

HuMoR: 3D Human Motion Model for Robust Pose Estimation

Davis Rempe¹

Tolga Birdal¹

Aaron Hertzmann²

Jimei Yang²

Srinath Sridhar³

Leonidas J. Guibas¹

¹Stanford University

²Adobe Research

³Brown University

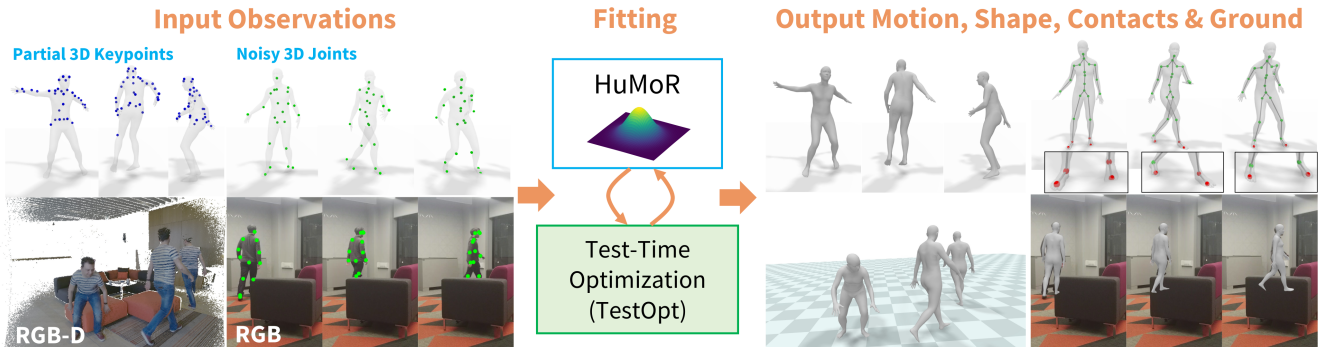


Figure 1: *Overview.* HuMoR is a 3D **H**uman **M**otion model for **R**obust estimation of temporal pose formulated as a conditional variational autoencoder. (Left) The proposed approach can operate on many input modalities and is designed to handle partial and noisy observations. (Middle/Right) A test-time optimization fits 3D motion and shape to an input sequence using HuMoR as a prior; additional outputs include the ground and person-ground contacts (colored as **ground plane** and **contacts**).

Abstract

We introduce HuMoR: a 3D **H**uman **M**otion Model for **R**obust Estimation of temporal pose and shape. Though substantial progress has been made in estimating 3D human motion and shape from dynamic observations, recovering plausible pose sequences in the presence of noise and occlusions remains a challenge. For this purpose, we propose an expressive generative model in the form of a conditional variational autoencoder, which learns a distribution of the change in pose at each step of a motion sequence. Furthermore, we introduce a flexible optimization-based approach that leverages HuMoR as a motion prior to robustly estimate plausible pose and shape from ambiguous observations. Through extensive evaluations, we demonstrate that our model generalizes to diverse motions and body shapes after training on a large motion capture dataset, and enables motion reconstruction from multiple input modalities including 3D keypoints and RGB(-D) videos. See the project page at geometry.stanford.edu/projects/humor.

1. Introduction

As humans, we are constantly moving in, interacting with, and manipulating the world around us. Thus, applications such as action recognition [99, 100] or holistic dy-

namic indoor scene understanding [20] require accurate perception of 3D human pose, shape, motion, contacts, and interaction. Extensive previous work has focused on estimating 2D or 3D human pose [16, 66, 67], shape [73, 32, 83], and motion [47] from videos. These are challenging problems due to the large space of articulations, body shape, and appearance variations. Even the best methods struggle to accurately capture a wide variety of motions from varying input modalities, producing noisy or overly-smoothed motions (especially at ground contact, *i.e.*, footskate), and struggle with occlusions (*e.g.*, walking behind a couch as in Fig. 1).

We focus on the problem of building a robust human motion model that can address these challenges. To date, most motion models directly represent sequences of likely poses — *e.g.*, in PCA space [71, 97, 86] or via future-predicting autoregressive processes [91, 96, 77]. However, purely pose-based predictions either make modeling environment interactions and generalization beyond training poses difficult, or quickly diverge from the space of realistic motions. On the other hand, explicit physical dynamics models [79, 53, 85, 78, 15, 14] are resource intensive and require knowledge of unobservable physical quantities. While generative models potentially offer the required flexibility, building an *expressive*, *generalizable* and *robust* model for *realistic* 3D human motions remains an open problem.

To address this, we introduce a learned, autoregressive, generative model that captures the *dynamics* of 3D human

motion, *i.e.*, how pose changes over time. Rather than describing likely poses, the **Human Motion Model for Robust Estimation** (HuMoR) models a probability distribution of possible *pose transitions*, formulated as a conditional variational autoencoder [88]. Though not explicitly physics-based, its components correspond to a physical model: the latent space can be interpreted as generalized forces, which are inputs to a dynamics model with numerical integration (the decoder). Moreover, ground contacts are explicitly predicted and used to constrain pose estimation at test time.

After training on the large AMASS motion capture dataset [63], we use HuMoR as a *motion prior* at test time for 3D human perception from noisy and partial observations across different input modalities such as RGB(-D) video and 2D or 3D joint sequences, as illustrated in Fig. 1 (left). In particular, we introduce a robust test-time optimization strategy which interacts with HuMoR to estimate the parameters of *3D motion*, *body shape*, the *ground plane*, and *contact points* as shown in Fig. 1 (middle/right). This interaction happens in two ways: (i) by parameterizing the motion in the latent space of HuMoR, and (ii) using HuMoR priors in order to regularize the optimization towards the space of plausible motions.

Comprehensive evaluations reveal that our method surpasses the state-of-the-art on a variety of visual inputs in terms of accuracy and physical plausibility of motions under partial and severe occlusions. We further demonstrate that our motion model generalizes to diverse motions and body shapes on common generative tasks like sampling and future prediction. In a nutshell, our contributions are:

- HuMoR, a generative 3D human motion prior modeled by a novel conditional VAE which enables expressive and general motion reconstruction and generation,
- A subsequent robust test-time optimization approach that uses HuMoR as a strong motion prior jointly solving for pose, body shape, and ground plane / contacts,
- The capability to operate on a variety of inputs, such as RGB(-D) video and 2D/3D joint position sequences, to yield accurate and plausible motions and contacts, exemplified through extensive evaluations.

Our work, more generally, suggests that neural nets for dynamics problems can benefit from architectures that model transitions, allowing control structures that emulate classical physical formulations.

2. Related Work

Much progress has been made on building methods to recover 3D joint locations [76, 67, 66] or parameterized 3D pose and shape (*i.e.*, SMPL [59]) from observations [98]. We focus primarily on motion and shape estimation.

Learning-Based Estimation. Deep learning approaches have shown success in regressing 3D shape and pose from a single image [49, 42, 74, 31, 30, 108, 21]. This has led to developments in predicting *motion* (pose sequences) and shape directly from RGB video [43, 110, 84, 90, 23]. Most recently, VIBE [47] uses adversarial training to encourage plausible outputs from a conditional recurrent motion generator. MEVA [62] maps a fixed-length image sequence to the latent space of a pre-trained motion autoencoder. These methods are fast and produce accurate root-relative joint positions for video, but motion is globally inconsistent and they struggle to generalize, *e.g.*, under severe occlusions. Other works have addressed occlusions but only on static images [10, 111, 80, 27, 48]. Our approach resolves difficult occlusions in video and other modalities by producing plausible and expressive motions with HuMoR.

Optimization-Based Estimation. One may directly optimize to more accurately fit to observations (images or 2D pose estimators [16]) using human body models [25, 7, 11]. SMPLify [11] uses the SMPL model [59] to fit pose and shape parameters to 2D keypoints in an image using priors on pose and shape. Later works consider body silhouettes [51] and use a learned variational pose prior [73]. Optimization for motion sequences has been explored by several works [6, 41, 57, 109, 104] which apply simple smoothness priors over time. These produce reasonable estimates when the person is fully visible, but with unrealistic dynamics, *e.g.*, overly smooth motions and footskate.

Some works employ human-environment interaction and contact constraints to improve shape and pose estimation [34, 57, 35] by assuming scene geometry is given. iMapper [69] recovers both 3D joints and a primitive scene representation from RGB video based on interactions by motion retrieval, which may differ from observations. In contrast, our approach optimizes for pose and shape by using an expressive generative model that produces more natural motions than prior work with realistic ground contact.

Human Motion Models. Early sophisticated motion models for pose tracking used a variety of approaches, including mixtures-of-Gaussians [39], linear embeddings of periodic motion [71, 97, 86], nonlinear embeddings [24], and nonlinear autoregressive models [91, 101, 96, 77]. These methods operate in pose space, and are limited to specific motions. Models based on physics can potentially generalize more accurately [79, 53, 85, 78, 15, 14, 107], while also estimating global pose and environmental interactions. However, general-purpose physics-based models are difficult to learn, computationally intensive at test-time, and often assume full-body visibility to detect contacts [79, 53, 85].

Many motion models have been learned for computer animation [13, 50, 82, 52, 56, 38, 89] including recent recurrent and autoregressive models [33, 29, 36, 105, 54]. These often focus on visual fidelity for a small set of characters

and periodic locomotions. Some have explored generating more general motion and body shapes [112, 75, 3, 22], but in the context of short-term future prediction. HuMoR is most similar to Motion VAE [54], however we make crucial contributions to enable generalization to unseen, non-periodic motions on novel body shapes.

3. HuMoR: 3D Human Dynamics Model

The goal of our work is to build an *expressive* and *generalizable* generative model of 3D human motion learned from real human motions, and to show that this can be used for robust test-time optimization (TestOpt) of pose and shape. In this section, we first describe the model, HuMoR.

State Representation. We represent the state of a moving person as a matrix $\mathbf{x}$ composed of a root translation $\mathbf{r} \in \mathbb{R}^3$, root orientation $\Phi \in \mathbb{R}^3$ in axis-angle form, body pose joint angles $\Theta \in \mathbb{R}^{3 \times 21}$ and joint positions $\mathbf{J} \in \mathbb{R}^{3 \times 22}$:

$$\mathbf{x} = \begin{bmatrix} \mathbf{r} & \dot{\mathbf{r}} & \Phi & \dot{\Phi} & \Theta & \mathbf{J} & \dot{\mathbf{J}} \end{bmatrix}, \quad (1)$$

where $\dot{\mathbf{r}}$, $\dot{\Phi}$ and $\dot{\mathbf{J}}$ denote the root and joint velocities, respectively, giving $\mathbf{x} \in \mathbb{R}^{3 \times 69}$. Part of the state, $(\mathbf{r}, \Phi, \Theta)$, parameterizes the SMPL body model [59, 81] which is a differentiable function $M(\mathbf{r}, \Phi, \Theta, \beta)$ that maps to body mesh vertices $\mathbf{V} \in \mathbb{R}^{3 \times 6890}$ and joints $\mathbf{J}^{\text{SMPL}} \in \mathbb{R}^{3 \times 22}$ given shape parameters $\beta \in \mathbb{R}^{16}$. Our over-parameterization allows for two ways to recover the joints: (i) explicitly from $\mathbf{J}$, (ii) implicitly through the SMPL map $M(\cdot)$.

Latent Variable Dynamics Model. We are interested in modeling the probability of a time sequence of states

$$p_\theta(\mathbf{x}_0, \mathbf{x}_1, \dots, \mathbf{x}_T) = p_\theta(\mathbf{x}_0) \prod_{t=1}^T p_\theta(\mathbf{x}_t | \mathbf{x}_{t-1}), \quad (2)$$

where each state is assumed to be dependent on only the previous one and θ are learned parameters. Then $p_\theta(\mathbf{x}_t | \mathbf{x}_{t-1})$ must capture the *plausibility* of a transition.

We propose a **conditional variational autoencoder (CVAE)** which formulates the motion $p_\theta(\mathbf{x}_t | \mathbf{x}_{t-1})$ as a latent variable model as shown in Fig. 2. Following the original CVAE derivation [88], our model contains two main components. First, conditioned on the previous state $\mathbf{x}_{t-1}$, the distribution over possible latent variables $\mathbf{z}_t \in \mathbb{R}^{48}$ is described by a learned **conditional prior**:

$$p_\theta(\mathbf{z}_t | \mathbf{x}_{t-1}) = \mathcal{N}(\mathbf{z}_t; \mu_\theta(\mathbf{x}_{t-1}), \sigma_\theta(\mathbf{x}_{t-1})), \quad (3)$$

which parameterizes a Gaussian distribution with diagonal covariance via a neural network. Intuitively, the latent variable $\mathbf{z}_t$ represents the transition to $\mathbf{x}_t$ and should therefore have different distributions given different $\mathbf{x}_{t-1}$. For example, an idle person has a large variation of possible next states while a person in midair is on a nearly deterministic

trajectory. Learning the conditional prior significantly improves the ability of the CVAE to generalize to diverse motions and empirically stabilizes both training and TestOpt.

Second, conditioned on $\mathbf{z}_t$ and $\mathbf{x}_{t-1}$, the **decoder** produces two outputs, Δ_θ and $\mathbf{c}_t$. The *change in state* Δ_θ defines the output distribution $p_\theta(\mathbf{x}_t | \mathbf{z}_t, \mathbf{x}_{t-1})$ through

$$\mathbf{x}_t = \mathbf{x}_{t-1} + \Delta_\theta(\mathbf{z}_t, \mathbf{x}_{t-1}) + \boldsymbol{\eta}, \quad \boldsymbol{\eta} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}). \quad (4)$$

We find the additive update Δ_θ improves predictive accuracy compared to direct next-step prediction. The person-ground contact $\mathbf{c}_t$ is the probability that each of 8 body joints (*left and right toes, heels, knees, and hands*) is in contact with the ground at time t . Contacts are *not* part of the input to the conditional prior, only an output of the decoder. The contacts enable environmental constraints in TestOpt.

The complete probability model for a transition is then:

$$p_\theta(\mathbf{x}_t | \mathbf{x}_{t-1}) = \int_{\mathbf{z}_t} p_\theta(\mathbf{z}_t | \mathbf{x}_{t-1}) p_\theta(\mathbf{x}_t | \mathbf{z}_t, \mathbf{x}_{t-1}). \quad (5)$$

Given an initial state $\mathbf{x}_0$, one can sample a motion sequence by alternating between sampling $\mathbf{z}_t \sim p_\theta(\mathbf{z}_t | \mathbf{x}_{t-1})$ and sampling $\mathbf{x}_t \sim p_\theta(\mathbf{x}_t | \mathbf{z}_t, \mathbf{x}_{t-1})$, from $t = 1$ to T . This model parallels a conventional stochastic physical model. The conditional prior can be seen as a controller, producing “forces” $\mathbf{z}_t$ as a function of state $\mathbf{x}_{t-1}$, while the decoder acts like a combined physical dynamics model and Euler integrator of generalized position and velocity in Eq. (4).

In addition to this nice physical interpretation, our model is motivated by Motion VAE (MVAE) [54], which has recently shown promising results for single-character locomotion animation, also using a VAE for $p_\theta(\mathbf{x}_t | \mathbf{x}_{t-1})$. However, we find that directly applying MVAE for estimation does not give good results (Sec. 5). We overcome this by additionally learning a *conditional* prior, modeling the *change in state* and contacts, and encouraging *consistency* between joint position and angle predictions (Sec. 3.1).

Rollout. We use our model to define a deterministic *rollout function*, which is key to TestOpt. Given an initial state $\mathbf{x}_0$ and a sequence of latent transitions $\mathbf{z}_{1:T}$, we define a function $\mathbf{x}_T = f(\mathbf{x}_0, \mathbf{z}_{1:T})$ that deterministically maps the motion “parameters” $(\mathbf{x}_0, \mathbf{z}_{1:T})$ to the resulting state at time T . This is done through autoregressive rollout which decodes and integrates $\mathbf{x}_t = \mathbf{x}_{t-1} + \Delta_\theta(\mathbf{z}_t, \mathbf{x}_{t-1})$ at each timestep.

Initial State GMM. We model $p_\theta(\mathbf{x}_0)$ with a Gaussian mixture model (GMM) containing $K = 12$ components with weights γ^i , so that $p_\theta(\mathbf{x}_0) = \sum_{i=1}^K \gamma^i \mathcal{N}(\mathbf{x}_0; \mu_\theta^i, \sigma_\theta^i)$.

