

Prior Knowledge Based Motion Model Representation

Angel D. Sappa* Niki Aifanti⁺ Sotiris Malassiotis⁺ Michael G. Strintzis⁺

** Computer Vision Center, Edifici O, Campus UAB, 08193 Bellaterra - Barcelona, Spain*

+ Informatics & Telematics Institute, 1st Km Thermi-Panorama Road, Thermi-Thessaloniki, Greece

Received 20 December 2004; accepted 15 March 2005

Abstract

This paper presents a new approach for human walking modeling from monocular image sequences. A kinematics model and a walking motion model are introduced in order to exploit prior knowledge. The proposed technique consists of two steps. Initially, an efficient feature point selection and tracking approach is used to compute feature points' trajectories. Peaks and valleys of these trajectories are used to detect *key frames*—frames where both legs are in contact with the floor. Secondly, motion models associated with each joint are locally tuned by using those key frames. Differently than previous approaches, this tuning process is not performed at every frame, reducing CPU time. In addition, the movement's frequency is defined by the elapsed time between two consecutive key frames, which allows handling walking displacement at different speed. Experimental results with different video sequences are presented.

Key Words: Model based representations; Human walking modeling; Human motion tracking.

1 Introduction

3D human body representations opens a new and attractive field of applications, from more realistic movies to interactive environments. Unfortunately, it is not possible to completely recover 3D information from 2D video sequences when no other extra information is given or can be estimated (i.e. intrinsic and extrinsic camera parameters should be provided to compute the real 3D data up to a scale). However, since several video sequences are populated with objects with known structure and motion such as humans, cars, etc, prior knowledge would arguably aid the recovery of the scene. Prior knowledge in the form of kinematics constraints (average size of an articulated structure, degrees of freedom (DOFs) for each articulation), or motion dynamics (physical laws ruling the objects' movements), is a commonplace solution to handle the aforementioned problem.

In this direction, our work is devoted to depth augmentation of common human walking video sequences. Although this work is only focused on walking modeling, which is the most frequent type of locomotion of persons, other kinds of cyclic movement could also be modeled with the proposed approach (e.g. movements such as running or climbing down/upstairs).

3D motion models are required for applications such as: intelligent video surveillance, pedestrian detection for traffic applications, gait recognition, medical diagnosis and rehabilitation, human-machine

Correspondence to: <angel.sappa@cvc.uab.es> This work was supported in part by the Government of Spain under the CICYT project TRA2004-06702/AUT. The first author was supported by *The Ramón y Cajal Program*.

Recommended for acceptance by Perales F., Drapper B.

ELCVIA ISSN: 1577-5097

Published by Computer Vision Center / Universitat Autònoma de Barcelona, Barcelona, Spain

interface [1]. Due to the widely interest it has generated, 3D human motion modeling is one of the most active area within the computer vision community.

In this paper a new approach to cope with the problem of human walking modeling is presented. The main idea is to search for a particular kinematics configuration throughout the frames of the given video sequence, and then to use the extracted information in order to tune a general motion model. Walking displacement involves the *synchronized* movements of each body part—the same is valid for any cyclic human body displacement (e.g., running, jogging). In this work, a set of curves, obtained from anthropometric studies [2], is used as a coarse walking model. These curves need to be individually tuned according to the walking attitude of each pedestrian. This tuning process is based on the observation that although each person walks with a particular style, there is an instant in which every human body structure achieves the same configuration. This instant happens when both legs are in contact with the floor. Then, the open articulated structure becomes a closed structure. This closed structure is a rich source of information useful to tune most of the motion model's parameters. The outline of this work is as follows. Related works are presented in the next section. The proposed technique is described in section 3. Experimental results using different video sequences are presented in section 4. Conclusions and further improvements are given in section 5.

2 Previous Works

Vision-based human motion modeling approaches usually combine several computer vision processing techniques (e.g. video sequence segmentation, object tracking, motion prediction, 3D object representation, model fitting, etc.). Different techniques have been proposed to find a model that matches a walking displacement. These approaches can be broadly classified into monocular or multi camera approaches.

A multicamera system was proposed by [3]. It consists of a stereoscopic technique able to cope not only with self-occlusions but also with fast movements and poor quality images. This approach incorporates physical forces to each rigid part of a kinematics 3D human body model consisting of truncated cones. These forces guide each 3D model's part towards a convergence with the body posture in the image. The model's projections are compared with the silhouettes extracted from the image by means of a novel approach, which combines the Maxwell's demons algorithm with the classical ICP algorithm. Although stereoscopic systems provide us with more information for the scanned scenes, 3D human motion systems with only one camera-view available is the most frequent case.

