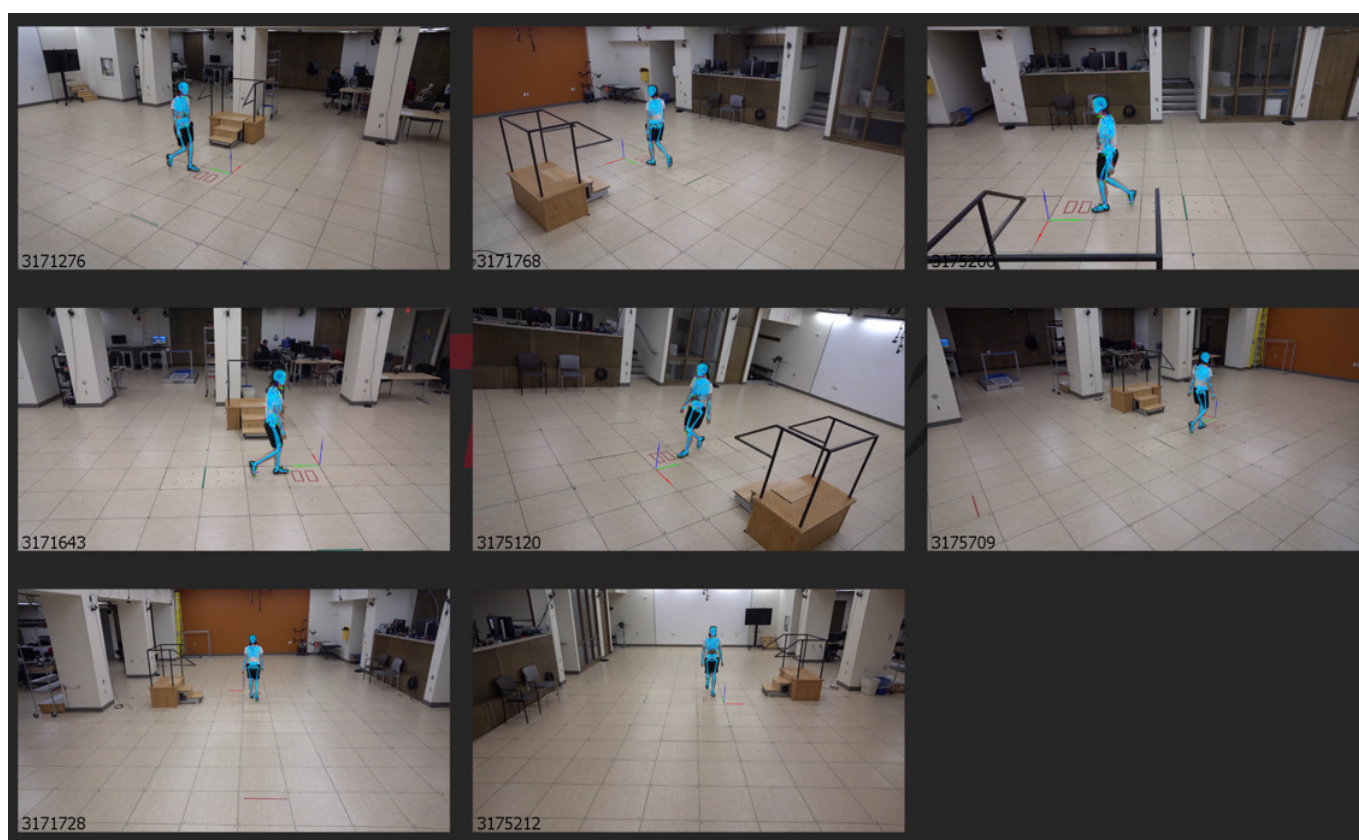
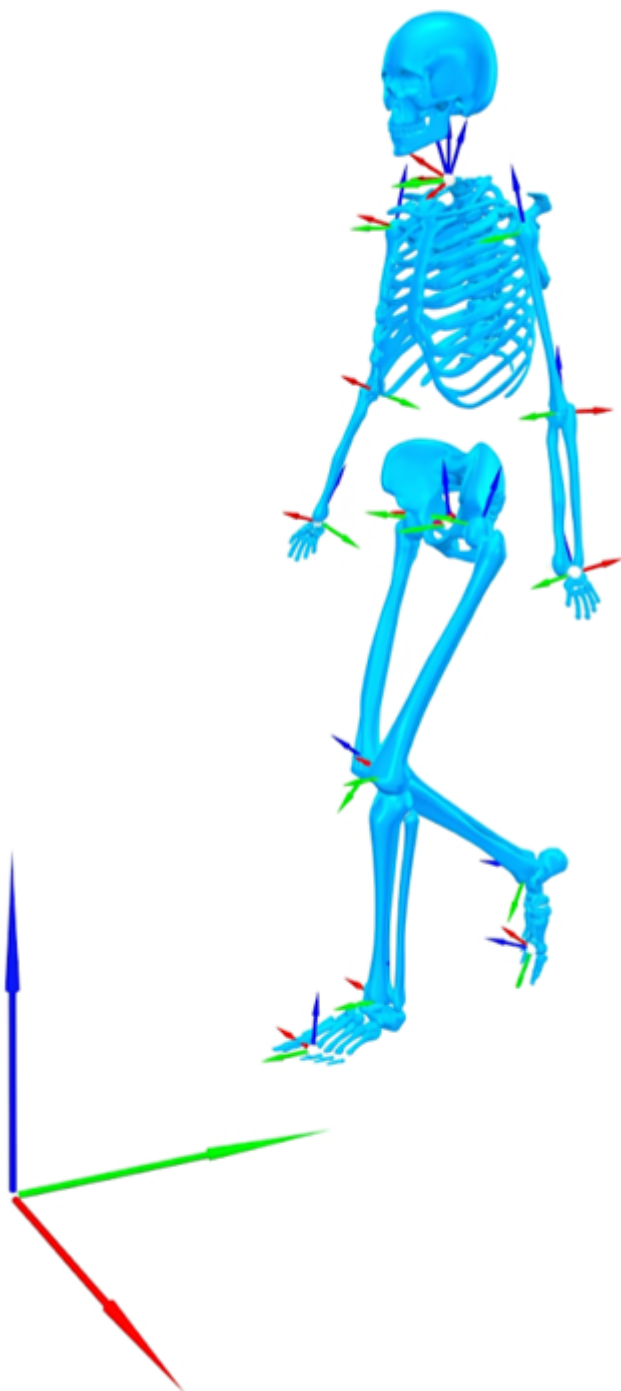