3.1. Training

Our CVAE is trained using pairs of $(\mathbf{x}_{t-1}, \mathbf{x}_t)$. We consider the usual variational lower bound:

$$\log p_\theta(\mathbf{x}_t | \mathbf{x}_{t-1}) \geq \mathbb{E}_{q_\phi} [\log p_\theta(\mathbf{x}_t | \mathbf{z}_t, \mathbf{x}_{t-1})] - D_{\text{KL}}(q_\phi(\mathbf{z}_t | \mathbf{x}_t, \mathbf{x}_{t-1}) \parallel p_\theta(\mathbf{z}_t | \mathbf{x}_{t-1})). \quad (6)$$

Input States

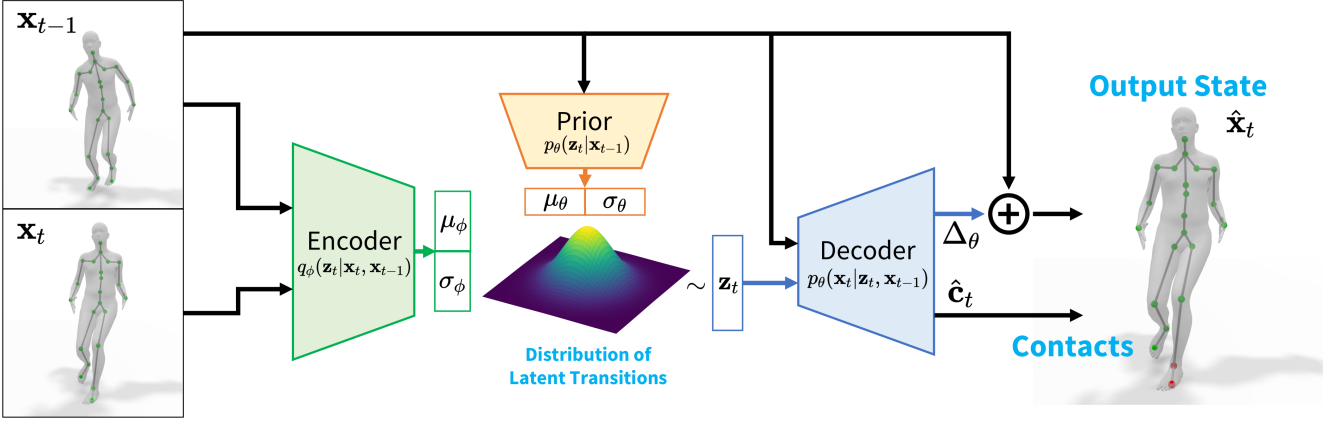


Figure 2: *HuMoR CVAE Architecture*. During training, given the previous state $\mathbf{x}_{t-1}$ and ground truth current state $\mathbf{x}_t$, the model reconstructs $\hat{\mathbf{x}}_t$ by sampling from the **encoder** distribution. At test time we can (i) *generate* the next state from $\mathbf{x}_{t-1}$ by sampling from the **prior** distribution and **decoding**, (ii) *infer* a latent transition $\mathbf{z}_t$ with the **encoder**, or (iii) evaluate the *likelihood* of a given $\mathbf{z}_t$ with the **conditional prior**.

The expectation term measures the reconstruction error of the **decoder**. The **encoder**, *i.e.* approximate posterior, is introduced for training and parameterizes a Gaussian distribution $q_\phi(\mathbf{z}_t|\mathbf{x}_t, \mathbf{x}_{t-1}) = \mathcal{N}(\mathbf{z}_t; \mu_\phi(\mathbf{x}_t, \mathbf{x}_{t-1}), \sigma_\phi(\mathbf{x}_t, \mathbf{x}_{t-1}))$. The KL divergence $D_{\text{KL}}(\cdot \| \cdot)$ regularizes its output to be near the **prior**. Therefore, we seek the parameters (θ, ϕ) that minimize the loss function

$$\mathcal{L}_{\text{rec}} + w_{\text{KL}}\mathcal{L}_{\text{KL}} + \mathcal{L}_{\text{reg}} \quad (7)$$

over all training pairs in our dataset, where $\mathcal{L}_{\text{rec}} + w_{\text{KL}}\mathcal{L}_{\text{KL}}$ is the lower bound in Eq. (6) with weight w_{KL} , and $\mathcal{L}_{\text{reg}}$ contains additional regularizers.

For a single training pair $(\mathbf{x}_{t-1}, \mathbf{x}_t)$, the reconstruction loss is computed as $\mathcal{L}_{\text{rec}} = \|\mathbf{x}_t - \hat{\mathbf{x}}_t\|^2$ from the decoder output $\hat{\mathbf{x}}_t = \mathbf{x}_{t-1} + \Delta_\theta(\mathbf{z}_t, \mathbf{x}_{t-1})$ with $\mathbf{z}_t \sim q_\phi(\mathbf{z}_t|\mathbf{x}_t, \mathbf{x}_{t-1})$. Gradients are backpropagated through this sample using the reparameterization trick [45]. The regularization loss contains two terms: $\mathcal{L}_{\text{reg}} = \mathcal{L}_{\text{SMPL}} + w_{\text{contact}}\mathcal{L}_{\text{contact}}$. The SMPL term $\mathcal{L}_{\text{SMPL}} = \mathcal{L}_{\text{joint}} + \mathcal{L}_{\text{vtx}} + \mathcal{L}_{\text{consist}}$ uses the output of the body model with the estimated parameters and ground truth shape $[\hat{\mathbf{J}}_t^{\text{SMPL}}, \hat{\mathbf{V}}_t] = M(\hat{\mathbf{r}}_t, \hat{\Phi}_t, \hat{\Theta}_t, \beta)$:

$$\mathcal{L}_{\text{joint}} = \|\mathbf{J}_t^{\text{SMPL}} - \hat{\mathbf{J}}_t^{\text{SMPL}}\|^2 \quad (8)$$

$$\mathcal{L}_{\text{vtx}} = \|\mathbf{V}_t - \hat{\mathbf{V}}_t\|^2 \quad \mathcal{L}_{\text{consist}} = \|\hat{\mathbf{J}}_t - \hat{\mathbf{J}}_t^{\text{SMPL}}\|^2. \quad (9)$$

The loss $\mathcal{L}_{\text{consist}}$ encourages consistency between regressed joints and those of the body model. The contact loss $\mathcal{L}_{\text{contact}} = \mathcal{L}_{\text{BCE}} + \mathcal{L}_{\text{vel}}$ contains two terms. The first supervises ground contact classification with a typical binary cross entropy; the second regularizes joint velocities to be consistent with contacts $\mathcal{L}_{\text{vel}} = \sum_j \hat{c}_t^j \|\hat{\mathbf{v}}_t\|^2$ with $\hat{\mathbf{v}}_t \in \hat{\mathbf{J}}_t$ and $\hat{c}_t^j \in \hat{\mathbf{c}}_t$ the predicted probability that joint j is in ground contact. We set $w_{\text{contact}} = 0.01$ and $w_{\text{KL}} = 4e^{-4}$.

The initial state GMM is trained separately with expectation-maximization on the same dataset used to train the CVAE.

Implementation Details. To ease learning and improve generalization, our model operates in an aligned canonical coordinate frame at each step. All networks are 4 or 5 layer MLPs with ReLU activations and group normalization [103]. To combat posterior collapse [61, 54, 88], we linearly anneal w_{KL} during training [12]. Following [54], we also use scheduled sampling [9] to enable long-term generation by making the model robust to its own errors. Additional details are available in the supplementary.

4. Test-time Motion Optimization

We next use the space of motion learned by HuMoR as a *prior* in TestOpt to recover pose and shape from noisy and partial observations while ensuring plausibility.

4.1. Optimization Variables

Given a sequence of observations $\mathbf{y}_{0:T}$, either as 2D/3D joints, 3D point clouds, or 3D keypoints, we seek the shape β and a sequence of SMPL pose parameters $(\mathbf{r}_{0:T}, \Phi_{0:T}, \Theta_{0:T})$ which describe the underlying motion being observed. We parameterize the optimized motion **using our CVAE** by the initial state $\mathbf{x}_0$ and a sequence of latent transitions $\mathbf{z}_{1:T}$. Then at T (and any intermediate steps) $\mathbf{x}_T = f(\mathbf{x}_0, \mathbf{z}_{1:T})$ is determined through model *roll-out* using the **decoder** as previously detailed. Compared to directly optimizing SMPL [6, 11, 41], this motion representation naturally encourages plausibility and is compact in the number of variables. To obtain the transformation between the canonical coordinate frame in which our CVAE is trained and the observation frame used for optimization, we

additionally optimize the ground plane of the scene $\mathbf{g} \in \mathbb{R}^3$. All together, we simultaneously optimize initial state $\mathbf{x}_0$, a sequence of latent variables $\mathbf{z}_{1:T}$, ground $\mathbf{g}$, and shape β . We assume a static camera with known intrinsics.

4.2. Objective & Optimization

The optimization objective can be formulated as a maximum a-posteriori (MAP) estimate (see supplementary), which seeks a motion that is plausible under our generative model while closely matching observations:

$$\min_{\mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta} \mathcal{E}_{\text{mot}} + \mathcal{E}_{\text{data}} + \mathcal{E}_{\text{reg}}. \quad (10)$$

We next detail each of these terms which are the motion prior, data, and regularization energies. In the following, λ are weights to determine the contribution of each term.

Motion Prior $\mathcal{E}_{\text{mot}}$. This energy measures the likelihood of the latent transitions $\mathbf{z}_{1:T}$ and initial state $\mathbf{x}_0$ under the HuMoR CVAE and GMM. It is $\mathcal{E}_{\text{mot}} = \mathcal{E}_{\text{CVAE}} + \mathcal{E}_{\text{init}}$ where

$$\begin{aligned} \mathcal{E}_{\text{CVAE}} &= -\lambda_{\text{CVAE}} \sum_{t=1}^T \log \mathcal{N}(\mathbf{z}_t; \mu_\theta(\mathbf{x}_{t-1}), \sigma_\theta(\mathbf{x}_{t-1})) \\ \mathcal{E}_{\text{init}} &= -\lambda_{\text{init}} \log \sum_{i=1}^K \gamma^i \mathcal{N}(\mathbf{x}_0; \mu_\theta^i, \sigma_\theta^i). \end{aligned} \quad (11)$$

$\mathcal{E}_{\text{CVAE}}$ uses the learned **conditional prior** and $\mathcal{E}_{\text{init}}$ uses the initial state GMM.

Data Term $\mathcal{E}_{\text{data}}$. This term is the *only* modality-dependent component of our approach, requiring different losses for different inputs: 3D joints, 2D joints, and 3D point clouds. All data losses operate on SMPL joints or mesh vertices obtained through the body model $[\mathbf{J}_t^{\text{SMPL}}, \mathbf{V}_t] = M(\mathbf{r}_t, \Phi_t, \Theta_t, \beta)$ using the current shape β along with the SMPL parameters $(\mathbf{r}_t, \Phi_t, \Theta_t)$ contained in $\mathbf{x}_t = f(\mathbf{x}_0, \mathbf{z}_{1:t})$ transformed from the canonical to observation (*i.e.* camera) frame. In the simplest case, the observations $\mathbf{y}_t$ are 3D joint positions (or keypoints with known correspondences) and our energy is

$$\mathcal{E}_{\text{data}} \triangleq \mathcal{E}_{\text{data}}^{\text{3D}} = \lambda_{\text{data}} \sum_{t=0}^T \sum_{j=1}^J \|\mathbf{p}_t^j - \mathbf{y}_t^j\|^2 \quad (12)$$

with $\mathbf{p}_t^j \in \mathbf{J}_t^{\text{SMPL}}$. For 2D joint positions, each with a detection confidence σ_t^j , we use a re-projection loss

$$\mathcal{E}_{\text{data}} \triangleq \mathcal{E}_{\text{data}}^{\text{2D}} = \lambda_{\text{data}} \sum_{t=0}^T \sum_{j=1}^J \sigma_t^j \rho(\Pi(\mathbf{p}_t^j) - \mathbf{y}_t^j) \quad (13)$$

with ρ the robust Geman-McClure function [11, 26] and Π the pinhole projection. If an estimated person segmentation mask is available, it is used to ignore spurious 2D joints.

Finally, if $\mathbf{y}_t$ is a 3D point cloud obtained from a depth map roughly masked around the person of interest, we use the mesh vertices to compute

$$\mathcal{E}_{\text{data}} \triangleq \mathcal{E}_{\text{data}}^{\text{PC3D}} = \lambda_{\text{data}} \sum_{t=0}^T \sum_{i=1}^{N_t} w_{\text{bs}} \min_{\mathbf{p}_t \in \mathbf{V}_t} \|\mathbf{p}_t - \mathbf{y}_t^i\|^2 \quad (14)$$

where w_{bs} is a robust bisquare weight [8] computed based on the Chamfer distance term.

Regularizers $\mathcal{E}_{\text{reg}}$. The additional regularization consists of four terms $\mathcal{E}_{\text{reg}} = \mathcal{E}_{\text{skel}} + \mathcal{E}_{\text{env}} + \mathcal{E}_{\text{gnd}} + \mathcal{E}_{\text{shape}}$. The first two terms encourage rolled-out motions from the CVAE to be plausible even when the initial state $\mathbf{x}_0$ is far from the optimum (*i.e.* early in optimization). The skeleton consistency term uses the joints $\mathbf{J}_t$ directly *predicted* by the decoder during rollout along with the SMPL joints:

$$\mathcal{E}_{\text{skel}} = \sum_{t=1}^T \left(\lambda_c \sum_{j=1}^J \|\mathbf{p}_t^j - \mathbf{p}_t^{j, \text{pred}}\|^2 + \lambda_b \sum_{i=1}^B (l_t^i - l_{t-1}^i)^2 \right)$$

with $\mathbf{p}_t^j \in \mathbf{J}_t^{\text{SMPL}}$ and $\mathbf{p}_t^{j, \text{pred}} \in \mathbf{J}_t$. The second summation uses bone lengths l computed from $\mathbf{J}_t$ at each step. The second regularizer $\mathcal{E}_{\text{env}}$ ensures consistency between predicted CVAE contacts, the motion, and the environment:

$$\mathcal{E}_{\text{env}} = \sum_{t=1}^T \sum_{j=1}^J \lambda_{\text{cv}} c_t^j \|\mathbf{p}_t^j - \mathbf{p}_{t-1}^j\|^2 + \lambda_{\text{ch}} c_t^j \max(|\mathbf{p}_{z,t}^j| - \delta, 0)$$

where $\mathbf{p}_t^j \in \mathbf{J}_t^{\text{SMPL}}$ and c_t^j is the contact probability output from the model for joint j . The contact height term weighted by λ_{ch} ensures the z -component of contacting joints are within δ of the floor in the canonical frame.

The final two regularizers are priors on the ground and shape. We assume the ground should stay close to initialization $\mathcal{E}_{\text{gnd}} = \lambda_{\text{gnd}} \|\mathbf{g} - \mathbf{g}^{\text{init}}\|^2$. Finally, β should stay near the neutral zero vector similar to [34, 73]: $\mathcal{E}_{\text{shape}} = \lambda_{\text{shape}} \|\beta\|^2$.

Initialization & Optimization. We initialize the temporal SMPL parameters $\mathbf{r}_{0:T}, \Phi_{0:T}, \Theta_{0:T}$ and shape β with an initialization optimization using $\mathcal{E}_{\text{data}}$ and $\mathcal{E}_{\text{shape}}$ along with two additional regularization terms. $\mathcal{E}_{\text{pose}} = \sum_t \|\mathbf{z}_t^{\text{pose}}\|^2$ is a pose prior where $\mathbf{z}_t^{\text{pose}} \in \mathbb{R}^{32}$ is the body joint angles represented in the latent space of the VPoser model [73, 34]. The smoothness term $\mathcal{E}_{\text{smooth}} = \sum_{t=1}^T \sum_{j=1}^J \|\mathbf{p}_t^j - \mathbf{p}_{t-1}^j\|^2$ with $\mathbf{p}_t^j \in \mathbf{J}_t^{\text{SMPL}}$ smooths 3D joint positions over time. Afterwards, the initial latent sequence $\mathbf{z}_{1:T}^{\text{init}}$ is computed through inference with the CVAE **encoder**. Our optimization is implemented in PyTorch [72] using L-BFGS and *autograd*; with batching, a 3s RGB video takes about 5.5 min to fit. We provide further details in the supplementary material.

5. Experimental Results

We evaluate HuMoR on (i) generative sampling tasks and (ii) as a prior in TestOpt to estimate motion from 3D

Model	Future Prediction			Diversity APD $\uparrow$
	Contact $\uparrow$	ADE $\downarrow$	FDE $\downarrow$	
MVAE [54]	-	25.8	50.6	85.4
HuMoR	0.88	21.5	42.1	94.9
HuMoR (Qual)	0.88	22.0	46.3	100.0

Table 1: (Left) Future prediction accuracy for 2s AMASS sequences. Contact classification accuracy, average displacement error (*cm*), and final displacement error (*cm*) are reported. (Right) Sampling diversity over 5s rollouts measured by average pairwise distance (*cm*).

and RGB(-D) inputs. We encourage viewing the **supplementary videos** to appreciate the qualitative improvement of our approach. Additional dataset and experiment details are available in the supplementary document.

5.1. Datasets

AMASS [63] is a large motion capture database containing diverse motions and body shapes on the SMPL body model. We sub-sample the dataset to 30 Hz and use the recommended training split to train the CVAE and initial state GMM in HuMoR. We evaluate on the held out Transitions and HumanEva [87] subsets (Sec. 5.3 and 5.4).

i3DB [69] contains RGB videos of person-scene interactions involving medium to heavy occlusions. It provides annotated 3D joint positions and a primitive 3D scene reconstruction which we use to fit a ground plane for computing plausibility metrics. We run off-the-shelf 2D pose estimation [16], person segmentation [18], and plane detection [55] models to obtain inputs for our optimization.

PROX [34] contains RGB-D videos of people interacting with indoor environments. We use a subset of the qualitative data to evaluate plausibility metrics using a floor plane fit to the provided ground truth scene mesh. We obtain 2D pose, person masks, and ground plane initialization in the same way as done for i3DB.