Motion modeling using monocular image sequences constitutes a complex and challenging problem. Similarly to approach [3], but in a 2D space and assuming a segmented video sequence is given as an input, [4] proposes a system that fits a projected body model with the contour of a segmented image. This boundary matching technique consists of an error minimization between the pose of the projected model and the pose of the real body—all in a 2D space. The main disadvantage of this technique is that it needs to find the correspondence between the projected body parts and the silhouette contour, before starting the matching approach. This means that it looks for the point of the silhouette contour that corresponds to a given projected body part, assuming that the model posture is not initialized. This problem is still more difficult to handle in those frames where self-occlusions appear or edges cannot be properly computed.

Differently than the previous approaches, the aspect ratio of the bounding box of the moving silhouette has been used in [5]. This approach is able to cope with both lateral and frontal views. In this case the contour is studied as a whole and body parts do not need to be detected. The aspect ratio is used to encode the pedestrian's walking way. However, although shapes are one of the most important semantic attributes of an image, problems appear in those cases where the pedestrian wears clothes not so tight or carries objects such as a suitcase, handbag or backpack. Carried objects distort the human body silhouette and therefore the aspect ratio of the corresponding bounding box.

In order to be able to tackle some of the problems mentioned above, some authors propose simplifying assumptions. In [6] for example, tight-fitting clothes with sleeves of contrasting colors have been used. Thus, the right arm is depicted with a different color than the left arm and edge detection is simplified especially in case of self-occlusions. [7] proposes an approach where the user selects some points on the image, which mainly correspond to the joints of the human body. Points of interest are also marked in [8] using infrared

diode markers. The authors present a physics-based framework for 3D shape and non-rigid motion estimation based on the use of a non-contact 3D motion digitizing system. Unfortunately, when a 2D video sequence is given, it is not likely to affect its content afterwards in such a way. Therefore, the usefulness of these approaches is restricted to cases where access in making the sequence is possible.

Recently, a novel approach based on feature point selection and tracking was proposed in [9]. This approach is closely related to the technique proposed in this work. However, a main difference is that in [9] feature points are triangulated together and similarity between triangles and body parts is studied, while in the current work, feature point's trajectories are plotted on the image plane and used to detect key frames. Robustness in feature point based approaches is considerably better than in those techniques based on silhouette, since silhouette does not only depend on walking style or direction but also on other external factors such as those mentioned above. Walking attitude is easier captured by studying the spatio-temporal motion of feature points.

3 The Proposed Approach

We may safely assume that the center of gravity of a walking person moves with approximately constant velocity. However, the speed of other parts of the body fluctuates. There is one instant per walking cycle (without considering the starting and ending positions) in which both feet are in contact with the floor, in other words with null velocity. This happens when the pedestrian changes from one pivot foot to the other. At that moment the articulated structure (Fig. 1) reaches the maximum hip angles. Frames containing these configurations will be called *key frames* and can be easily detected by extracting static points (i.e. pixels defining the boundary of the body shape contained in the segmented frames that remain static at least in three consecutive frames) through the given video sequence. Information provided by these frames is used to tune motion model parameters. In addition, the elapsed time between two consecutive key frames defines the duration of a half walking period—indirectly the speed of the movement. Motion models are tuned and used to perform the movement from the current key frame to the next one. This iterative process is applied until all key frames are covered by the walking model (input video sequence).

Human motion modeling based on tracking of point features has recently been used in [9]. Human motion is modeled by the joint probability density function of the position and velocity of a collection of body parts, while no information about kinematics or dynamics of the human body structure is considered. This technique has been tested only on video sequences containing pedestrians walking on a plane orthogonal to the camera's viewing direction.

Assuming that a segmented video sequence is given as an input (in the current implementation some segmented images were provided by the authors of [10] and others were computed by using the algorithm presented in [11]) the proposed technique consists of two stages. In the first stage feature points are selected and tracked throughout the whole video sequence in order to find key frames' positions. In the second stage a generic motion model is locally tuned by using kinematics information extracted from the key frames. The main advantage comparing with previous approaches is that matching between the projection of the 3D model and the body silhouette image features is not performed at every frame (e.g., hip tuning is performed twice per walking cycle). The algorithm's stages are fully described below together with a brief description of the 3D representation used to model the human body.

