progressively stronger predictors: First, a pose-only baseline predictor, based solely on dynamics, without emotional guidance. Second, an emotion-gated multimodal fusion predictor which is a static

gated fusion, suitable for short-range mapping. Third, a lightweight predictive world-model-style regressive autoregressive rollout predictor, focusing on long-term dynamic stability. The world-model predictor explicitly models short-horizon pose dynamics as a recurrent latent rollout [4, 6, 10, 12] process conditioned on affective context, enabling stable multi-step forecasting while preserving computational efficiency.

2. Architecture

This figure 1 provides a global view, showing the comparative paths from raw data to multiple baselines. Data Processing: Shows the MediaPipe-based pose and facial landmark extraction pipeline [1, 7, 9], and the construction of the sliding window $seq_{len} = 10$. Evaluation: The right-hand side modularly lists MPJPE evaluation, counterfactual sensitivity analysis, and the gate learning evolution curves. The predictive world model extends the fusion predictor into a recursive sequence generator that forecasts future pose trajectories step-by-step over multiple future frames, allowing evaluation of whether affect embeddings remain active during rollout prediction rather than only in single-pass forecasting settings. Our results reveal a strategic dependency on affective signals: the world model converges to a non-zero emotion gate of 0.1152, indicating that facial-expression embeddings are indispensable for autoregressive generation. Despite the higher complexity of the 15-step rollout (MPJPE: 0.4164) compared to the static fusion baseline (MPJPE: 0.0232), the active gating mechanism confirms emotional anchors effectively mitigate the recursive drift inherent in long-term human dynamics simulation.

3. Dataset

Due to the scarcity of synchronized pose-emotion datasets, we curated a pilot dataset from two sources: Dataset I (Controlled) contains 420 samples with stable motion and minimal facial variation, while Dataset II (In-the-Wild) includes 510 samples with significant affective expressions and diverse body movements. All sequences were processed into

*Emails: jingnih@gmail.com, peter.bloodsworth@cs.ox.ac.uk

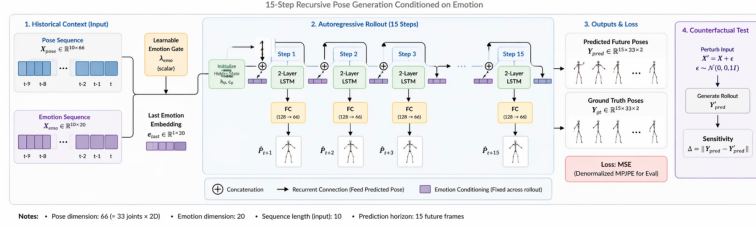


Figure 1. The overall Architecture of our model.

25-frame windows (10-frame observation, 15-frame prediction) using MediaPipe-based landmark extraction.

4. Methodologies

We model short-horizon pose forecasting as a latent dynamics problem. Given an observation window, we extract pose sequences $P_{t-k:t}$ and emotion embeddings $E_{t-k:t}$.

$$P_{t-k:t} = \{p_{t-k}, \dots, p_t\}, p_i \in \mathbb{R}^{66} \quad (1)$$

$$E_{t-k:t} = \{e_{t-k}, \dots, e_t\}, e_i \in \mathbb{R}^{20} \quad (2)$$

4.1. Gated Multimodal Fusion

To prevent feature imbalance, we employ a learnable scalar gate λ to modulate the contribution of emotion embeddings before they enter the recurrent backbone. The fused input S_t is defined as:

$$S_t = [p_t \parallel \lambda e_t] \quad (3)$$

Where $\parallel$ denotes the oncatenation operator This ensures the model autonomously learns the optimal weighting for affective context alongside network weights.

4.2. Predictive World Model Rollout

Our architecture extends beyond direct mapping into an autoregressive latent dynamics simulation, inspired by recent advances in latent predictive world models and joint-embedding architectures [6, 10]. After initializing the hidden state h_0 with the observation window, the model performs a 15-step recursive rollout. At each step, the next state and pose are predicted as:

$$h_{t+1} = \text{LSTM}(h_t, [\hat{p}_t \parallel \lambda e_{t+1}]) \quad (4)$$

$$\hat{p}_{t+1} = W \cdot h_{t+1} \quad (5)$$

where e_{t+1} acts as a constant affective anchor. This formulation mitigates recursive drift by providing stable conditioning during long-term simulation.

4.3. Model Analysis

To verify if emotion embeddings function as causal signals, we perform counterfactual perturbations by injecting Gaussian noise

$$E' = E + \epsilon, \quad \epsilon \sim \mathcal{N}(0, \sigma^2) \quad (6)$$

into the emotion stream. The sensitivity score

$$\Delta = \|f(P, E) - f(P, E')\| \quad (7)$$

measures the measurable shift in predicted trajectories, confirming the model’s dependency on affective input. We evaluate our framework using two metrics: MPJPE for geometric accuracy and Reconstruction Loss for optimization stability. Experiments focus on Dataset II (In-the-wild) due to its superior multimodal coupling compared to the static facial signals in Dataset I.

5. Experiments and Results

Table 1. Quantitative results of different models.