5.2. Baselines and Evaluation Metrics

Motion Prior Baselines. We ablate the proposed CVAE to analyze its core components: *No Delta* directly predicts the next state from the decoder rather than the change in state, *No Contacts* does not classify ground contacts, *No $\mathcal{L}_{SMPL}$* does not use SMPL regularization in training, and *Standard Prior* uses $\mathcal{N}(\mathbf{0}, \mathbf{I})$ rather than our learned **conditional prior**. All of these ablated together recovers MVAE [54].

Motion Estimation Baselines. *VPoser-t* is the initialization phase of our optimization. It uses VPoser [73] and 3D joint smoothing similar to previous works [6, 41, 109]. *PROX-(RGB/D)* [34] are optimization-based methods which operate on individual frames of RGB and RGB-D videos, respectively. Both assume the full scene mesh is given to en-

force contact and penetration constraints. *VIBE* [47] is a recent learned method to recover shape and pose from video.

Error Metrics. 3D positional errors are measured on joints, keypoints, or mesh vertices (**Vtx**) and compute *global* mean per-point position error unless otherwise specified. We report positional errors for all (**All**), occluded (**Occ**), and visible (**Vis**) observations separately. Finally, we report binary classification accuracy of the 8 person-ground contacts (**Contact**) predicted by HuMoR.

Plausibility Metrics. We use additional metrics to measure qualitative motion characteristics that joint errors cannot capture. Smoothness is evaluated by mean per-joint accelerations (**Accel**) [43]. Another important indicator of plausibility is ground penetration [79]. We use the true ground plane to compute the frequency (**Freq**) of foot-floor penetrations: the fraction of frames for both the left and right toe joints that penetrate more than a threshold. We measure frequency at 0, 3, 6, 9, 12, and 15 *cm* thresholds and report the mean. We also report mean penetration distance (**Dist**), where non-penetrating frames contribute a distance of 0 to make values comparable across differing frequencies.

5.3. Generative Model Evaluation

We first evaluate HuMoR as a standalone generative model and show improved generalization to unseen motions and bodies compared to MVAE for two common tasks (see Table 1): future prediction and diverse sampling. We use 2s AMASS sequences and start generation from the first step. Results are shown for *HuMoR* and a modified *HuMoR (Qual)* that uses $\mathbf{J}^{\text{SMPL}}$ as input to each step during rollout instead of $\mathbf{J}$, thereby enforcing skeleton consistency. This version produces *qualitatively* superior results for generation, but is too expensive to use during TestOpt.

For prediction, we report average displacement error (**ADE**) and final displacement error (**FDE**) [106], which measure mean joint errors over all steps and at the final step, respectively. We sample 50 2s motions for each initial state and the one with lowest **ADE** is considered the prediction. For diversity, we sample 50 5s motions and compute the average pairwise distance (**APD**) [4], *i.e.* the mean joint distance between all pairs of samples.

As seen in Tab. 1, the base MVAE [54] does not generalize well when trained on the large AMASS dataset; our proposed CVAE improves both the accuracy and diversity of samples. *HuMoR (Qual)* hinders prediction accuracy, but gives better diversity and visual quality (see supplement).

5.4. Estimation from 3D Observations

Next, we show that HuMoR also generalizes better when used in TestOpt for fitting to 3D data, and that using a motion prior is crucial to plausibly handling occlusions. 3s AMASS sequences are used to demonstrate key abilities: (i) fitting to partial data and (ii) denoising. For the former,

Method	Input	Positional Error			Joints Legs	Mesh Vtx	Contact	Ground Pen		
		Vis	Occ	All				Accel	Freq	Dist
VPoser-t	Occ Keypoints	0.67	20.76	9.22	21.08	7.95	-	5.71	16.77%	2.28
MVAE [54]	Occ Keypoints	2.39	19.15	9.52	16.86	8.90	-	7.12	3.15%	0.30
HuMoR (Ours)	Occ Keypoints	1.46	17.40	8.24	15.42	7.56	0.89	5.38	3.31%	0.26
VPoser-t	Noisy Joints	-	-	3.67	4.47	4.98	-	4.61	1.35%	0.07
MVAE [54]	Noisy Joints	-	-	2.68	3.21	4.42	-	6.5	1.75%	0.11
HuMoR (Ours)	Noisy Joints	-	-	2.27	2.61	3.55	0.97	5.23	1.18%	0.05

Table 2: Motion and shape estimation from 3D observations: partially occluded keypoints (top) and noisy joints (bottom). *Positional Error (cm)* is reported w.r.t. the input modality. Acceleration is m/s^2 and penetration distance in cm .

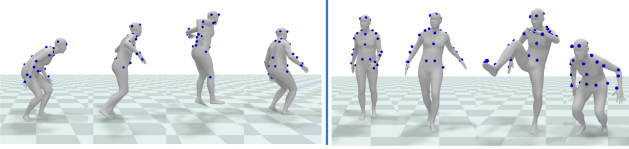


Figure 3: Fitting to partial 3D **keypoints**. HuMoR captures non-periodic motions like jumping, crouching, and kicking.

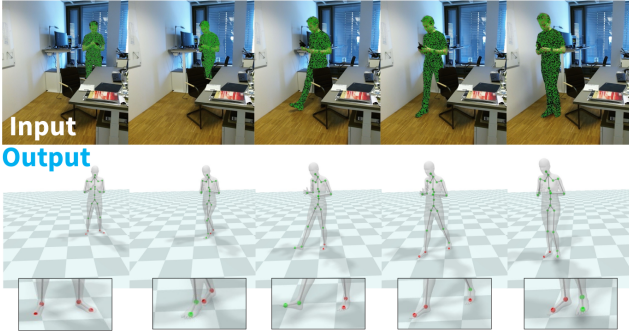


Figure 4: From RGB-D (top) TestOpt with HuMoR outputs 3D motion, the ground plane, and **contacts** (bottom).

TestOpt fits to 43 keypoints on the body that resemble motion capture markers; keypoints that fall below $0.9m$ at each timestep are “occluded”, leaving the legs unobservable at most steps. For denoising, Gaussian noise with $4cm$ standard deviation is added to 3D joint position observations.

Tab. 2 compares to *VPoser-t* and to using *MVAE* as the motion prior during optimization rather than HuMoR. We report leg joint errors (toes, ankles, and knees), which are often occluded, separately. The right side of the table reports plausibility metrics. HuMoR gives more accurate poses, especially for occluded keypoints and leg joints. It also estimates smoother motions with fewer and less severe ground penetrations. For denoising, *VPoser-t* oversmooths which gives the lowest acceleration but least accurate motion. TestOpt with HuMoR gives inherently smooth results while still allowing for necessarily large accelerations to fit dynamic observations. Notably, HuMoR predicts person-ground contact with 97% accuracy even under severe noise. Qualitative results are shown in Fig. 1 and Fig. 3.

5.5. Estimation from RGB(-D) Observations

Finally, we show that TestOpt with HuMoR can be applied to real-world RGB and RGB-D observations, and outperforms baselines on positional and plausibility metrics especially from partial and noisy data. We use $3s$ (90 frame) clips from i3DB [69] and PROX [34]. Tab. 3 shows results on i3DB which affords quantitative 3D joint evaluation. The top half compares to baseline estimation methods; the bottom uses ablations of HuMoR in TestOpt rather than the full model. Mean per-joint position errors are reported for **global** joint positions and after root alignment.

As seen in Tab. 3, *VIBE* gives locally accurate predictions for visible joints, but large global errors and unrealistic accelerations due to occlusions and temporal inconsistency (see Fig. 5). *VPoser-t* gives reasonable global errors, but suffers frequent penetrations as shown for sitting in Fig. 5. Using *MVAE* or ablations of HuMoR as the motion prior in TestOpt fails to effectively generalize to real-world data and performs worse than the full model. The conditional prior and $\mathcal{L}_{SMPL}$ have the largest impact, while performance even without using contacts still outperforms the baselines.

The top half of Tab. 4 evaluates plausibility on additional RGB results from PROX compared to *VIBE* and *PROX-RGB*. Since *PROX-RGB* uses the scene mesh as input to enforce environment constraints, it is a very strong baseline and its performance on penetration metrics is expectedly good. HuMoR comparatively increases penetration frequency since it only gets a rough ground plane as initialization, but gives much smoother motions.

The bottom half of Tab. 4 shows results fitting to RGB-D for the same PROX data, which uses both $\mathcal{E}_{data}^{2D}$ and $\mathcal{E}_{data}^{PC3D}$ in TestOpt. This improves performance using HuMoR, slightly outperforming *PROX-D* which is less robust to issues with 2D joint detections and 3D point noise causing large errors. Qualitative examples are in Fig. 1 and Fig. 4.

Thanks to the generalizability of HuMoR, TestOpt is also effective in recovering very dynamic motions like dancing from RGB video when the full body is visible (see supplementary material for examples).

Method	Global Joint Error				Root-Aligned Joint Error				Ground Pen		
	Vis	Occ	All	Legs	Vis	Occ	All	Legs	Accel	Freq	Dist
VIBE [47]	90.05	192.55	116.46	121.61	12.06	23.78	15.08	21.65	243.36	7.98%	3.01
VPoser-t	28.33	40.97	31.59	35.06	12.77	26.48	16.31	25.60	4.46	9.28%	2.42
MVAE [54]	37.54	50.63	40.91	44.42	16.00	28.32	19.17	26.63	4.96	7.43%	1.55
No Delta	27.55	35.59	29.62	32.14	11.92	23.10	14.80	21.65	3.05	2.84%	0.58
No Contacts	26.65	39.21	29.89	35.73	12.24	23.36	15.11	22.25	2.43	5.59%	1.70
No $\mathcal{L}_{\text{SMPL}}$	31.09	43.67	34.33	36.84	12.81	25.47	16.07	23.54	3.21	4.12%	1.31
Standard Prior	77.60	146.76	95.42	99.01	18.67	39.40	24.01	34.02	5.98	8.30%	6.47
HuMoR (Ours)	26.00	34.36	28.15	31.26	12.02	21.70	14.51	20.74	2.43	2.12%	0.68

Table 3: Motion and shape from RGB video (*i.e.* 2D joints) on i3DB [69]. Joint errors are in *cm* and acceleration is m/s^2 . Top shows results from motion estimation baselines while bottom uses ablations of HuMoR during optimization.

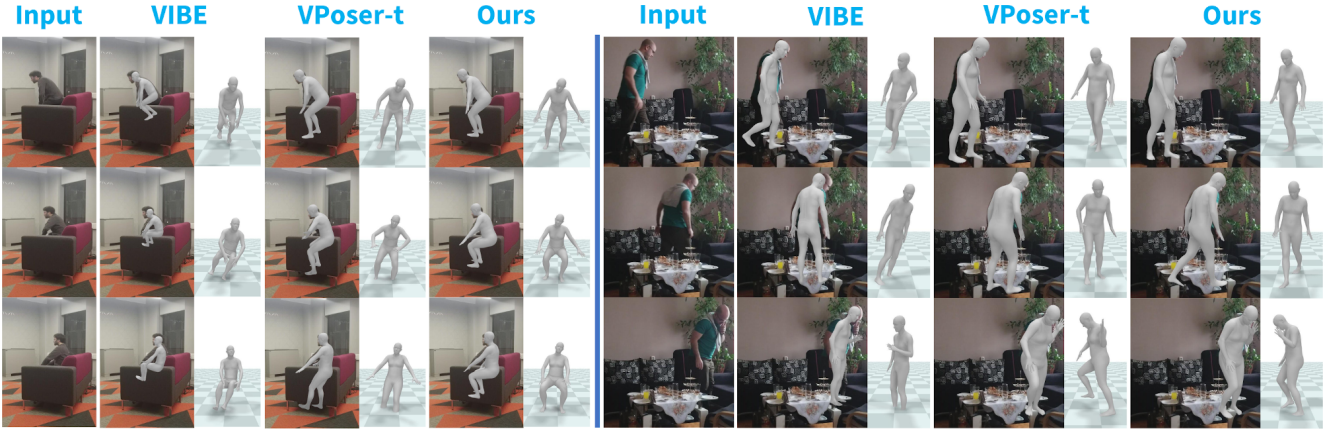


Figure 5: Qualitative comparison for fitting to RGB video (*i.e.* 2D joints) from i3DB [69]. Optimization using HuMoR (Ours) outputs natural and plausible sitting and walking motions under heavy occlusions compared to baseline approaches.

Method	Input	Ground Pen		
		Accel	Freq	Dist
VIBE [47]	RGB	86.06	23.46%	4.71
PROX-RGB [34]	RGB	196.07	2.55%	0.32
VPoser-t	RGB	3.14	13.38%	2.82
HuMoR (Ours)	RGB	1.73	9.99%	1.56
PROX-D [34]	RGB-D	46.59	8.95%	1.19
VPoser-t	RGB-D	3.27	10.66%	2.18
HuMoR (Ours)	RGB-D	1.61	5.19%	0.85

Table 4: Plausibility evaluation on videos in PROX [34]. Acceleration is m/s^2 and penetration distance in *cm*.

6. Discussion

We have introduced HuMoR, a learned generative model of 3D human motion leveraged during test-time optimization to robustly recover pose and shape from 3D, RGB, and RGB-D observations. We have demonstrated that the key components of our model enable generalization to novel motions and body shapes for both generative tasks and downstream optimization. Compared to strong learning and

optimization-based baselines, HuMoR excels at estimating plausible motion under heavy occlusions, and simultaneously produces consistent ground plane and contact outputs.

Limitations & Future Work. HuMoR leaves ample room for future studies. The static camera and ground plane assumptions are reasonable for indoor scenes but true *in-the-wild* operation demands methods handling dynamic cameras and complex terrain. Our rather simplistic contact model should be upgraded to capture scene-person interactions for improved motion and scene perception. Lastly, we plan to *learn* motion estimation directly from partial observations which will be faster than TestOpt and enable *sampling* multiple plausible motions rather than relying on a single local minimum.

Acknowledgments. This work was supported by the Toyota Research Institute (“TRI”) under the University 2.0 program, grants from the Samsung GRO program and the Ford-Stanford Alliance, a Vannevar Bush faculty fellowship, NSF grant IIS-1763268, and NSF grant CNS-2038897. TRI provided funds to assist the authors with their research but this article solely reflects the opinions and conclusions of its authors and not TRI or any other Toyota entity.