3.1 Body Modeling

Modeling the human body implies firstly the definition of a 3D articulated structure, which represents the body's biomechanical features; and secondly the definition of a motion model, which governs the movement of that structure.

Several 3D articulated representations have been proposed in the literature. Generally, a human body model is represented as a chain of rigid bodies, called links, interconnected to one another by *joints*. Links are generally represented by means of sticks [7], polyhedron [12], generalized cylinders [13] or superquadrics [6]. A joint interconnects two links by means of rotational motions about the axes. The number of independent rotation parameters defines the DOFs associated with that joint.

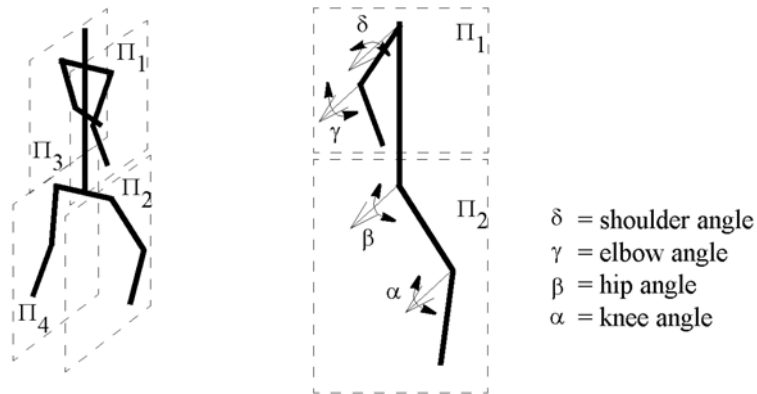


Figure 1. Simplified articulated structure defined by 12 DOFs, arms and legs rotations are contained in planes parallel to the walking's direction.

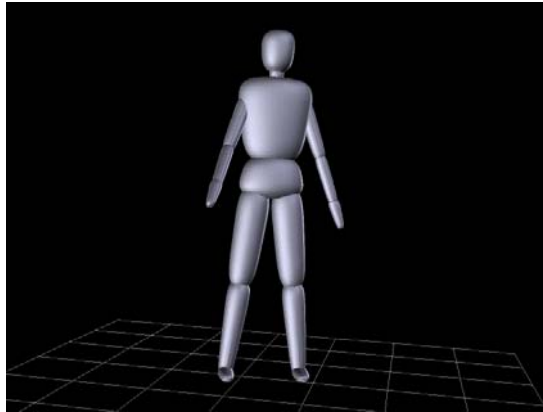


Figure 2. Illustration of a 22 DOF model built with superquadric.

Considering that a human body has about six hundred muscles, forty just for a hand, the development of a highly realistic model is a computational expensive task, involving a high dimensionality problem. In computer vision, where models with only medium precision are required, articulated structures with less than thirty degrees of freedom are generally adequate. (e.g. [6], [3]). In this work, an articulated structure defined by 16 links is initially considered. This model consists of 22 DOF, without modeling the palm of the hand or the foot and using a rigid head-torso approximation (four for each arm and leg and six for the torso, which are three for orientation and three for position). However, in order to reduce the complexity, a simplified model of 12 DOF has been finally chosen. This simplification assumes that in walking, legs' and arms' movements are contained in parallel planes (see illustration in Fig. 1). In addition, the body orientation is always orthogonal to the floor, thus the orientation is reduced to only one DOF. Hence, the final model is defined by two DOF for each arm and leg and four for the torso (three for the position plus one for the orientation).

The simplest 3D articulated structure is a stick representation with no associated volume or surface. Planar 2D representations, such as a cardboard model, have been also widely used. However, volumetric representations are preferred in order to generate more realistic models. Volumetric representations such as parallelepipeds, cylinders, or superquadrics have been extensively used. For example, [3], proposes to model a person by means of truncated cones (arms and legs), spheres (neck, joints and head) and right parallelepipeds (hands, feet and body). Most of these shapes can be modeled by means of superquadrics [14]. Superquadrics are a compact and accurate representation generally used to model human body parts.