Markerless Motion Capture Overview

What is Markerless?

What most people think of when they hear “markerless motion capture”, is the system which uses vision-based deep learning to detect and identify anatomical landmarks (similar to - but not necessarily the same as - the physical markers in marker-based methods), in 2D image space (per camera, per frame). The 3D position of the subject is found through triangulation across multiple calibrated cameras. The output of these systems is therefore generally 4x4 pose matrices that fully represent the pose (3D position and orientation) of the associated [segment](#). Examples of vision-based markerless systems include [Theia](#) and [Fujitsu](#).





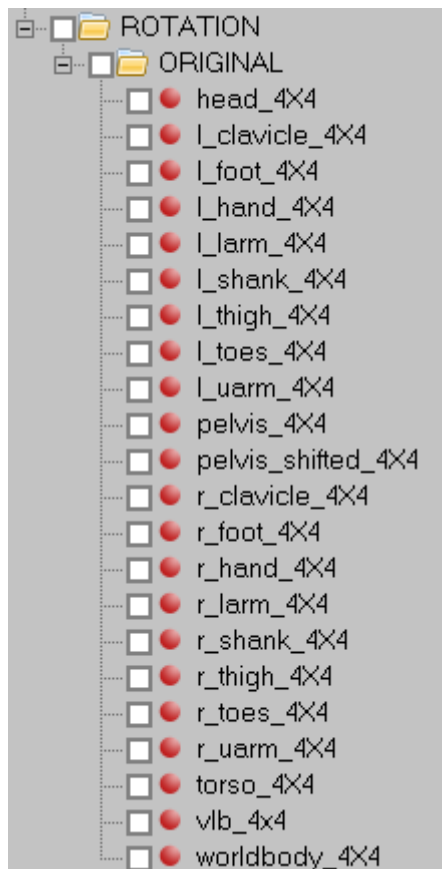
Some people also consider IMU's to be a markerless system. Check out the [IMU Overview Page](#) to learn more.

Markerless data in Visual3D

Automatic Model Building

Visual3D has automatic model building for [Theia](#), [Fujitsu](#), [XSens](#), [Kinatrax](#), and Hawkeye.

Visual3D builds the model from the body segment position and orientations, in the form of 4×4 transformation matrices. Segment pose can be found in the ROTATION folder under the main file tree.



Traditional [segments](#) are built using the proximal and distal (often the proximal end of the adjacent body segment) pose matrices, with the origin of the segment coordinate systems at the proximal end of the segment. The segment Z-Axis is defined from distal to proximal, using the right hand rule. All these traditional segments are used in the kinematic chain for [inverse dynamics](#).

[Virtual segments](#) enable an alternative definition for the segment. Though, the virtual segments cannot be included in the kinetic chain for inverse dynamics, as the virtual segment will behave differently than the traditional segment. To address this, Theia and XSens models have been implemented in Visual3D with the addition of [shadow segments](#). Shadow segments follow the definition of the virtual segment, but are tracked using a transformation of the original segments pose, so that it can be included in the kinetic chain.

Tutorials

[Manage File Merge](#) : combine motion capture files from the same recording session.

[Assesing Stability During Gait](#)

Learn More

[Deep learning](#)

Notes

Check the [release notes](#) for specific updates for how the systems models are being built. For example, shadow segments are added automatically to the model for the thorax and feet for Theia and XSens data.

From:
<https://has-motion.com/wiki/> - **Wiki Documentation Site**

Permanent link:
https://has-motion.com/wiki/doku.php?id=visual3d:documentation:markerless_motion_capture&rev=1783541900

Last update: **2026/07/08 20:18**