References

- [1] Advanced Computing Center for the Arts and Design. AC-CAD MoCap Dataset. 16
- [2] Ijaz Akhter and Michael J. Black. Pose-conditioned joint angle limits for 3D human pose reconstruction. In *2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 1446–1455, June 2015. 16
- [3] Emre Aksan, Manuel Kaufmann, and Otmar Hilliges. Structured prediction helps 3d human motion modelling. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 7144–7153, 2019. 3
- [4] Sadegh Aliakbarian, Fatemeh Sadat Saleh, Mathieu Salzmann, Lars Petersson, and Stephen Gould. A stochastic conditioning scheme for diverse human motion prediction. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 5223–5232, 2020. 6
- [5] Andreas Aristidou, Ariel Shamir, and Yiorgos Chrysanthou. Digital dance ethnography: Organizing large dance collections. *J. Comput. Cult. Herit.*, 12(4), Nov. 2019. 20, 25
- [6] Anurag Arnab, Carl Doersch, and Andrew Zisserman. Exploiting temporal context for 3d human pose estimation in the wild. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 3395–3404, 2019. 2, 4, 6
- [7] Andreas Baak, Meinard Müller, Gaurav Bharaj, Hans-Peter Seidel, and Christian Theobalt. A data-driven approach for real-time full body pose reconstruction from a depth camera. In *Consumer Depth Cameras for Computer Vision*, pages 71–98. Springer, 2013. 2
- [8] Albert E Beaton and John W Tukey. The fitting of power series, meaning polynomials, illustrated on band-spectroscopic data. *Technometrics*, 16(2):147–185, 1974. 5, 17
- [9] Samy Bengio, Oriol Vinyals, Navdeep Jaitly, and Noam Shazeer. Scheduled sampling for sequence prediction with recurrent neural networks. In *Proceedings of the 28th International Conference on Neural Information Processing Systems - Volume 1, NIPS’15*, page 1171–1179, Cambridge, MA, USA, 2015. MIT Press. 4
- [10] Benjamin Biggs, David Novotny, Sebastien Ehrhardt, Hanbyul Joo, Ben Graham, and Andrea Vedaldi. 3d multi-bodies: Fitting sets of plausible 3d human models to ambiguous image data. *Advances in Neural Information Processing Systems (NeurIPS)*, 33, 2020. 2
- [11] Federica Bogo, Angjoo Kanazawa, Christoph Lassner, Peter Gehler, Javier Romero, and Michael J. Black. Keep it SMPL: Automatic estimation of 3D human pose and shape from a single image. In *Computer Vision – ECCV 2016*, Lecture Notes in Computer Science. Springer International Publishing, Oct. 2016. 2, 4, 5
- [12] Samuel R Bowman, Luke Vilnis, Oriol Vinyals, Andrew M Dai, Rafal Jozefowicz, and Samy Bengio. Generating sentences from a continuous space. In *20th SIGNLL Conference on Computational Natural Language Learning, CoNLL 2016*, pages 10–21. Association for Computational Linguistics (ACL), 2016. 4
- [13] Matthew Brand and Aaron Hertzmann. Style machines. In *ACM SIGGRAPH*, pages 183–192, July 2000. 2
- [14] Marcus A. Brubaker, David J. Fleet, and Aaron Hertzmann. Physics-based person tracking using the anthropomorphic walker. *IJCV*, (1), 2010. 1, 2
- [15] Marcus A. Brubaker, Leonid Sigal, and David J. Fleet. Estimating contact dynamics. In *Proc. ICCV*, 2009. 1, 2
- [16] Z. Cao, G. Hidalgo Martinez, T. Simon, S. Wei, and Y. A. Sheikh. Openpose: Realtime multi-person 2d pose estimation using part affinity fields. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2019. 1, 2, 6, 19
- [17] Carnegie Mellon University. CMU MoCap Dataset. 16
- [18] Liang-Chieh Chen, George Papandreou, Florian Schroff, and Hartwig Adam. Rethinking atrous convolution for semantic image segmentation. *arXiv preprint arXiv:1706.05587*, 2017. 6, 19, 20
- [19] Ricky TQ Chen, Yulia Rubanova, Jesse Bettencourt, and David Duvenaud. Neural ordinary differential equations. In *Proceedings of the 32nd International Conference on Neural Information Processing Systems*, pages 6572–6583, 2018. 14
- [20] Yixin Chen, Siyuan Huang, Tao Yuan, Siyuan Qi, Yixin Zhu, and Song-Chun Zhu. Holistic++ scene understanding: Single-view 3d holistic scene parsing and human pose estimation with human-object interaction and physical commonsense. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 8648–8657, 2019. 1
- [21] Vasileios Choutas, Georgios Pavlakos, Timo Bolkart, Dimitrios Tzionas, and Michael J Black. Monocular expressive body regression through body-driven attention. In *European Conference on Computer Vision*, pages 20–40. Springer, 2020. 2
- [22] Enric Corona, Albert Pumarola, Guillem Alenya, and Francesc Moreno-Noguer. Context-aware human motion prediction. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, June 2020. 3
- [23] Carl Doersch and Andrew Zisserman. Sim2real transfer learning for 3d human pose estimation: motion to the rescue. *Advances in Neural Information Processing Systems (NeurIPS)*, 32:12949–12961, 2019. 2
- [24] Ahmed Elgammal and Chan-Su Lee. Separating style and content on a nonlinear manifold. In *IEEE Conf. Comp. Vis. and Pattern Recognition*, pages 478–485, 2004. Vol. 1. 2
- [25] Varun Ganapathi, Christian Plagemann, Daphne Koller, and Sebastian Thrun. Real time motion capture using a single time-of-flight camera. In *2010 IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pages 755–762. IEEE, 2010. 2
- [26] S. Geman and D. McClure. Statistical methods for tomographic image reconstruction. *Bulletin of the International Statistical Institute*, 52(4):5–21, 1987. 5, 17
- [27] Georgios Georgakis, Ren Li, Srikrishna Karanam, Terrence Chen, Jana Košecká, and Ziyang Wu. Hierarchical kinematic human mesh recovery. In *European Conference on Computer Vision*, pages 768–784. Springer, 2020. 2

- [28] Saeed Ghorbani, Kimia Mahdavian, Anne Thaler, Konrad Kording, Douglas James Cook, Gunnar Blohm, and Nikolaus F. Troje. MoVi: A large multipurpose motion and video dataset, 2020. 16
- [29] Saeed Ghorbani, Calden Wloka, Ali Etemad, Marcus A Brubaker, and Nikolaus F Troje. Probabilistic character motion synthesis using a hierarchical deep latent variable model. In *Computer Graphics Forum*, volume 39. Wiley Online Library, 2020. 2, 14
- [30] Riza Alp Güler and Iasonas Kokkinos. Holopose: Holistic 3d human reconstruction in-the-wild. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 10884–10894, 2019. 2
- [31] Rıza Alp Güler, Natalia Neverova, and Iasonas Kokkinos. Densepose: Dense human pose estimation in the wild. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 7297–7306, 2018. 2
- [32] Marc Habermann, Weipeng Xu, Michael Zollhofer, Gerard Pons-Moll, and Christian Theobalt. Deepcap: Monocular human performance capture using weak supervision. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 5052–5063, 2020. 1
- [33] Ikhsanul Habibie, Daniel Holden, Jonathan Schwarz, Joe Yearsley, and Taku Komura. A recurrent variational autoencoder for human motion synthesis. In *28th British Machine Vision Conference*, 2017. 2
- [34] Mohamed Hassan, Vasileios Choutas, Dimitrios Tzionas, and Michael J. Black. Resolving 3D human pose ambiguities with 3D scene constraints. In *International Conference on Computer Vision*, pages 2282–2292, 2019. 2, 5, 6, 7, 8, 19
- [35] Mohamed Hassan, Partha Ghosh, Joachim Tesch, Dimitrios Tzionas, and Michael J Black. Populating 3d scenes by learning human-scene interaction. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 14708–14718, 2021. 2
- [36] Gustav Eje Henter, Simon Alexanderson, and Jonas Beskow. Moglow: Probabilistic and controllable motion synthesis using normalising flows. *ACM Transactions on Graphics (TOG)*, 39(6):1–14, 2020. 2
- [37] I. Higgins, Loïc Matthey, A. Pal, C. Burgess, Xavier Glorot, M. Botvinick, S. Mohamed, and Alexander Lerchner. beta-vae: Learning basic visual concepts with a constrained variational framework. In *International Conference on Learning Representations*, 2017. 14
- [38] Daniel Holden, Taku Komura, and Jun Saito. Phase-functioned neural networks for character control. *ACM Transactions on Graphics (TOG)*, 36(4):1–13, 2017. 2
- [39] Nicholas R. Howe, Michael E. Leventon, and William T. Freeman. Bayesian reconstruction of 3D human motion from single-camera video. In *Advances in Neural Information Processing Systems 12*, pages 820–826, 2000. 2
- [40] Ludovic Hoyet, Kenneth Ryall, Rachel McDonnell, and Carol O’Sullivan. Sleight of hand: Perception of finger motion from reduced marker sets. In *Proceedings of the ACM SIGGRAPH Symposium on Interactive 3D Graphics and Games*, I3D ’12, page 79–86, 2012. 16
- [41] Yinghao Huang, Federica Bogo, Christoph Lassner, Angjoo Kanazawa, Peter V Gehler, Javier Romero, Ijaz Akhter, and Michael J Black. Towards accurate marker-less human shape and pose estimation over time. In *2017 International Conference on 3D Vision*, pages 421–430. IEEE, 2017. 2, 4, 6
- [42] Angjoo Kanazawa, Michael J. Black, David W. Jacobs, and Jitendra Malik. End-to-end recovery of human shape and pose. In *Computer Vision and Pattern Recognition*, 2018. 2
- [43] Angjoo Kanazawa, Jason Y Zhang, Panna Felsen, and Jitendra Malik. Learning 3d human dynamics from video. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2019. 2, 6
- [44] Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. In *International Conference on Learning Representations*, 2015. 16
- [45] Diederik P Kingma and Max Welling. Auto-encoding variational bayes. In *Proceedings of the International Conference on Learning Representations (ICLR)*, 2014. 4
- [46] Ivan Kobyzev, Simon J.D. Prince, and Marcus A. Brubaker. Normalizing flows: An introduction and review of current methods. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2020. 14
- [47] Muhammed Kocabas, Nikos Athanasiou, and Michael J. Black. Vibe: Video inference for human body pose and shape estimation. In *The IEEE Conference on Computer Vision and Pattern Recognition*, June 2020. 1, 2, 6, 8
- [48] Muhammed Kocabas, Chun-Hao P Huang, Otmar Hilliges, and Michael J Black. Pare: Part attention regressor for 3d human body estimation. *arXiv preprint arXiv:2104.08527*, 2021. 2
- [49] Nikos Kolotouros, Georgios Pavlakos, Michael J. Black, and Kostas Daniilidis. Learning to reconstruct 3D human pose and shape via model-fitting in the loop. In *Proceedings International Conference on Computer Vision (ICCV)*, pages 2252–2261. IEEE, Oct. 2019. ISSN: 2380-7504. 2
- [50] Lucas Kovar, Michael Gleicher, and Frédéric Pighin. Motion graphs. In *ACM Transactions on Graphics 21(3), Proc. SIGGRAPH*, pages 473–482, July 2002. 2
- [51] Christoph Lassner, Javier Romero, Martin Kiefel, Federica Bogo, Michael J Black, and Peter V Gehler. Unite the people: Closing the loop between 3d and 2d human representations. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 6050–6059, 2017. 2
- [52] Yan Li, Tianshu Wang, and Heung-Yeung Shum. Motion texture: A two-level statistical model for character motion synthesis. In *ACM Transactions on Graphics 21(3), Proc. SIGGRAPH*, pages 465–472, July 2002. 2
- [53] Zongmian Li, Jiri Sedlar, Justin Carpentier, Ivan Laptev, Nicolas Mansard, and Josef Sivic. Estimating 3d motion and forces of person-object interactions from monocular video. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2019. 1, 2
- [54] Hung Yu Ling, Fabio Zinno, George Cheng, and Michiel van de Panne. Character controllers using motion vaes. In *ACM Transactions on Graphics (Proceedings of ACM SIGGRAPH)*, volume 39. ACM, 2020. 2, 3, 4, 6, 7, 8, 14, 16, 17, 19

- [55] Chen Liu, Kihwan Kim, Jinwei Gu, Yasutaka Furukawa, and Jan Kautz. Planercnn: 3d plane detection and reconstruction from a single image. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 4450–4459, 2019. 6, 18, 19
- [56] C. Karen Liu, Aaron Hertzmann, and Zoran Popović. Learning physics-based motion style with nonlinear inverse optimization. *ACM Trans. Graph.*, 2005. 2
- [57] Miao Liu, Dexin Yang, Yan Zhang, Zhaopeng Cui, James M Rehg, and Siyu Tang. 4d human body capture from egocentric video via 3d scene grounding. *arXiv preprint arXiv:2011.13341*, 2020. 2
- [58] Matthew Loper, Naureen Mahmood, and Michael J. Black. MoSh: Motion and Shape Capture from Sparse Markers. *ACM Trans. Graph.*, 33(6), Nov. 2014. 16
- [59] Matthew Loper, Naureen Mahmood, Javier Romero, Gerard Pons-Moll, and Michael J. Black. SMPL: A skinned multi-person linear model. *ACM Trans. Graphics (Proc. SIGGRAPH Asia)*, 34(6):248:1–248:16, Oct. 2015. 2, 3
- [60] Eyes JAPAN Co. Ltd. Eyes Japan MoCap Dataset. 16
- [61] James Lucas, George Tucker, Roger B Grosse, and Mohammad Norouzi. Don't blame the elbow! a linear vae perspective on posterior collapse. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d'Alché-Buc, E. Fox, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 32. Curran Associates, Inc., 2019. 4, 17
- [62] Zhengyi Luo, S Alireza Golestaneh, and Kris M Kitani. 3d human motion estimation via motion compression and refinement. In *Proceedings of the Asian Conference on Computer Vision*, 2020. 2
- [63] Naureen Mahmood, Nima Ghorbani, Nikolaus F. Troje, Gerard Pons-Moll, and Michael J. Black. AMASS: Archive of motion capture as surface shapes. In *International Conference on Computer Vision*, 2019. 2, 6, 14, 15, 16, 19, 21
- [64] Naureen Mahmood, Nima Ghorbani, Nikolaus F. Troje, Gerard Pons-Moll, and Michael J. Black. AMASS: Archive of motion capture as surface shapes. In *2019 IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 5441–5450, Oct. 2019. 20, 25
- [65] C. Mandery, Ö. Terlemez, M. Do, N. Vahrenkamp, and T. Asfour. The KIT whole-body human motion database. In *2015 International Conference on Advanced Robotics (ICAR)*, pages 329–336, July 2015. 16
- [66] Dushyant Mehta, Oleksandr Sotnychenko, Franziska Mueller, Weipeng Xu, Mohamed Elgharib, Pascal Fua, Hans-Peter Seidel, Helge Rhodin, Gerard Pons-Moll, and Christian Theobalt. Xnect: Real-time multi-person 3d motion capture with a single rgb camera. *ACM Transactions on Graphics (TOG)*, 39(4):82–1, 2020. 1, 2
- [67] Dushyant Mehta, Srinath Sridhar, Oleksandr Sotnychenko, Helge Rhodin, Mohammad Shafiei, Hans-Peter Seidel, Weipeng Xu, Dan Casas, and Christian Theobalt. Vnect: Real-time 3d human pose estimation with a single rgb camera. *ACM Transactions on Graphics (TOG)*, 36(4):1–14, 2017. 1, 2
- [68] Mehdi Mirza and Simon Osindero. Conditional generative adversarial nets. *arXiv preprint arXiv:1411.1784*, 2014. 14
- [69] Aron Monszpart, Paul Guerrero, Duygu Ceylan, Ersin Yumer, and Niloy J. Mitra. iMapper: Interaction-guided scene mapping from monocular videos. *ACM SIGGRAPH*, 2019. 2, 6, 7, 8, 19, 21
- [70] M. Müller, T. Röder, M. Clausen, B. Eberhardt, B. Krüger, and A. Weber. Documentation mocap database HDM05. Technical Report CG-2007-2, Universität Bonn, 2007. 16
- [71] Dirk Ormoneit, Hedvig Sidenbladh, Michael J. Black, and Trevor Hastie. Learning and tracking cyclic human motion. In *Advances in Neural Information Processing Systems 13*, pages 894–900, 2001. 1, 2
- [72] Adam Paszke, Sam Gross, Soumith Chintala, Gregory Chanan, Edward Yang, Zachary DeVito, Zeming Lin, Alban Desmaison, Luca Antiga, and Adam Lerer. Automatic differentiation in pytorch. In *Advances in Neural Information Processing Systems*, 2017. 5, 18
- [73] Georgios Pavlakos, Vasileios Choutas, Nima Ghorbani, Timo Bolkart, Ahmed A. A. Osman, Dimitrios Tzionas, and Michael J. Black. Expressive body capture: 3D hands, face, and body from a single image. In *Proceedings IEEE Conf. on Computer Vision and Pattern Recognition (CVPR)*, pages 10975–10985, 2019. 1, 2, 5, 6, 17, 19
- [74] Georgios Pavlakos, Luyang Zhu, Xiaowei Zhou, and Kostas Daniilidis. Learning to estimate 3d human pose and shape from a single color image. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 459–468, 2018. 2
- [75] Dario Pavllo, Christoph Feichtenhofer, Michael Auli, and David Grangier. Modeling human motion with quaternion-based neural networks. *International Journal of Computer Vision*, pages 1–18, 2019. 3
- [76] Dario Pavllo, Christoph Feichtenhofer, David Grangier, and Michael Auli. 3d human pose estimation in video with temporal convolutions and semi-supervised training. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 7753–7762, 2019. 2
- [77] Vladimir Pavlović, James M. Rehg, and John MacCormick. Learning switching linear models of human motion. In *Advances in Neural Information Processing Systems 13*, pages 981–987, 2001. 1, 2
- [78] Xue Bin Peng, Angjoo Kanazawa, Jitendra Malik, Pieter Abbeel, and Sergey Levine. Sfv: Reinforcement learning of physical skills from videos. *ACM Trans. Graph.*, 2018. 1, 2
- [79] Davis Rempe, Leonidas J. Guibas, Aaron Hertzmann, Bryan Russell, Ruben Villegas, and Jimei Yang. Contact and human dynamics from monocular video. In *Proceedings of the European Conference on Computer Vision (ECCV)*, 2020. 1, 2, 6, 15
- [80] Chris Rockwell and David F Fouhey. Full-body awareness from partial observations. In *Computer Vision—ECCV 2020: 16th European Conference, Glasgow, UK, August 23–28, 2020, Proceedings, Part XVII 16*, pages 522–539, 2020. 2
- [81] Javier Romero, Dimitrios Tzionas, and Michael J. Black. Embodied hands: Modeling and capturing hands and bodies together. *ACM Transactions on Graphics, (Proc. SIGGRAPH Asia)*, 36(6), Nov. 2017. 3, 16, 18