Through this work the articulated structure will be represented by 16 superquadrics, see Fig. 2 ([15]). A superquadric surface is defined by the following parametric equation:

$$x(\theta, \phi) = \begin{bmatrix} \alpha_1 \cos^{\varepsilon_1}(\theta) \cos^{\varepsilon_2}(\phi) \\ \alpha_2 \cos^{\varepsilon_1}(\theta) \sin^{\varepsilon_2}(\phi) \\ \alpha_3 \sin^{\varepsilon_1}(\theta) \end{bmatrix} \quad (1)$$

where $(-\pi/2) \leq \theta \leq (\pi/2)$, $-\pi \leq \phi \leq \pi$. The parameters α_1 , α_2 and α_3 define the size of the superquadric along the x , y and z axis respectively, while ε_1 is the squareness parameter in the latitude plane and ε_2 is the squareness parameter in the longitudinal plane. Furthermore, superquadric shapes can be deformed with tapering, bending and cavities. In our model, the different body parts are represented with superquadrics tapered along the y -axis. The parametric equation is then written as:

$$x'(\theta, \phi) = \begin{bmatrix} \left(\frac{t_1}{\alpha_2} x_2 + 1 \right) x_1 \\ x_2 \\ \left(\frac{t_3}{\alpha_2} x_2 + 1 \right) x_3 \end{bmatrix} \quad (2)$$

where $-1 \leq t_1, t_3 \leq 1$ are the tapering parameters and x_1 , x_2 and x_3 are the elements of the vector in equation (1). Parameters α_1 , α_2 and α_3 were defined in each body part according to anthropometric measurements. An example can be seen in Fig. 2.

The movements of the limbs are based on a hierarchical approach (the torso is considered the root) using Euler angles. The body posture is synthesized by concatenating the transformation matrices associated with the joints, starting from the root.

3.2 Feature Point Selection and Tracking

Feature point selection and tracking approaches were chosen because they allow capturing the motion's parameters by using as prior knowledge the kinematics of the body structure. In addition, point-based approaches seem to be more robust in comparison with silhouette based approaches. Next, a brief description of the techniques used is given.

3.2.1 Feature Point Selection

In this work, the feature points are used to capture human body movements and are selected by using a corner detector algorithm. Let $I(x,y)$ be the first frame of a given video sequence. Then, a pixel (x,y) is a corner feature if at all pixels in a window W_s around (x,y) the smallest singular value of G is bigger than a predefined σ ; in the current implementation W_s was set to 5×5 and $\sigma = 0.05$. G is defined as:

$$G = \begin{bmatrix} \sum I_x^2 & \sum I_x I_y \\ \sum I_x I_y & \sum I_y^2 \end{bmatrix} \quad (3)$$

and (I_x, I_y) are the gradients obtained by convolving the image I with the derivatives of a pair of Gaussian filters. More details about corner detection can be found in [16]. Assuming that at the beginning there is no information about the pedestrian's position in the given frame, and in order to enforce a homogeneous feature sampling, input frames are partitioned into 4 regular tiles (2×2 regions of 240×360 pixels each in the illustration presented in Fig. 3).