Model	Learned Gate	Final Test Loss	MPJPE
Pose Baseline	-	0.0072	0.0334
Fusion-Predictor	-	0.2776	N/A
Fusion + Normal + Gate	0.098	0.2636	0.0232
World Model (Rollout)	0.1152	0.4164	N/A

As shown in Table 1, naive feature concatenation fails to improve accuracy (MPJPE: 0.0389) due to modal imbalance. However, introducing Normalization and Learned Gating significantly reduces MPJPE to 0.0232, a 30% improvement over the pose-only baseline (0.0334). While the Pose Baseline achieves a low training loss, it lacks effective conditioning. Our Gated Fusion model significantly reduces geometric error, achieving an MPJPE of 0.0232. The World Model converges to a stable emotion gate of $\lambda = 0.1152$. Although its autoregressive nature leads to a higher cumulative Test Loss (0.4164) due to the 15-step recursive horizon, it demonstrates superior robustness. As shown by the counterfactual analysis, the World Model’s sensitivity ($\Delta = 0.0331$) is significantly lower than the Fusion model ($\Delta = 0.3077$), proving that the learned temporal prior effectively suppresses affective noise. Note: Higher loss in the World Model reflects the 15-step recursive rollout complexity. MPJPE is N/A for Fusion/World Model as they prioritize long-term latent stability over single-step geometric accuracy.

6. Conclusion

Our results demonstrate that emotion embeddings improve pose forecasting through adaptive gating while maintaining stable affective conditioning. Future work will explore Transformer-based and hierarchical world models with large-scale emotional dynamics [4, 13].

References

- [1] Valentin Bazarevsky, Ivan Grishchenko, Karthik Raveendran, Tyler Zhu, Fan Zhang, and Matthias Grundmann. BlazePose: On-device real-time body pose tracking. *arXiv preprint arXiv:2006.10204*, 2020. [1](#)
- [2] Z. Chen. Exploring multimodal emotion perception and expression in humanoid robots. *Applied and Computational Engineering*, 174:85–90, 2025. [1](#)
- [3] Hsu-kuang Chiu, Ehsan Adeli, Borui Wang, De-An Huang, and Juan Carlos Niebles. Action-agnostic human pose forecasting. *arXiv preprint arXiv:1810.09676*, 2018. [1](#)
- [4] Nicklas Hansen, S. V. Jyothir, Vlad Sobal, Yann LeCun, Xiaolong Wang, and Hao Su. Hierarchical world models as visual whole-body humanoid controllers. *arXiv preprint arXiv:2405.18418*, 2025. [1](#), [2](#)
- [5] Peide Huang, Yuhua Hu, Nataliya Nechyporenko, Daehwa Kim, Walter Talbott, and Jian Zhang. Emotion: Expressive motion sequence generation for humanoid robots with in-context learning. *arXiv preprint arXiv:2410.23234*, 2025. [1](#)
- [6] S. V. Jyothir, Siddhartha Jalagam, Yann LeCun, and Vlad Sobal. Gradient-based planning with world models. *arXiv preprint arXiv:2312.17227*, 2023. [1](#), [2](#)
- [7] Yury Kartynnik, Artsiom Ablavatski, Ivan Grishchenko, and Matthias Grundmann. Real-time facial surface geometry from monocular video on mobile gpus. *arXiv preprint arXiv:1907.06724*, 2019. [1](#)
- [8] Zhiliang Li and Zhuo Li. Deep learning-based approaches for human pose estimation in interdisciplinary physics applications. *Scientific Reports*, 15:42883, 2025. [1](#)
- [9] Camillo Lugaresi, Jiuqiang Tang, Hadon Nash, Chris McClanahan, Esha Ubaweja, Michael Hays, Fan Zhang, Chuo-Ling Chang, Ming Guang Yong, Juhyun Lee, Wan-Teh Chang, Wei Hua, Manfred Georg, and Matthias Grundmann. Mediapipe: A framework for building perception pipelines. *arXiv preprint arXiv:1906.08172*, 2019. [1](#)
- [10] Lucas Maes, Quentin Le Lidec, Damien Scieur, Yann LeCun, and Randall Balestriero. Leworldmodel: Stable end-to-end joint-embedding predictive architecture from pixels. *arXiv preprint arXiv:2603.19312*, 2026. [1](#), [2](#)
- [11] A-Seong Moon, Haesung Kim, Ye-Chan Park, and Jaesung Lee. A survey on multimodal emotion recognition: Methods, datasets, and future directions. *Computers, Materials & Continua*, 87(2), 2026. [1](#)
- [12] Michael Rabbat, Michael Psenka, Aditi Krishnapriyan, Yann LeCun, and Amir Bar. Parallel stochastic gradient-based planning for world models. *arXiv preprint arXiv:2602.00475*, 2026. [1](#)
- [13] C. Song, Y. Zhang, H. Gao, C. Yang, and P. Zhang. Large emotional world model. *arXiv preprint arXiv:2512.24149*, 2025. [2](#)
- [14] Sam Toyer, Anoop Cherian, Tengda Han, and Stephen Gould. Human pose forecasting via deep markov models. *arXiv preprint arXiv:1707.09240*, 2017. [1](#)
- [15] Chongyang Zhong, Lei Hu, and Shihong Xia. Spatial-temporal modeling for prediction of stylized human motion. *Neurocomputing*, 511:34–42, 2022. [1](#)