- [82] Charles Rose, Michael F. Cohen, and Bobby Bodenheimer. Verbs and adverbs: Multidimensional motion interpolation. *IEEE Computer Graphics and Applications*, 18(5):32–40, 1998. [2](#)
- [83] Shunsuke Saito, Zeng Huang, Ryota Natsume, Shigeo Morishima, Angjoo Kanazawa, and Hao Li. Pifu: Pixel-aligned implicit function for high-resolution clothed human digitization. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV)*, October 2019. [1](#)
- [84] Mingyi Shi, Kfir Aberman, Andreas Aristidou, Taku Komura, Dani Lischinski, Daniel Cohen-Or, and Baoquan Chen. Motionet: 3d human motion reconstruction from monocular video with skeleton consistency. *ACM Transactions on Graphics (TOG)*, 40(1):1–15, 2020. [2](#)
- [85] Soshi Shimada, Vladislav Golyanik, Weipeng Xu, and Christian Theobalt. Physcap: Physically plausible monocular 3d motion capture in real time. *ACM Trans. Graph.*, 39(6), Nov. 2020. [1](#), [2](#)
- [86] Hedvig Sidenbladh, Michael J. Black, and David J. Fleet. Stochastic tracking of 3D human figures using 2D image motion. In *ECCV*, pages 702–718, 2000. Part II. [1](#), [2](#)
- [87] Leonid Sigal, Alexandru O Balan, and Michael J Black. Humaneva: Synchronized video and motion capture dataset and baseline algorithm for evaluation of articulated human motion. *International journal of computer vision*, 87(1-2):4, 2010. [6](#), [16](#), [19](#)
- [88] Kihyuk Sohn, Honglak Lee, and Xinchen Yan. Learning structured output representation using deep conditional generative models. In C. Cortes, N. Lawrence, D. Lee, M. Sugiyama, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 28, pages 3483–3491. Curran Associates, Inc., 2015. [2](#), [3](#), [4](#), [17](#)
- [89] Sebastian Starke, He Zhang, Taku Komura, and Jun Saito. Neural state machine for character-scene interactions. *ACM Trans. Graph.*, 38(6):209–1, 2019. [2](#)
- [90] Yu Sun, Yun Ye, Wu Liu, Wenpeng Gao, Yili Fu, and Tao Mei. Human mesh recovery from monocular images via a skeleton-disentangled representation. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 5349–5358, 2019. [2](#)
- [91] Graham W. Taylor, Geoffrey E. Hinton, and Sam T. Roweis. Modeling human motion using binary latent variables. In *Proc. NIPS*, 2007. [1](#), [2](#)
- [92] Nikolaus F. Troje. Decomposing biological motion: A framework for analysis and synthesis of human gait patterns. *Journal of Vision*, 2(5):2–2, Sept. 2002. [16](#)
- [93] Matthew Trumble, Andrew Gilbert, Charles Malleson, Adrian Hilton, and John Collomosse. Total capture: 3d human pose estimation fusing video and inertial sensors. In *Proceedings of the British Machine Vision Conference (BMVC)*, Sept. 2017. [16](#)
- [94] Shuhei Tsuchida, Satoru Fukayama, Masahiro Hamasaki, and Masataka Goto. Aist dance video database: Multi-genre, multi-dancer, and multi-camera database for dance information processing. In *Proceedings of the 20th International Society for Music Information Retrieval Conference, ISMIR 2019*, pages 501–510, Delft, Netherlands, Nov. 2019. [20](#), [21](#)
- [95] Simon Fraser University and National University of Singapore. SFU Motion Capture Database. [16](#)
- [96] Raquel Urtasun, David J. Fleet, and Pascal Fua. 3D people tracking with Gaussian process dynamical models. In *IEEE Conf. Comp. Vis. & Pattern Rec.*, pages 238–245, 2006. Vol. 1. [1](#), [2](#)
- [97] Raquel Urtasun, David J. Fleet, and Pascal Fua. Temporal motion models for monocular and multiview 3D human body tracking. *CVIU*, 104(2):157–177, 2006. [1](#), [2](#)
- [98] Gül Varol, Duygu Ceylan, Bryan Russell, Jimei Yang, Ersin Yumer, Ivan Laptev, and Cordelia Schmid. BodyNet: Volumetric inference of 3D human body shapes. In *ECCV*, 2018. [2](#)
- [99] Gül Varol, Ivan Laptev, Cordelia Schmid, and Andrew Zisserman. Synthetic humans for action recognition from unseen viewpoints. *International Journal of Computer Vision*, 129(7):2264–2287, 2021. [1](#)
- [100] Gül Varol, Javier Romero, Xavier Martin, Naureen Mahmood, Michael J Black, Ivan Laptev, and Cordelia Schmid. Learning from synthetic humans. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 109–117, 2017. [1](#)
- [101] Jack M. Wang. Gaussian process dynamical models for human motion. Master’s thesis, University of Toronto, 2005. [2](#)
- [102] Jianqiao Wangni, Dahua Lin, Ji Liu, Kostas Daniilidis, and Jianbo Shi. Towards statistically provable geometric 3d human pose recovery. *SIAM Journal on Imaging Sciences*, 14(1):246–270, 2021. [15](#)
- [103] Yuxin Wu and Kaiming He. Group normalization. In *Proceedings of the European conference on computer vision (ECCV)*, pages 3–19, 2018. [4](#), [16](#)
- [104] Donglai Xiang, Hanbyul Joo, and Yaser Sheikh. Monocular total capture: Posing face, body, and hands in the wild. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 10965–10974, 2019. [2](#)
- [105] Dongseok Yang, Doyeon Kim, and Sung-Hee Lee. LoB-STR: Real-time Lower-body Pose Prediction from Sparse Upper-body Tracking Signals. *Computer Graphics Forum*, 2021. [2](#)
- [106] Ye Yuan and Kris Kitani. Dlow: Diversifying latent flows for diverse human motion prediction. In *European Conference on Computer Vision*, pages 346–364. Springer, 2020. [6](#)
- [107] Ye Yuan, Shih-En Wei, Tomas Simon, Kris Kitani, and Jason Saragih. Simpoe: Simulated character control for 3d human pose estimation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021. [2](#)
- [108] Andrei Zanfir, Eduard Gabriel Bazavan, Hongyi Xu, William T Freeman, Rahul Sukthankar, and Cristian Sminchisescu. Weakly supervised 3d human pose and shape reconstruction with normalizing flows. In *European Conference on Computer Vision*, pages 465–481. Springer, 2020. [2](#)

- [109] Andrei Zanfir, Elisabeta Marinoiu, and Cristian Sminchisescu. Monocular 3d pose and shape estimation of multiple people in natural scenes-the importance of multiple scene constraints. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 2148–2157, 2018. [2](#), [6](#)
- [110] Jason Y Zhang, Panna Felsen, Angjoo Kanazawa, and Jitendra Malik. Predicting 3d human dynamics from video. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 7114–7123, 2019. [2](#)
- [111] Tianshu Zhang, Buzhen Huang, and Yangang Wang. Object-occluded human shape and pose estimation from a single color image. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2020. [2](#)
- [112] Yan Zhang, Michael J Black, and Siyu Tang. We are more than our joints: Predicting how 3d bodies move. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 3372–3382, 2021. [3](#), [14](#)

Appendices

Here we provide details and extended evaluations omitted from the main paper¹ for brevity. App. A provides extended discussions, App. B and App. C give method details regarding the HuMoR model and test-time optimization (TestOpt), App. D derives our optimization energy from a probabilistic perspective, App. E provides experimental details from the main paper, and App. F contains extended experimental evaluations.

We encourage the reader to view the **supplementary videos** on the project webpage² and supplementary webpage³ for extensive qualitative results. We further discuss these results in App. F.

A. Discussions

State Representation. Our state representation is somewhat redundant to include both explicit joint positions $\mathbf{J}$ and SMPL parameters (which also give joint positions $\mathbf{J}^{\text{SMPL}}$). This is motivated by recent works [54, 112] which show that using an extrinsic representation of body keypoints (*e.g.* joint positions or mesh vertices) helps in learning motion characteristics like static contact, thereby improving the visual quality of generated motions. The over-parameterization, unique to our approach, additionally allows for consistency losses leveraged during CVAE training and in TestOpt.

Another noteworthy property of our state is that it does not explicitly represent full-body shape – only bone proportions are implicitly encoded through joint locations. During training, we use shape parameters β provided in AMASS [63] to compute $\mathcal{L}_{\text{SMPL}}$, but otherwise the CVAE is shape-unaware. Extending our formulation to include full-body shape is an important direction for improved generalization and should be considered in future work.

Conditioning on More Time Steps. Alternatively, we could condition the dynamics learned by the CVAE with additional previous steps, *i.e.* $p_{\theta}(\mathbf{x}_t | \mathbf{x}_{t-1}, \dots, \mathbf{x}_{t-p})$, however since $\mathbf{x}_{t-1}$ includes velocities this is unnecessary and only increases the chances of *overfitting* to training motions. It would additionally increase the necessary computation for both generation and TestOpt.

Why CVAE? Our use of a CVAE to model motion is primarily motivated by recent promising results in the graphics community [54, 29]. Not only is it a simple solution, but also affords the physical interpretation presented in the main paper. Other deep generative models could be considered for $p_{\theta}(\mathbf{x}_t | \mathbf{x}_{t-1})$, however each have potential issues compared to our CVAE. The conditional generative adversarial

network [68] would use standard normal noise for $\mathbf{z}_t$, which we show is insufficient in multiple experiments. Furthermore, it does not allow for inferring a latent transition $\mathbf{z}_t$. Past works have had success with recurrent and variational-recurrent architectures [112]. As discussed previously, the reliance of these networks on multiple timesteps increases overfitting which is especially dangerous for our estimation application which requires being able to represent arbitrary observed motions. Finally, normalizing flows [46] and neural ODEs [19] show exciting potential for modeling human motion, however conditional generation with these models is not yet well-developed.

In comparison to Motion VAE (MVAE) [54]. Our proposed CVAE is inspired by MVAE, but introduces a number of key improvements that enable generalization and expressivity: (i) HuMoR uses a neural network to learn a conditional prior $p_{\theta}(\mathbf{z}_t | \mathbf{x}_{t-1})$ rather than assuming $p_{\theta}(\mathbf{z}_t) = \mathcal{N}(\mathbf{z}_t; \mathbf{0}, \mathbf{I})$, (ii) the decoder predicts the *change in state* Δ_{θ} rather than the next state $\mathbf{x}_t$ directly, (iii) the decoder outputs person-ground contacts $\mathbf{c}_t$ which MVAE does not model, and (iv) HuMoR trains using $\mathcal{L}_{\text{SMPL}}$ regularization to encourage joint position/angle consistency whereas MVAE uses the typical ELBO. Tab. 3 in the main paper shows that these differences are crucial to achieve good results as MVAE does not work well. The conditional prior and $\mathcal{L}_{\text{SMPL}}$ are particularly important: learning $p_{\theta}(\mathbf{z}_t | \mathbf{x}_{t-1})$ is theoretically justified when deriving the CVAE and is intuitive for human motion, while $\mathcal{L}_{\text{SMPL}}$ provides a strong supervision to improve stability of model rollout.

Furthermore, the pose state employed by HuMoR and MVAE differ slightly. In MVAE, the root is defined by projecting the pelvis onto the ground, giving 2D linear and 1D angular velocities. In HuMoR, the root is at the pelvis giving full 3D velocities.

A Note on β -VAE [37]. The KL weight w_{KL} in Eq. 7 of the main paper is not directly comparable to a typical β -VAE [37] due to various implementation details. First, $\mathcal{L}_{\text{rec}}$ is the mean-squared error (MSE) of the unnormalized state rather than the true log-likelihood. The use of additional regularizers $\mathcal{L}_{\text{reg}}$ that are not formulated probabilistically to be part of the reconstruction loss further compounds the difference. Furthermore, in practice losses are averaged over both the feature and batch dimensions as not to depend on chosen dimensionalities. All these differences result in setting $w_{\text{KL}} = 4e^{-4}$.

The Need for $\mathcal{E}_{\text{reg}}$ in Optimization. The motion prior term $\mathcal{E}_{\text{mot}} = \mathcal{E}_{\text{CVAE}} + \mathcal{E}_{\text{init}}$, which leverages our learned conditional prior and GMM, nicely falls out of the MAP derivation (see App. D below) and is by itself reasonable to ensure motion is plausible. However, in practice it can be prone to local minima and slow to converge without any regularization. This is primarily because HuMoR is trained

¹In the rest of this document we refer to the main paper briefly as *paper*.

²<https://geometry.stanford.edu/projects/humor/>

³<https://geometry.stanford.edu/projects/humor/supp.html>

on clean motion capture data from AMASS [63], but in the early stages of optimization the initial state $\mathbf{x}_0$ will be far from this domain. This means rolled out motions using the CVAE decoder will be implausible and the likelihood output from learned conditional prior is not necessarily meaningful (since inputs will be well outside the training distribution). The additional regularizers presented in the main paper, mainly $\mathcal{E}_{\text{skel}}$ and $\mathcal{E}_{\text{env}}$, allow us to resolve this issue by reflecting expected behavior of the motion model when it is producing truly plausible motions (*i.e.* $\mathbf{x}_0$ is similar to the training data).

On Evaluation Metrics. As discussed in prior work [79], traditional positional metrics used to evaluate root-relative pose estimates do not capture the accuracy of the absolute (“global”) motion nor its physical/perceptual plausibility. This is why we use a range of metrics to capture both the global joint accuracy, local joint accuracy (after aligning root joints), and plausibility of a motion. However, these metrics still have flaws and there is a need to develop more informative motion estimation evaluation metrics for both absolute accuracy and plausibility. This is especially true in scenarios of severe occlusions where there is not a single correct answer: even if the “ground truth” 3D joints are available, there may be multiple motions that explain the partial observations equally well.

On Convergence. Our multi-objective optimization uses a mixture of convex and non-convex loss functions. As we utilize L-BFGS, the minimum energy solution we report is only locally optimal. While simulated annealing or MCMC / HMC (Markov Chain Monte Carlo / Hamiltonian Monte Carlo) type of exploration approaches can be deployed to search for the global optimum, such methods would incur heavy computational load and hence are prohibitive in our setting. Thanks to the accurate initialization, we found that most of the time TestOpt converges to a good minimum. This observation is also lightly supported by recent work arguing that statistically provable convergence can be attained for the human pose problem under convex and non-convex regularization using a multi-stage optimization scheme [102].

A.1. Assumptions and Limitations

On the Assumption of a Ground Plane. We use *the ground* during TestOpt to obtain a transformation to the canonical reference frame where our prior is trained. While this is a reasonable assumption in a majority of scenarios, we acknowledge that certain applications might require *in-the-wild* operation where a single ground plane does not exist *e.g.* climbing up stairs or moving over complex terrain. In such scenarios, we require a consistent reference frame, which can be computed from: (i) an accelerometer if a mobile device is used, (ii) pose of static, rigid objects if an

object detector is deployed, (iii) fiducial tags or any other means of obtaining a gravity direction.

Note that the ground plane is not an essential piece in the test-time optimization. It is a requirement only because of the way our CVAE is trained: on motions with a ground plane at $z = 0$, gravity in the $-z$ direction, and without complex terrain interactions. Although we empirically noticed that convergence of training necessitates this assumption, other architectures or the availability of larger *in-the-wild* motion datasets might make training HuMoR possible under arbitrary poses. This perspective should clarify why our method *can work* when the ground is invisible: TestOpt might converge from a bad initialization as long as our prior (HuMoR) is able to account for the observation.

On the Assumption of a Static Camera. While a static camera is assumed in all of our evaluations, recent advances in 3D computer vision make it possible to overcome this limitation. Our method, backed by either a structure from motion / SLAM pipeline or a camera relocalization engine, can indeed work in scenarios where the camera moves as well as the human targets. A more sophisticated solution could leverage our learned motion model to disambiguate between camera and human motion. Expectedly, this requires further investigation, making room for future studies as discussed at the end of the main paper.

Other Limitations and Failure Cases. As discussed in Sec. 6 of the main paper, HuMoR has limitations that motivate multiple future directions. First, optimization is generally slow compared to learning-based (direct prediction) methods. This also reflects on our test-time optimization. Approaches for *learning to optimize* can come handy in increasing the efficiency of our method. Additionally, our current formulation of TestOpt allows only for a single output, the local optimum. Therefore, future work may explore learned approaches yielding multi-hypothesis output, which can be used to characterize uncertainty.

Specific failure cases (as shown in the supplementary videos and Fig. 6) further highlight areas of future improvement. First, extreme occlusions (*e.g.* only a few visible points as in Fig. 6 left), especially at the first frame which determines $\mathbf{x}_0$, makes for a difficult optimization that often lands in local minima with implausible motions. Second, uncommon motions that are rare during CVAE training, such as laying down in Fig. 6 (middle), can cause spurious ground plane outputs as TestOpt attempts to make the motion more likely. Leveraging more holistic scene understanding methods and models of human-environment interaction will help in these cases. Finally, our method is dependent on motion in order to resolve ambiguity, which is usually very helpful but has corner cases as shown in Fig. 6 (right). For example, if the observed person is nearly static, the optimization may produce implausible poses due to ambiguous occlusions (*e.g.* standing when really the person is

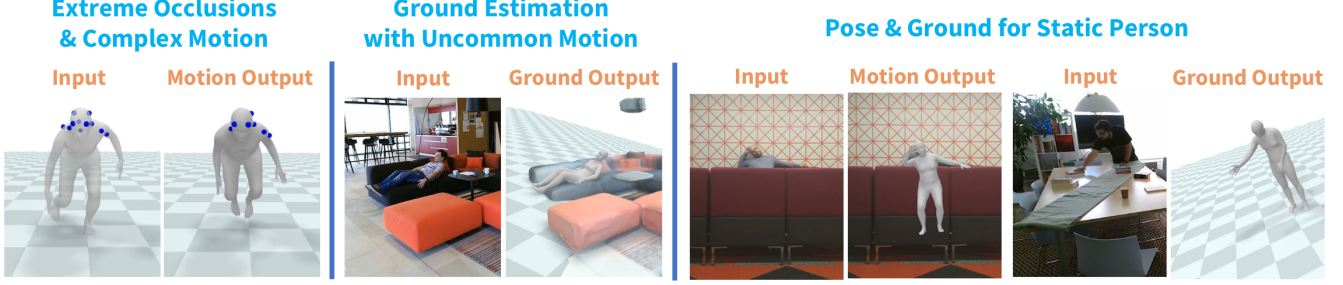


Figure 6: Failure cases of TestOpt using HuMoR. Please see App. A.1 or the supplementary videos for details of each.

sitting) and/or incorrect ground plane estimations.