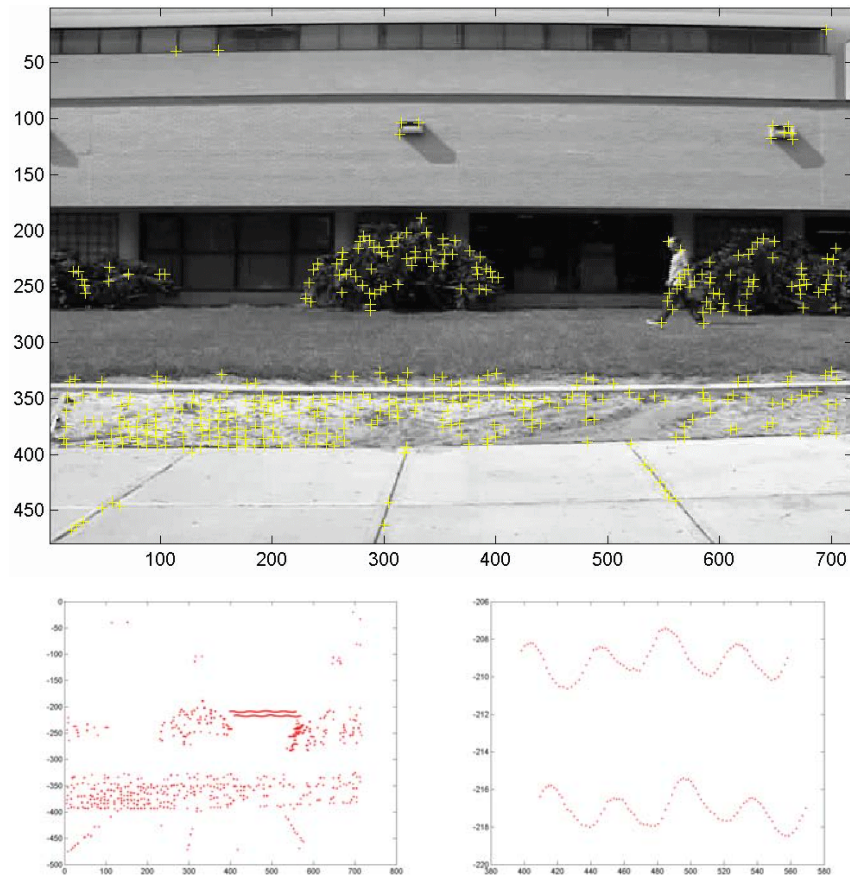


Figure 3. (top) Feature points from the first frame of the video sequence used in [17]. (bottom-left) Feature points' trajectories. (bottom-right) Feature points' trajectories after removing static points.

3.2.2 Feature Point Tracking

After selecting a set of feature points and setting a tracking window W_T (3×3 in the current implementation) an iterative feature tracking algorithm has been used [16]. Assuming a small interframe motion, feature points are tracked by minimizing the sum of squared differences between two consecutive frames.

Points, lying on the head or shoulders, are the best candidates to satisfy the aforementioned assumption. Most of the other points (e.g. points over the legs, arms or hands, are missed after a couple of frames). Fig. 3(top) illustrates feature points detected in the first frame of the video sequence used in [17]. Fig. 3(bottom-left) depicts the trajectories of the feature points when all frames are considered. On the contrary, Fig. 3(bottom-right) shows the trajectories after removing static points. In the current implementation we only use one feature point's trajectory. Further improvements could be to merge feature points' trajectories in order to generate a more robust approach.

3.3 Motion Model Tuning

The outcome of the previous stage is the trajectory of a feature point (Fig. 4(top)) consisting of peaks and valleys. Firstly, the first-order derivative of the curve is computed to find peaks' and valleys' positions by seeking the positive-to-negative zero-crossing points. Peaks correspond to those frames where the pedestrian reaches the maximum height, which happens in that moment of the half walking cycle when the hip angles