B. HuMoR Model Details

In this section, we provide additional implementation details for the HuMoR motion model described in Sec. 3 of the main paper.

B.1. CVAE Architecture and Implementation

Body Model. We use the SMPL+H body model [81] since it is used by the AMASS [63] dataset. However, our focus is on modeling body motion, so HuMoR and TestOpt do not consider the hand joints (leaving the 22 body joints including the root). Hand joints could be straightforwardly optimized with body motion, but was not in our current scope.

Canonical Coordinate Frame. To ease learning and improve generalization, our network operates on inputs in a canonical coordinate frame. Specifically, based on $\mathbf{x}_{t-1}$ we apply a rotation around the up (+z) axis and translation in x, y such that the x and y components of $\mathbf{r}_{t-1}$ are 0 and the person’s body right axis (w.r.t. Φ_{t-1}) is facing the +x direction.

Architecture. The encoder and prior networks are identical multi-layer perceptrons (MLP) with 5 layers and hidden size 1024. The decoder is a 4-layer MLP with hidden sizes (1024, 1024, 512). The latent transition $\mathbf{z}_t \in \mathbb{R}^{48}$ is skip-connected to every layer of the decoder in order to emphasize its importance and help avoid posterior collapse [54]. ReLU non-linearities and group normalization [103] with 16 groups are used between all layers except outputs in each network. Input rotations are represented as matrices, while the network outputs the axis-angle representation in $\mathbb{R}^3$. In total, the CVAE network contains ~ 9.7 million parameters.

B.2. CVAE Training

Losses. The loss function used for training is primarily described in the main paper (see Eq. 7). For a training pair $(\mathbf{x}_{t-1}, \mathbf{x}_t)$, the KL divergence loss term is computed between the output distributions of the encoder and condi-

tional prior as

$$\begin{aligned} \mathcal{L}_{\text{KL}} &= D_{\text{KL}}(q_{\phi}(\mathbf{z}_t | \mathbf{x}_t, \mathbf{x}_{t-1}) \parallel p_{\theta}(\mathbf{z}_t | \mathbf{x}_{t-1})) \\ &= D_{\text{KL}}(\mathcal{N}(\mathbf{z}_t; \mu_{\phi}(\mathbf{x}_t, \mathbf{x}_{t-1}), \sigma_{\phi}(\mathbf{x}_t, \mathbf{x}_{t-1})) \parallel \mathcal{N}(\mathbf{z}_t; \mu_{\theta}(\mathbf{x}_{t-1}), \sigma_{\theta}(\mathbf{x}_{t-1}))). \end{aligned} \quad (15)$$

The SMPL loss $\mathcal{L}_{\text{SMPL}}$ is computed using the ground truth shape parameters β provided in AMASS on the ground truth gendered body model.

Dataset. For training, we use AMASS [63]: a large, publicly-available motion capture (mocap) database containing over 11k motion sequences from 344 different people fit to SMPL. The database aggregates and standardizes many mocap datasets into one. We pre-process AMASS by cropping the middle 80% of each motion sequence, sub-sampling to 30 Hz, estimating velocities with finite differences, and using automated heuristics based on foot contacts to remove sequences with substantial terrain interaction (e.g. stairs, ramps, or platforms). We automatically annotate ground contacts for 8 body joints (*left and right toes, heels, knees, and hands*) based on velocity and height. In particular, if a joint has moved less than 0.5cm in the last timestep and its z component is within 8cm of the floor, it is considered to be in contact. For toe joints, we use a tighter height threshold of 4cm.

For training the CVAE, we use the recommended training split (save for TCD Hands [40] which contains mostly hand motions): CMU [17], MPI Limits [2], TotalCapture [93], Eyes Japan [60], KIT [65], BMLrub [92], BML-movi [28], EKUT [65], and ACCAD [1]. For validation during training we use MPI HDM05 [70], SFU [95], and MPI MoSh [58]. Finally for evaluations (Sec. 5.3 of the main paper), we use HumanEva [87] and Transitions [63].

Training Procedure. We train using 10-step sequences sampled on-the-fly from the training set (in order to use scheduled sampling as detailed below). To acquire a training sequence, a full mocap sequence is randomly (uniformly) chosen from AMASS and then a random 10-step window within that sequence is (uniformly) sampled. Training is performed using batches of 2000 sequences for 200 epochs with Adamax [44] and settings $\beta_1 = 0.9$,

$\beta_2 = 0.999$, and $\epsilon = 1e^{-8}$. We found this to be more stable than using Adam. The learning rate starts at $1e^{-4}$ and decays to $5e^{-5}$, $2.5e^{-5}$, and $1.25e^{-5}$ at epochs 50, 80, and 140, respectively. We use early stopping by choosing the network parameters that result in the best validation split performance throughout training.

A common difficulty in training VAEs is posterior collapse [61] – when the learned latent encoding $\mathbf{z}_t$ is effectively ignored by the decoder. This problem is exacerbated in CVAEs since the decoder receives additional conditioning [54, 88]. To combat collapse, we linearly anneal w_{KL} from 0.0 to its full value of $4e^{-4}$ over the first 50 epochs. We also found that our full model, which uses a learned conditional prior, was less susceptible to posterior collapse than the baselines that assume $p_\theta(\mathbf{z}_t) = \mathcal{N}(\mathbf{z}_t; \mathbf{0}, \mathbf{I})$.

Training Computational Requirements. We train our CVAE on a single Tesla V100 16GB GPU, which takes approximately 4 days.

Scheduled Sampling. As explained in the main paper, our scheduled sampling follows [54]. In particular, at each training epoch i we define a probability $s_i \in [0.0, 1.0]$ of using the ground truth state input $\mathbf{x}_{t-1}$ at each timestep t in a training sequence, as opposed to the model’s own previous output $\hat{\mathbf{x}}_{t-1}$. Training is done using a curriculum that includes $s_i = 1.0$ (regular supervised training), $s_i \in (0.0, 1.0)$ (mix of true and self inputs at each step), and finally $s_i = 0.0$ (always use full generated rollouts). Importantly for training stability, if using the model’s own prediction $\hat{\mathbf{x}}_{t-1}$ as input to t , we do not backpropagate gradients from the loss on $\hat{\mathbf{x}}_t$ back through $\hat{\mathbf{x}}_{t-1}$.

For CVAE training, we use 10 epochs of regular supervised training, 10 of mixed true and self inputs, and the rest using full self-rollouts.

B.3. Initial State GMM

State Representation. Since the GMM models a single state, we use a modified representation that is minimal (*i.e.* avoids redundancies) in order to be useful during test-time optimization. In particular the GMM state is

$$\mathbf{x}^{\text{GMM}} = \begin{bmatrix} \dot{\mathbf{r}} & \dot{\Phi} & \mathbf{J} & \dot{\mathbf{J}} \end{bmatrix} \quad (16)$$

with $\dot{\mathbf{r}}, \dot{\Phi} \in \mathbb{R}^3$ the root linear and angular velocities, and joint positions and velocities $\mathbf{J}, \dot{\mathbf{J}} \in \mathbb{R}^{3 \times 22}$. During TestOpt, joints are determined from the current SMPL parameters $\mathbf{J} = \mathbf{J}^{\text{SMPL}} = M(\mathbf{r}, \Phi, \Theta, \beta)$ so that gradients of the GMM log-likelihood (Eq. 11 in the main paper) will be used to update the initial state SMPL parameters.

Implementation Details. The GMM uses full covariance matrices for each of the 12 components and operates in the same canonical coordinate frame as the CVAE. It trains using expectation maximization⁴ on every state in the same

AMASS training set used for the CVAE.

C. Test-Time Optimization Details

In this section, we give additional details of the motion and shape optimization detailed in Sec. 4 of the main paper.

State Representation. In practice, for optimizing $\mathbf{x}_0$ we slightly modify the state from Eq. 1 in the main paper. First, we remove the joint positions $\mathbf{J}$ to avoid the previously discussed redundancy, which is good for training the CVAE but bad for test-time optimization. Instead, we use $\mathbf{J}^{\text{SMPL}}$ whenever needed at t_0 . Second, we represent body pose Θ in the latent space of the VPoser [73] pose prior with $\mathbf{z}_0^{\text{pose}} \in \mathbb{R}^{32}$. Whenever needed, we can map between full joint angles and latent pose using the VPoser encoder and decoder. Finally, in our implementation state variables $\mathbf{x}_t$ are by default represented in the coordinate frame of the given observations, *e.g.* relative to the camera, to allow easily fitting to data – they are transformed back into the canonical CVAE frame when necessary as discussed below.

Floor Parameterization. As detailed in the main paper, to obtain the transformation between the canonical coordinate frame in which our CVAE is trained and the observation frame used for optimization, we additionally optimize the floor plane of the scene $\mathbf{g} \in \mathbb{R}^3$. This parameterization is $\mathbf{g} = d\hat{\mathbf{n}}$ where $\hat{\mathbf{n}}$ is the ground unit normal vector and d the plane offset. To disambiguate the normal vector direction $\hat{\mathbf{n}}$ given $\mathbf{g}$, we assume that the y -component of the normal vector must be negative, *i.e.* it points upward in the camera coordinate frame. This assumes the camera is not severely tilted such that the observed scene is “upside down”.

Observation-to-Canonical Transformation. We assume that gravity is orthogonal to the ground plane. Therefore, given the current floor $\mathbf{g}$ and root state $\mathbf{r}, \Phi$ (in the observation frame) we compute a rotation and translation to the canonical CVAE frame: after the transformation, $\hat{\mathbf{n}}$ is aligned with $+z$ and $d = 0$, Φ faces body right towards $+x$, and the x, y components of $\mathbf{r}$ are 0. With this ability, we can always compute the (observed) state at time $\mathbf{x}_t$ from $\mathbf{z}_{1:t}$, $\mathbf{x}_0$, and $\mathbf{g}$ by (i) transforming $\mathbf{x}_0$ to the canonical frame, (ii) using the CVAE to rollout $\mathbf{x}_t = f(\mathbf{x}_0, \mathbf{z}_{1:t})$, and (iii) transforming $\mathbf{x}_t$ back to the observation frame.

Optimization Objective Details. The optimization objective is detailed in Sec. 4.2 of the main paper. To compensate for the heavy tailed behavior of real data, we use robust losses for multiple data terms. $\mathcal{E}_{\text{data}}^{2\text{D}}$ uses the Geman-McClure function [26] which for our purposes is defined as $\rho(r, \sigma) = (\sigma^2 r^2) / (\sigma^2 + r^2)$ for a residual r and scaling factor σ . We use $\sigma = 100$ for all experiments. $\mathcal{E}_{\text{data}}^{\text{PC3D}}$ uses robust bisquare weights [8]. These weights are computed based on the one-way chamfer distance term (see Eq. 14 in the main paper): residuals over the whole sequence are first normalized using a robust estimate of the standard deviation

⁴using scikit-learn

based on the median absolute deviation (MAD), then each weight is computed as

$$w_{bs} = \begin{cases} (1 - (\hat{r}/\kappa)^2)^2 & |\hat{r}/\kappa| < 1 \\ 0.0 & \text{else} \end{cases} \quad (17)$$

In this equation, $\hat{r}$ is a normalized residual and κ is a tuning constant which we set to 4.6851.

In the $\mathcal{E}_{\text{env}}$ energy term, we use $\delta = 8 \text{ cm}$ to ensure the z -component of contacting joints are within 8 cm of the floor when in contact (since joints are inside the body) in the canonical frame.

Initialization. As detailed in Sec. 4.2 of the main paper, our optimization is initialized by directly optimizing SMPL pose and shape parameters using $\mathcal{E}_{\text{data}}$ and $\mathcal{E}_{\text{shape}}$ along with a pose prior $\mathcal{E}_{\text{pose}}$ and joint smoothing $\mathcal{E}_{\text{smooth}}$. The latter are weighted by λ_{pose} and λ_{smooth} . This two-stage initialization first optimizes global translation and orientation for 30 optimization steps, followed by full pose and shape for 80 steps. At termination, we estimate velocities using finite differences, which allows direct initialization of the state $\mathbf{x}_0^{\text{init}}$. To get $\mathbf{z}_{1:T}$, the CVAE encoder is used to infer the latent transition between every pair of frames. The initial shape parameters β_{init} are a direct output of the initialization optimization. Finally, for fitting to RGB(-D) the ground plane is initialized from video with PlaneRCNN [55], though we found simply setting the floor to $y = 0$ (i.e. the normal is aligned with the camera up axis) works just as well in most cases.

Optimization (TestOpt) Details. Our optimization is implemented in PyTorch [72] using L-BFGS with a step size of 1.0 and *autograd*. For all experiments, we optimize using the neutral SMPL+H [81] body model in 3 stages. First, only the initial state $\mathbf{x}_0$ and first 15 frames of the latent sequence $\mathbf{z}_{1:15}$ are optimized for 30 iterations in order to quickly reach a reasonable initial state. Next, $\mathbf{x}_0$ is fixed while the full latent dynamics sequence $\mathbf{z}_{1:T}$ is optimized for 25 iterations, and then finally the full sequence and initial state are tuned together for another 15 iterations. The ground $\mathbf{g}$ and shape β are optimized in every stage.

The energy weights used for each experiment in the main paper are detailed in Tab. 5. The left part of the table indicates weights for the initialization phase (i.e. the *VPoser-t* baseline), while the right part is our full proposed optimization. A dash indicates the energy is not relevant for that data modality and therefore not used. Weights were manually tuned using the presented evaluation metrics and qualitative assessment. Note that for similar modalities (e.g. 3D joints and keypoints, or RGB and RGB-D) weights are quite similar and so only slight tuning should be necessary to transfer to new data. The main tradeoff comes between reconstruction accuracy and motion plausibility: e.g. the motion prior is weighted higher for i3DB, which contains many severe

occlusions, than for PROX RGB where the person is often nearly fully visible.

D. MAP Objective Derivation

In this section, we formulate the core of the pose and shape optimization objective (Eq. 10 in the main paper) from a probabilistic perspective. Recall, we want to optimize the initial state $\mathbf{x}_0$, a sequence of latent variables $\mathbf{z}_{1:T}$, ground $\mathbf{g}$, and shape β based on a sequence of observations $\mathbf{y}_{0:T}$. We are interested in the maximum a-posteriori (MAP) estimate:

$$\max_{\mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta} p(\mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta | \mathbf{y}_{0:T}) \quad (18)$$

$$= \max_{\mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta} p(\mathbf{y}_{0:T} | \mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta) p(\mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta) \quad (19)$$

Assuming $\mathbf{y}_0$ is independent of $\mathbf{g}$, the left term is written

$$p(\mathbf{y}_0 | \mathbf{x}_0, \beta) \prod_{t=1}^T p(\mathbf{y}_t | \mathbf{z}_{\leq t}, \mathbf{x}_0, \mathbf{g}, \beta) = \prod_{t=0}^T p(\mathbf{y}_t | \mathbf{x}_t, \beta) \quad (20)$$

where $\mathbf{y}_t$ is assumed to only be dependent on the initial state and *past* transitions. Additionally, $\{\mathbf{z}_{\leq t}, \mathbf{x}_0, \mathbf{g}\}$ is replaced with $\mathbf{x}_t = f(\mathbf{x}_0, \mathbf{z}_{1:t})$ using CVAE rollout as detailed previously. The right term in Eq. (19) is written as

$$p(\mathbf{x}_0, \mathbf{g}, \beta) \prod_{t=1}^T p(\mathbf{z}_t | \mathbf{z}_{< t}, \mathbf{x}_0, \mathbf{g}, \beta) \quad (21)$$

$$= p(\mathbf{x}_0 | \mathbf{g}) p(\mathbf{g}) p(\beta) \prod_{t=1}^T p(\mathbf{z}_t | \mathbf{x}_{t-1}) \quad (22)$$

where $\mathbf{x}_0$, $\mathbf{z}_t$, and $\mathbf{g}$ are assumed to be independent of β . We then use these results within Eq. (19) to optimize the log-likelihood:

$$\begin{aligned} & \max_{\mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta} \log p(\mathbf{y}_{0:T} | \mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta) + \log p(\mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta) \\ &= \min_{\mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta} - \sum_{t=0}^T \log p(\mathbf{y}_t | \mathbf{x}_t, \beta) - \sum_{t=1}^T \log p(\mathbf{z}_t | \mathbf{x}_{t-1}) \\ & \quad - \log p(\mathbf{x}_0 | \mathbf{g}) - \log p(\mathbf{g}) - \log p(\beta) \end{aligned} \quad (23)$$

$$\begin{aligned} &= \min_{\mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta} \mathcal{E}_{\text{data}} + \mathcal{E}_{\text{CVAE}} + \mathcal{E}_{\text{init}} + \mathcal{E}_{\text{gnd}} + \mathcal{E}_{\text{shape}} \\ &= \min_{\mathbf{x}_0, \mathbf{z}_{1:T}, \mathbf{g}, \beta} \mathcal{E}_{\text{mot}} + \mathcal{E}_{\text{data}} + \mathcal{E}_{\text{gnd}} + \mathcal{E}_{\text{shape}}. \end{aligned} \quad (24)$$

Assuming each energy presented in the main paper can be written as the log-likelihood of a distribution, this formulation recovers our optimization objective besides the additional regularizers $\mathcal{E}_{\text{skel}}$ and $\mathcal{E}_{\text{env}}$ (these terms could, in principle, be written as part of a more complex motion prior term $\mathcal{E}_{\text{mot}}$, however for simplicity we do not do this). Next,

Dataset	Initialization						Full Optimization										
	$\lambda_{\text{data}}^{3\text{D}}$	$\lambda_{\text{data}}^{2\text{D}}$	$\lambda_{\text{data}}^{\text{PC3D}}$	λ_{shape}	λ_{pose}	λ_{smooth}	$\lambda_{\text{data}}^{3\text{D}}$	$\lambda_{\text{data}}^{2\text{D}}$	$\lambda_{\text{data}}^{\text{PC3D}}$	λ_{shape}	λ_{CVAE}	λ_{init}	λ_{c}	λ_{b}	λ_{cv}	λ_{ch}	λ_{gnd}
AMASS (occ keypoints)	1.0	-	-	0.015	$2e^{-4}$	0.1	1.0	-	-	0.015	$5e^{-4}$	$5e^{-4}$	1.0	10.0	1.0	1.0	-
AMASS (noisy joints)	1.0	-	-	0.015	$2e^{-4}$	10.0	1.0	-	-	0.015	$1e^{-3}$	$1e^{-3}$	1.0	10.0	1.0	1.0	-
i3DB (RGB)	-	$1e^{-3}$	-	4.5	0.04	100.0	-	$1e^{-3}$	-	4.5	0.075	0.075	100.0	$2e^3$	0.0	10.0	15.0
PROX (RGB)	-	$1e^{-3}$	-	4.5	0.04	100.0	-	$1e^{-3}$	-	4.5	0.05	0.05	100.0	$2e^3$	100.0	10.0	15.0
PROX (RGB-D)	-	$1e^{-3}$	1.0	3.0	0.1	100.0	-	$1e^{-3}$	1.0	3.0	0.075	0.075	100.0	$2e^3$	100.0	10.0	90.0

Table 5: Energy weightings used in test-time optimization for each experiment in Sec. 5 of the main paper.

we connect each energy term as presented in Sec. 4.2 of the paper to the probabilistic perspective.

Motion Prior $\mathcal{E}_{\text{mot}}$. This term is already the log-likelihood of our HuMoR motion model (Eq. 11 of the paper), which exactly aligns with the MAP derivation.

Data Term $\mathcal{E}_{\text{data}}$. The form of $p(\mathbf{y}_t|\mathbf{x}_t, \beta)$ is modality-dependent. In the simplest case the observations $\mathbf{y}_t$ are 3D joints (or keypoints with known correspondences) and $p(\mathbf{y}_t|\mathbf{x}_t, \beta)$ is defined by $\mathbf{y}_t = \mathbf{J}_t^{\text{SMPL}} + \epsilon$ with $\epsilon \sim \mathcal{N}(0, \sigma_{\text{data}})$. Then the energy is as written in Eq. 12 of the paper. For other modalities (Eq. 13 and 14 in the paper), the data term can be seen as resulting from a more sophisticated noise model.

Ground Prior $\mathcal{E}_{\text{gnd}}$. We assume the ground should stay close to initialization so $p(\mathbf{g}) = \mathcal{N}(\mathbf{g}; \mathbf{g}^{\text{init}}, \sigma_{\text{gnd}})$ corresponding to the objective in the paper $\mathcal{E}_{\text{gnd}} = \lambda_{\text{gnd}} \|\mathbf{g} - \mathbf{g}^{\text{init}}\|^2$.

Shape Prior $\mathcal{E}_{\text{shape}}$. The shape β should stay near neutral zero and so $p(\beta) = \mathcal{N}(0, \mathbf{I})$ which gives the energy $\mathcal{E}_{\text{shape}} = \lambda_{\text{shape}} \|\beta\|^2$.

E. Experimental Evaluation Details

In this section, we provide details of the experimental evaluations in Sec. 5 of the main paper.

E.1. Datasets

AMASS [63] We use the same processed AMASS dataset as described in App. B.2 for experiments. Experiments in Sec. 5.3 and 5.4 of the main paper use the held out Transitions and HumanEva [87] subsets which together contain 4 subjects and about 19 minutes of motion.

i3DB [69] is a dataset of RGB videos captured at 30 Hz containing numerous person-environment interactions involving medium to heavy occlusions. It contains annotated 3D joint positions at 10 Hz along with a primitive cuboid 3D scene reconstruction. We run off-the-shelf 2D pose estimation (OpenPose) [16], person segmentation [18], and plane detection [55] models to obtain inputs and initialization for our test-time optimization. We evaluate our method in Sec. 5.5 of the main paper on 6 scenes (scenes 5, 7, 10, 11, 13, and 14) containing 2 people which totals about 1800 evaluation frames. From the annotated 3D objects, we fit a ground plane which is used to compute plausibility metrics.

PROX [34] is a large-scale dataset of RGB-D videos captured at 30 Hz containing person-scene interactions in a variety of environments with light to medium occlusions. We use a subset of the qualitative part of the dataset to evaluate the plausibility of our method’s estimations. The data does *not* have pose annotations, but does contain the scanned scene mesh to which we fit a ground plane for plausibility evaluation. We obtain 2D pose, person masks, and ground plane initialization in the same way as for i3DB. We evaluate in Sec. 5.5 of the main paper on all videos from 4 chosen scenes (N3Office, N3Library, N0Sofa, and MPH1Library) that tend to have more dynamic motions and occlusions. In total, these scenes contain 12 unique people and about 19 minutes of video.

E.2. Baselines and Evaluation Metrics

Motion Prior Baselines. To be usable in our whole framework (*e.g.* test-time optimization with SMPL), the *MVAE* baseline is our proposed CVAE with all ablations applied simultaneously (no delta step prediction, no contact prediction, no SMPL losses, and no learned conditional prior). Note that this differs slightly from the model as presented in [54]: the decoder is an MLP rather than a mixture-of-experts and the layer sizes are larger to provide the necessary representational capacity for training on AMASS. All ablations and *MVAE* are trained in the exact same way as the full model. Additionally, when used in test-time optimization we use the same energy weightings as described in Tab. 5 but with irrelevant energies removed (*e.g.* the *No Contacts* ablation does not allow the use of $\mathcal{E}_{\text{env}}$). Note that $\mathcal{E}_{\text{init}}$ is still used with *MVAE* and all ablations, the only thing that changes is the prior in $\mathcal{E}_{\text{CVAE}}$.

Motion Estimation Baselines. The *VPoser-t* baseline is exactly the initialization phase of our proposed test-time optimization, *i.e.* we use weightings in Tab. 5.

The *PROX-RGB* baseline fits the neutral SMPL-X [73] body model to the same 2D OpenPose detections used by our method. It does not use the face or hand keypoints for fitting similar to our approach. The *PROX-D* baseline uses the fittings provided with the PROX dataset, which are on the known gendered SMPL-X body model and use face/hand 2D keypoints for fitting.

The *VIBE* baseline uses the same 2D OpenPose detections as our method in order to define bounding boxes for

inference. We found this makes for a more fair comparison since the real-time trackers used in their implementation⁵ often fail for medium to heavy occlusions common in our evaluation datasets.

Evaluation Metrics. In order to report occluded (**Occ**) and visible (**Vis**) positional errors separately, we must determine which joints/keypoints are occluded during evaluation. This is easily done for 3D tasks where “occlusions” are synthetically generated. For RGB data in i3DB, we use the person segmentation mask obtained with DeepLabv3 [18] to determine if a ground truth 3D joint is visible after projecting it to the camera.

For a joint $\mathbf{p}_t \in \mathbf{J}_t^{\text{SMPL}}$ at time t the acceleration magnitude (**Accel**) is computed as

$$a = \|(\mathbf{p}_{t-1} - 2\mathbf{p}_t + \mathbf{p}_{t+1})/h^2\| \quad (25)$$

where $h = 1/30$ for all datasets. Ground penetration frequency (**Freq**) for a given penetration threshold g_{thresh} is computed over all D frames in a dataset as

$$\frac{\sum_D \mathbf{1}(d_{\text{pen}}^{\text{ltoe}} > g_{\text{thresh}}) + \mathbf{1}(d_{\text{pen}}^{\text{rtoe}} > g_{\text{thresh}})}{2D} \quad (26)$$

where $\mathbf{1}(\cdot)$ is the indicator function and $d_{\text{pen}}^{\text{ltoe}}, d_{\text{pen}}^{\text{rtoe}}$ are the penetration distances (shortest distance to the ground plane) for the left and right toe joints at the current frame.

E.3. Estimation from 3D Observations

For fitting to 3D data, as presented in Sec. 5.4 of the main paper, the observation and canonical coordinate frames are identical since AMASS data is used, therefore TestOpt does not optimize the ground plane $\mathbf{g}$.

E.4. Estimation from RGB(-D) Observations

i3DB. Positional joint errors are computed using a 12-joint subset of the ground truth 3D joint annotations which correspond to the SMPL joints used by our method and all baselines. These include ankles, knees, wrists, elbows, shoulders, the neck, and the root. The leg joints reported in Tab. 3 of the main paper include the ankles and knees.

In fitting to i3DB videos, we mask 2D joint observations using the person segmentation mask which we found beneficial under the numerous severe occlusions where OpenPose may predict spurious 2D pose with incorrectly high confidence.

PROX. For fitting to PROX RGB(-D) videos we found it best *not* to mask out 2D joints based on the person segmentation mask, as occlusions are typically minor and OpenPose is relatively accurate. However, the segmentation is always used on the point cloud back-projected from the depth map to ignore points far from the person.

⁵see Github

F. Extended Evaluations

In this section, we present experimental evaluations to supplement those in Sec. 5 of the main paper, which were omitted due to space constraints.

F.1. Qualitative Evaluation

Please see the **supplementary videos** for extensive qualitative results corresponding to each experiment in Sec. 5 of the main paper. In this document, we show various representative examples from these videos and summarize important results.

Fig. 9 shows results fitting to occluded 3D keypoints (Sec. 5.4 of the main paper). Performance of TestOpt with HuMoR is compared to the *VPoser-t* and *MVAE* baselines on two sequences. *VPoser-t* fails to produce any plausible lower-body motion since it uses only a pose prior, while using *MVAE* as the motion prior often gives unnatural and implausible motions that don’t align well with the observed keypoints.

Fig. 10 shows results using TestOpt with HuMoR for fitting to noisy 3D joints (Sec. 5.4 of the main paper). Both the estimated motion and contacts are shown for a crawling sequence. Note that HuMoR recovers complex contact patterns involving not only the feet, but also hands and knees.

Fig. 11 demonstrates fitting performance on RGB videos from PROX compared to the *PROX-RGB* baseline (Sec. 5.5 of the main paper). *PROX-RGB* produces temporally incoherent results since it operates on single frames. However, it also uses the scene mesh as input which allows for plausible poses when the person is fully visible. This does not greatly improve results under occlusions, though, often reverting to a mean leg pose similar to *VPoser-t* and *VIBE*.

Fig. 12 demonstrates fitting to RGB-D videos from PROX compared to the *PROX-D* baseline (Sec. 5.5 of the main paper). Using motion as a prior allows for natural interaction within the scene, as detailed in the figure caption.

Fig. 13 shows ground plane estimations when fitting to RGB-D data for each of the scenes in our PROX dataset. The estimated floor is rendered within the true scene mesh for reference.

Finally, we evaluate TestOpt with HuMoR on highly dynamic dancing data to demonstrate the generalization ability of the CVAE motion model. Fig. 14 shows a sample of frames from motions fit to the DanceDB [5] subset of AMASS [64]. In this case, the observations are full-body 3D keypoints. Though HuMoR is trained on data with few dancing motions, it is able to capture these difficult motions at test time since it only operates on pairs of frames. Additionally, Fig. 7 shows fitting results on RGB videos from the AIST dance dataset [94]. Since HuMoR allows for large accelerations, it accurately generalizes to fast motions (top - note motion blur). Moreover, it is able to recover from poor

Method	Global Joint Error				Root-Aligned Joint Error				Ground Pen		
	Vis	Occ	All	Legs	Vis	Occ	All	Legs	Accel	Freq	Dist
No $\mathcal{E}_{init}$	28.04	39.37	30.96	34.91	12.11	22.27	14.73	21.40	2.57	2.80%	0.85
No $\mathcal{E}_c$	29.17	40.62	32.12	35.83	12.99	24.34	15.92	22.66	2.75	1.98%	0.67
No $\mathcal{E}_b$	28.29	39.77	31.25	34.37	13.00	24.36	15.93	22.31	2.93	4.43%	1.13
No $\mathcal{E}_{skel}$	32.87	44.16	35.78	37.55	14.67	26.82	17.80	24.18	3.81	4.35%	1.29
No $\mathcal{E}_{env}$	26.83	37.62	29.61	32.91	12.09	22.04	14.66	20.96	2.51	2.10%	0.66
Full Energy	26.00	34.36	28.15	31.26	12.02	21.70	14.51	20.74	2.43	2.12%	0.68

Table 6: Motion and shape from RGB video (*i.e.* 2D joints) on i3DB [69]. Joint errors are in *cm* and acceleration is *m/s²*. Results use TestOpt with HuMoR. The top part shows various energy terms ablated.

2D joint detections from OpenPose due to the cartwheel motion (bottom).

F.2. Optimization Objective Ablation

In this experiment, we analyze the effect of the energy terms and regularizers in our test-time optimization (TestOpt) formulation. Tab. 6 reports results on the i3DB [69] dataset using TestOpt with HuMoR for different energy ablations.

No $\mathcal{E}_{init}$ does not use the initial state Gaussian mixture model (GMM) as part of the motion prior (*i.e.* assumes a uniform prior over the initial state). This means the input $\mathbf{x}_0$ to CVAE rollout may not be plausible, especially early in optimization, leading to degraded performance. No $\mathcal{E}_c$ and No $\mathcal{E}_b$ remove individual terms of the skeleton regularization: the joint and bone length consistency. No $\mathcal{E}_{skel}$ removes the entire skeleton regularization (both $\mathcal{E}_c$ and $\mathcal{E}_b$), severely affecting final performance. This term is important to ensuring the CVAE is actually rolling out realistic motions. Finally, No $\mathcal{E}_{env}$ removes the contact velocity and height terms, increasing errors particularly for occluded joints while resulting in similar plausibility metrics.

F.3. Sensitivity to Occlusions and Noise

Next, we look at performance on estimation from 3D data under increasing levels of occlusions and noise. Similar to Sec. 5.4 in the main paper, we consider fitting to occluded 3D keypoints (points under a given height threshold are unobserved) and 3D joint locations with added Gaussian noise. For this experiment, we use the held out Transitions subset of AMASS [63].

Mean keypoint errors for **all** and **occluded** points are shown for increasing occlusions in Fig. 8(a)(b). From left to right the occluded height threshold is 0.0, 0.3, 0.6, 0.9, and 1.2 *m* which roughly corresponds to the lower body being occluded from the floor up through the body parts on the *x*-axis. With no occlusions (*None*) *VPoser-t* most closely fits the clean points, while HuMoR outperforms *MVAE* due to improved expressiveness. As occlusions increase, HuMoR and *MVAE* perform similarly while occluded errors increase

Dataset	Batch Size	Mean Seq Time
AMASS (occ keypoints)	12	2.95
AMASS (noisy joints)	12	2.45
i3DB (RGB)	6	6.42
PROX (RGB)	6	4.48
PROX (RGB-D)	6	4.85

Table 7: Mean per-sequence optimization times (in minutes) for evaluations on each dataset. Optimizations are done on batches of 3s (90 frame) sequences.

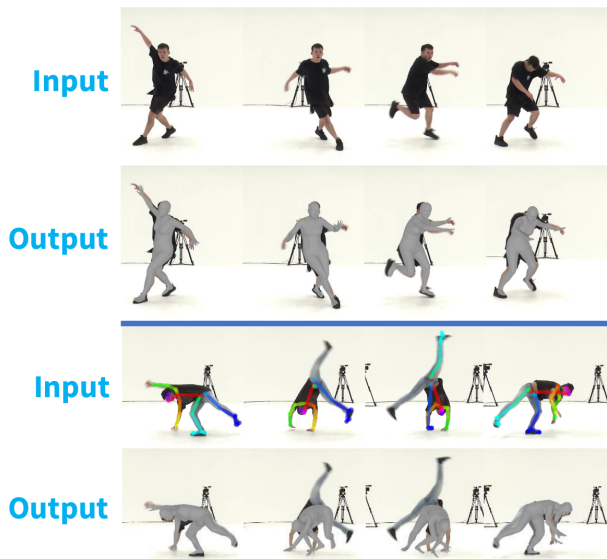


Figure 7: Example sequences using TestOpt with HuMoR to fit to 2D joints in AIST dance videos [94]. HuMoR generalizes to these highly dynamic motions and robustly recovers from inaccurate 2D joint detections (bottom).

greatly for *VPoser-t* as also observed in other experiments. Note that after the knees become occluded, errors for occluded keypoints tend to saturate as performance is dependent nearly entirely on the motion prior.