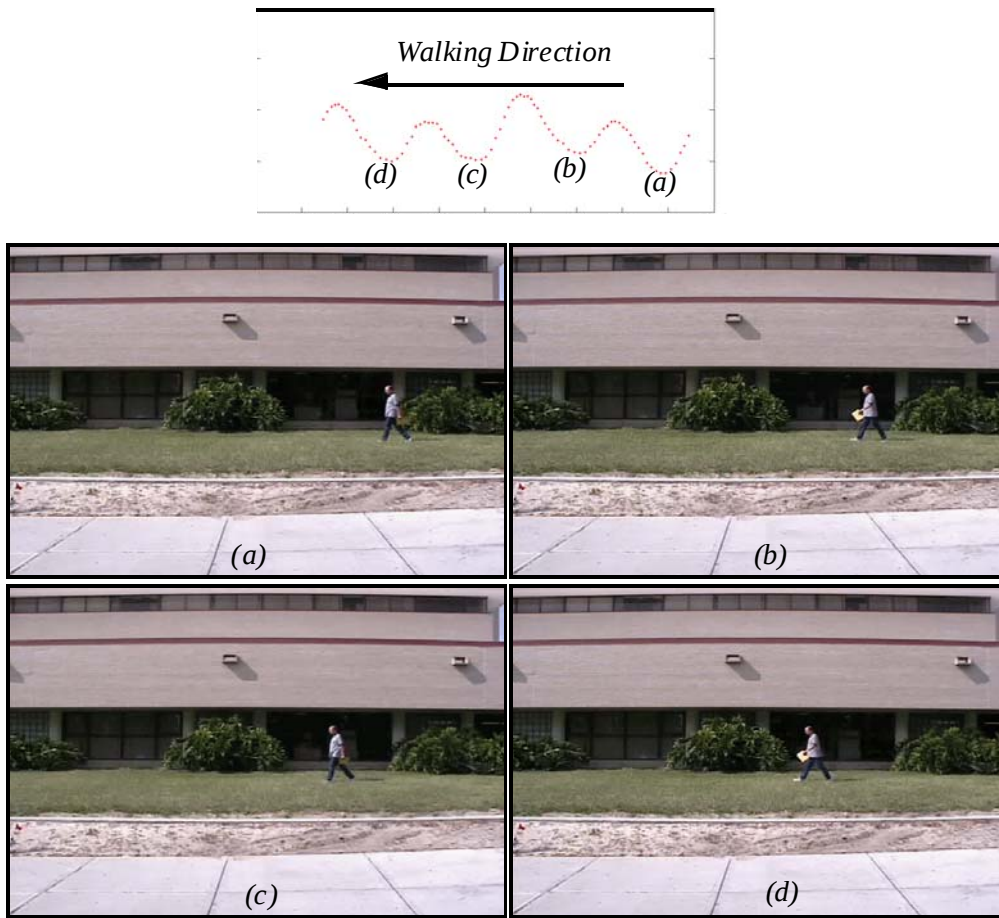


Figure 4. (top) A single feature point's trajectory. (middle and bottom) Key frames associated with the valleys of a feature point's trajectory.

are minimum. On the contrary, the valleys correspond to those frames where the two legs are in contact with the floor and then, the hip angles are maximum. So, the valleys are used to find key frames, while the peaks are used for footprint detection. The frames corresponding to each valley of Fig. 4(top) are presented in Fig. 4(middle) and (bottom). An interesting point of the proposed approach is that in this video sequence, in spite of the fact that the pedestrian is carrying a folder, key frames are correctly detected and thus, the 3D human body configuration can be computed. On the contrary, with an approach such as [4], it will be difficult since the matching error will try to minimize the whole shape (including folder).

After detecting key frames, which correspond to the valleys of the trajectory, it is necessary to define also the footprints of the pedestrian throughout the sequence. In order to achieve this, body silhouettes were computed using an image segmentation algorithm [11]. For some video sequences the segmented images were provided by [10]. Footprint positions are computed as follows.

Throughout a walking displacement sequence, there is always, at least, one foot in contact with the ground, with null velocity (pivot foot). In addition, there is one instant per walking cycle in which both feet are in contact with the floor (both with null velocity). The foot that is in contact with the floor can be easily detected by extracting its defining *static points*. A point is considered as a static point $spt_{(i,j)}^F$ in frame F , if it remains as a boundary point $bp_{(i,j)}^F$ (silhouette point, Fig. 5(left)) in at least three consecutive frames—value computed experimentally $stp_{(i,j)}^F \Rightarrow (bp_{(i,j)}^{F-1}, bp_{(i,j)}^F, bp_{(i,j)}^{F+1})$.

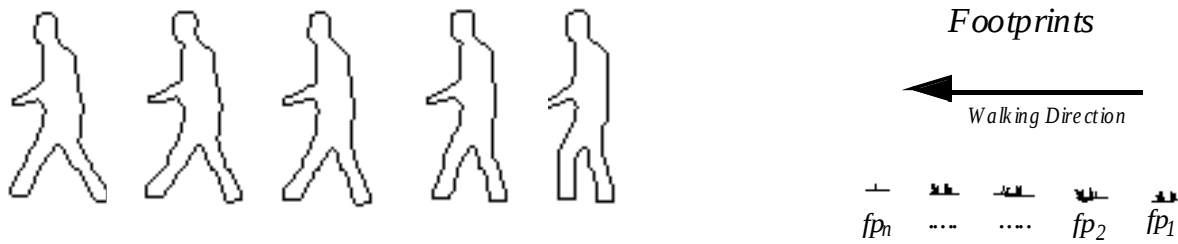


Figure 5. (left) Five consecutive frames used to detect static points. (right) Footprints computed after clustering static points generated by the same foot (peaks in a feature point's trajectory (Fig. 4(top))).

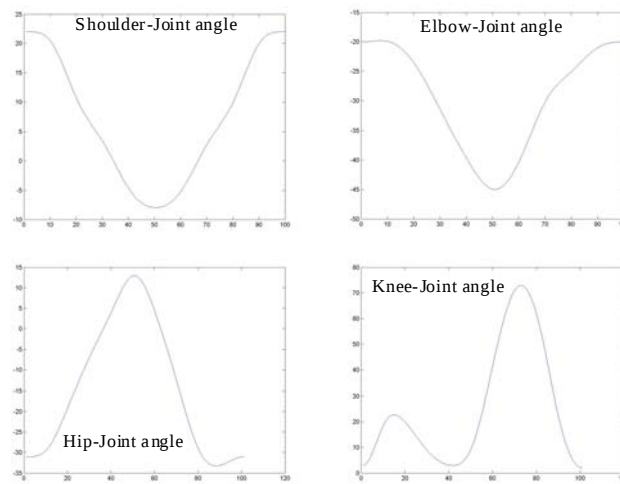


Figure 6. Motion curves of the joints at the shoulder, elbow, hip and knee (computed from [2]).

The result of the previous stage is a set of static points distributed along the pedestrian's path. Now, the problem is to cluster those points belonging to the same foot. Static points defining a single footprint are easily clustered by studying the peaks' positions in the feature point's trajectory. All those static points in a neighborhood of $F \pm 3$ from the frame corresponding to a peak position (F) will be clustered together and will define the same footprint (fp_i). Fig. 5(right) shows an illustration of static points detected after processing consecutive frames.

As it was introduced above, key frames are defined as those frames where both feet are in contact with the floor. At every key frame, the articulated human body structure reaches a posture with maximum hip angles. In the current implementation, hip angles are defined by the legs and the vertical axis containing the hip joints. This maximum value, together with the maximum value of the hip motion model (Fig. 6) are used to compute a scale factor κ . This factor is utilized to adjust the hip motion model to the current pedestrian's walking. Actually, it is used for half the walking cycle, which does not start from the current key frame but from a quarter of the walking cycle before the current key frame until halfway to the next one. The maximum hip angle in the next key frame is used to update this scale factor.

This local tuning, within a half walking cycle, is illustrated with the 2D articulated structure shown in Fig. 7, from Posture 1 to Posture 3. A 2D articulated structure was chosen in order to make the understanding easier, however the tuning process is carried out in the 3D space. The two footprints of the first key frame are represented by the points **A** and **B**, while the footprints of the next key frame are the corresponding points **A''** and **B''**. During this half walking cycle one foot is always in contact with the floor (so points **A** =

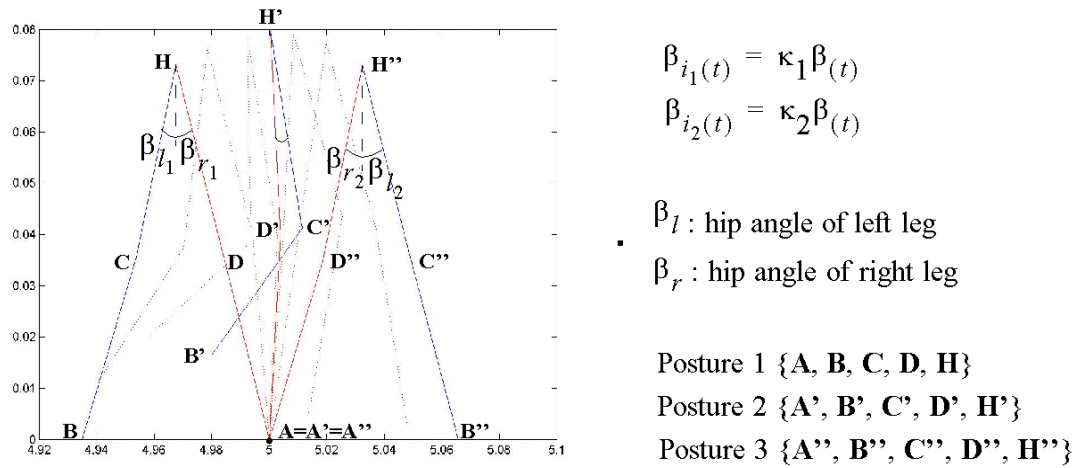


Figure 7. Half walking cycle executed by using scale factors (κ_1, κ_2) over the hip motion curve presented in Fig. 6 (knee motion curve is not tuned at this stage). Spatial positions of points (D, H, C and B) are computed by using angles from the motion curves and trigonometric relationships.