Mean joint acceleration magnitude is shown for increasingly noisy 3D joint observations in Fig. 8(c). From left

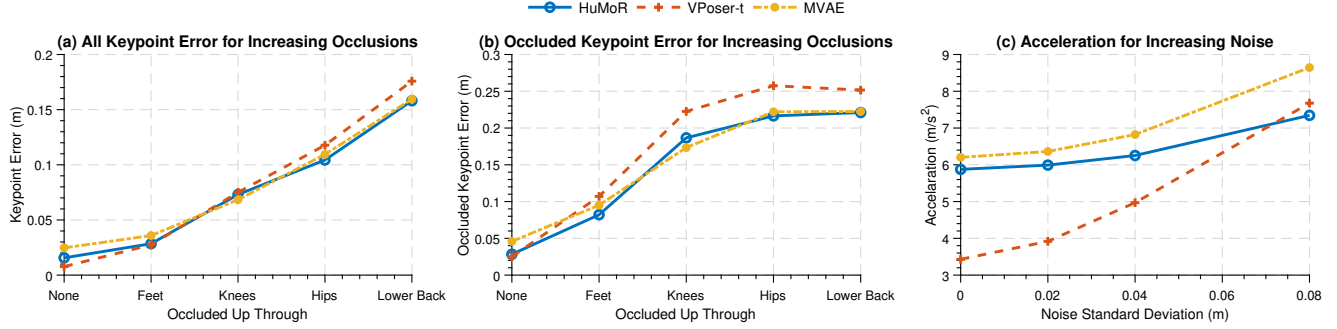


Figure 8: Keypoint errors for **all** (a) and **occluded** (b) points for increasing levels of occlusion when fitting to 3D keypoints with TestOpt. (c) Joint acceleration magnitude for increasing levels of noise when fitting to 3D joints with TestOpt.

to right noise increases to 8 *cm* standard deviation. Importantly, the performance of HuMoR stays relatively stable. HuMoR increases only 25% while *VPoser-t* and *MVAE* increase 124% and 57%, respectively.

F.4. Computational Requirements of TestOpt

For all experiments presented in the main paper, we perform optimization on batches of 3*s* sequences (90 frames). Tab. 7 shows the mean per-sequence optimization times for each of these experiments where *AMASS* corresponds to experiments in Sec. 5.4 of the paper and *i3DB* and *PROX* correspond to Sec. 5.5. The per-sequence time is computed by taking the total time to optimize over the whole dataset divided by the number of 3*s* sequences. Note that batching speeds up this large scale optimization significantly, *i.e.* optimizing a single sequence will only be slightly faster than a batch of sequences since the primary bottleneck is CVAE rollout. All batched optimizations were performed on a 24 GB Titan RTX GPU, though an 8 GB GPU is sufficient to optimize a single sequence.

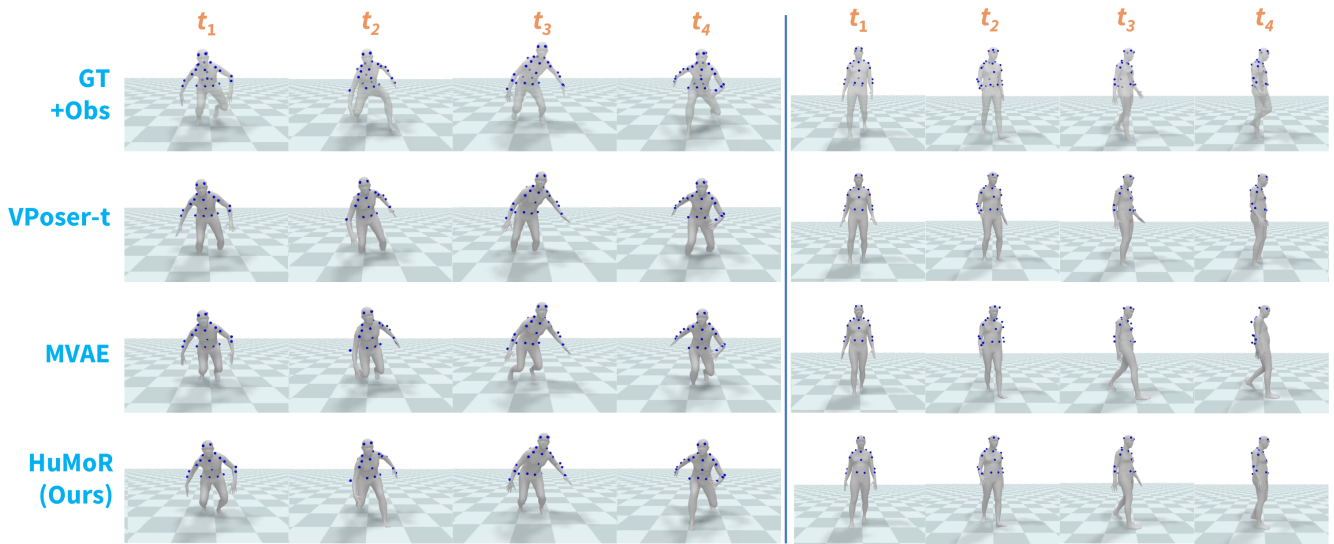


Figure 9: Comparison to baselines when fitting to 3D keypoints from held out sequences in the AMASS dataset. *GT+Obs* shows the ground truth body motion and observed keypoints in blue, while each method output shows the predicted motion with the observed keypoints for reference.

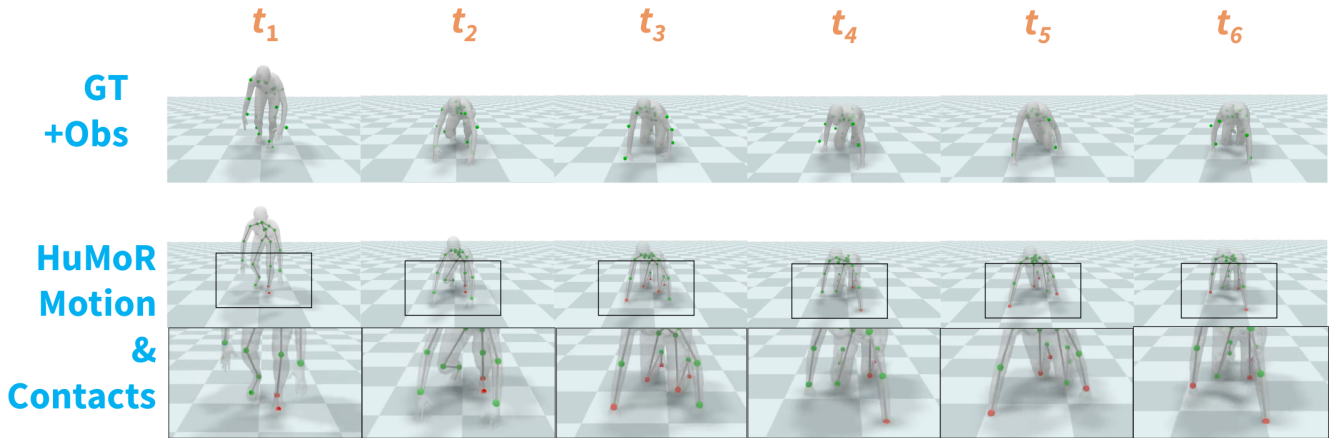


Figure 10: TestOpt with HuMoR recovers complex contact patterns involving feet, knees, and hands from noisy 3D joint observations (shown on top along with the ground truth motion).

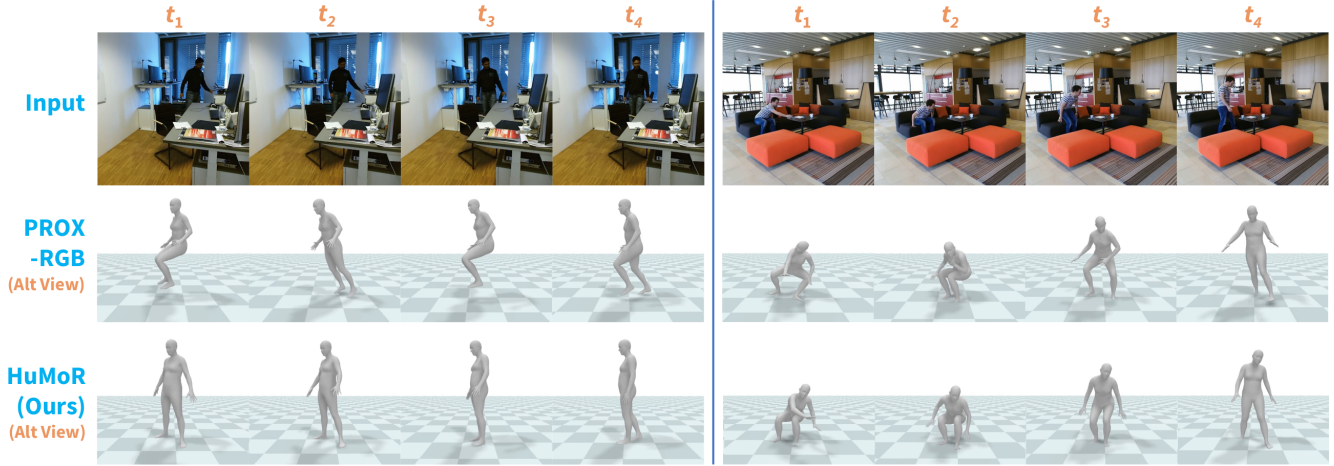


Figure 11: Comparison to *PROX-RGB* on videos from the PROX dataset. Predicted motion is shown from an alternate viewpoint for both methods. In both sequences, *PROX-RGB* is temporally inconsistent. In the left example, the lower body occlusion causes implausible neutral standing or sitting poses. The predicted ground plane from HuMoR is shown for reference.

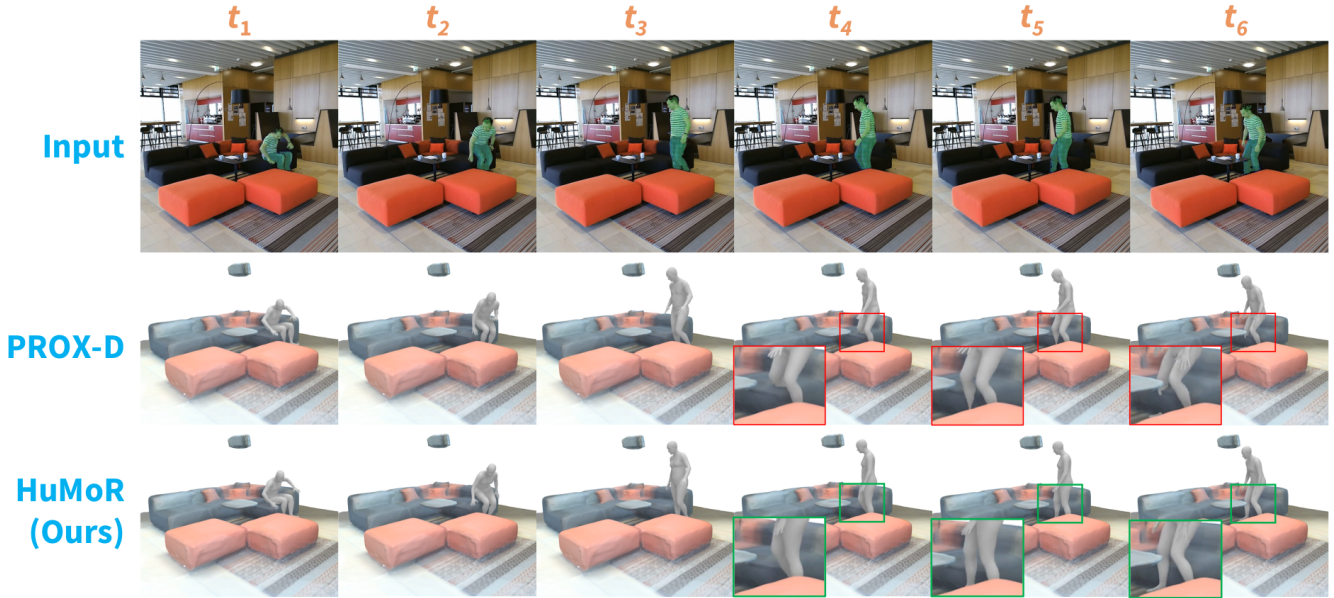


Figure 12: Comparison to *PROX-D* on PROX RGB-D data. Motion and shape results are rendered within the ground truth scene mesh for reference. Though our method does not use the scene mesh as input like *PROX-D*, it still produces motions that are plausible within the environment by using HuMoR as the motion prior, sometimes better than *PROX-D* as indicated by the red and green boxes here.

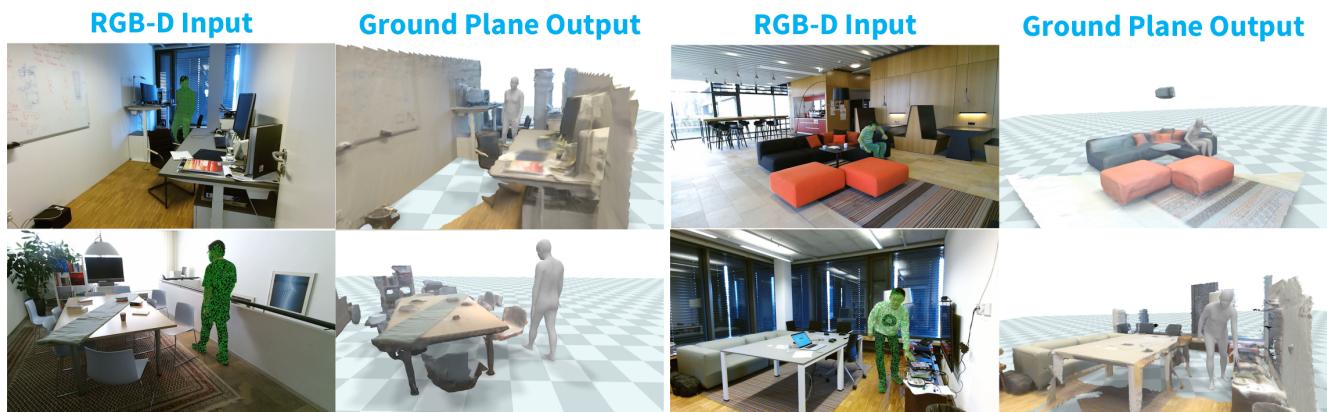


Figure 13: Ground estimation examples from each scene in PROX from RGB-D. The scene mesh is shown for reference only, it is not an input or output of TestOpt with HuMoR.

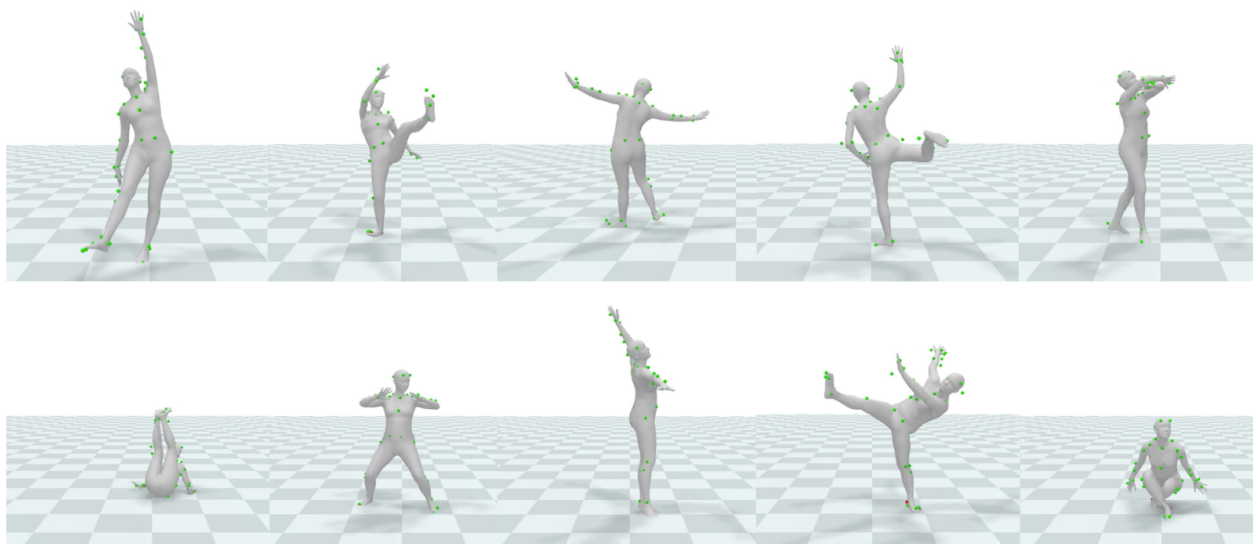


Figure 14: Example frames from using TestOpt with HuMoR to fit to full-body 3D keypoint motions from dynamic DanceDB [5] data in AMASS [64]. Estimated body shape and pose is shown along with ground truth keypoints in green. Despite not training on this dance data, HuMoR is able to effectively generalize to capture these complex motions.