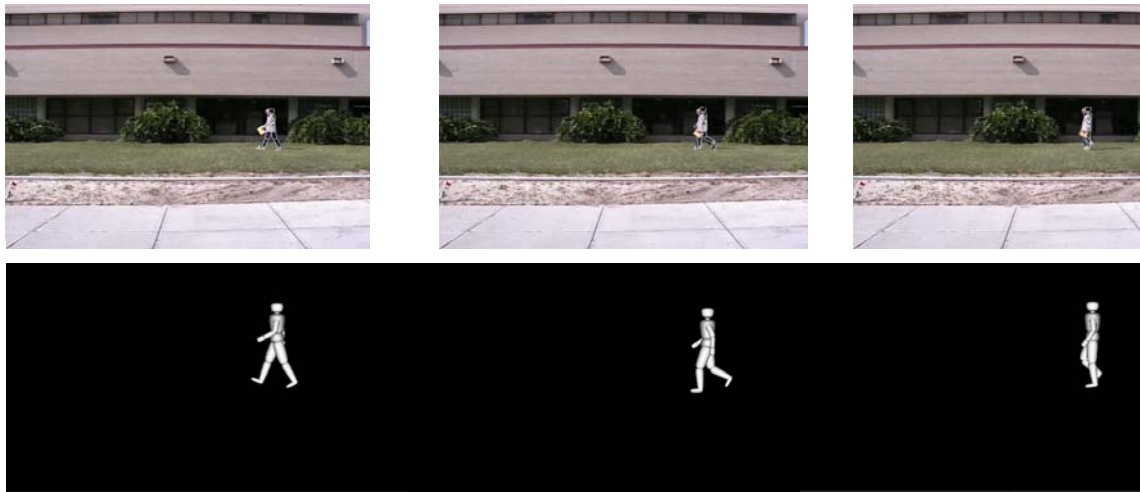


Figure 8. (top) Three different frames of the video sequence used in [17]. (bottom) The corresponding 3D walking models.

$A' = A''$), while the other leg is moving from point **B** to point **B''**. In halfway to **B''**, the moving leg crosses the other one (null hip angle values). Points **C**, **C'**, **C''** and **D**, **D'**, **D''** represent the left and right knee, while the points **H**, **H'**, **H''** represent the hip joints.

Given the first key frame, the scale factor κ_1 is computed and used to perform the motion ($\beta_{i_1(t)}$) through the first quarter of the walking cycle. The second key frame (**A''**, **B''**) is used to compute the scale factor κ_2 . At each iteration of this half walking cycle, the spatial positions of the points **B**, **C**, **D** and **H** are calculated using the position of point **A**, which remains static, the hip angles of Fig. 6 scaled by the corresponding factor κ_i and the knee angles of Fig. 6. The number of frames in between the two key frames defines the sampling rate of the motion curves presented on Fig. 6. This allows handling variations in the walking speed.

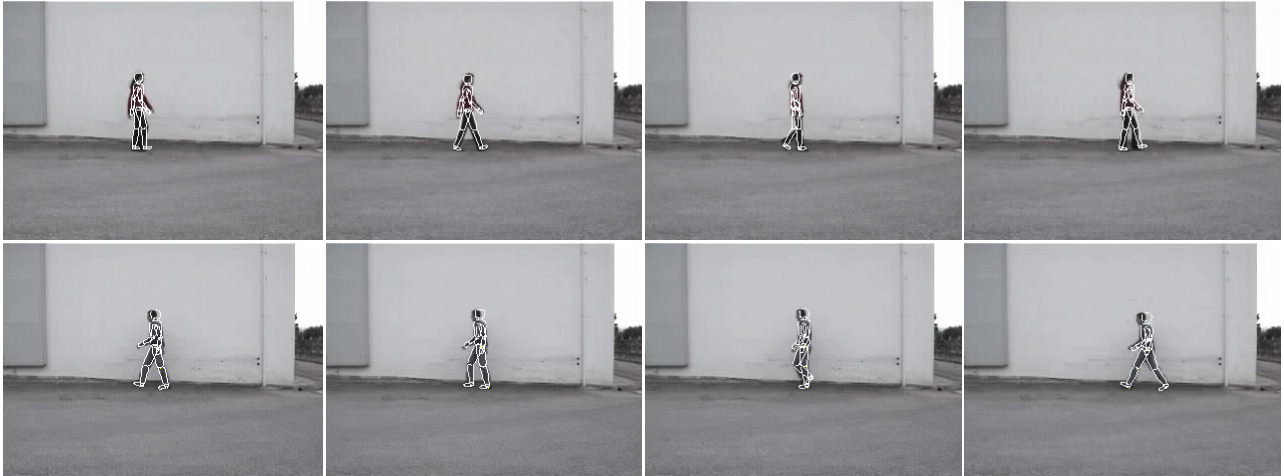


Figure 9. Input frames of two video sequences (240x320 each).

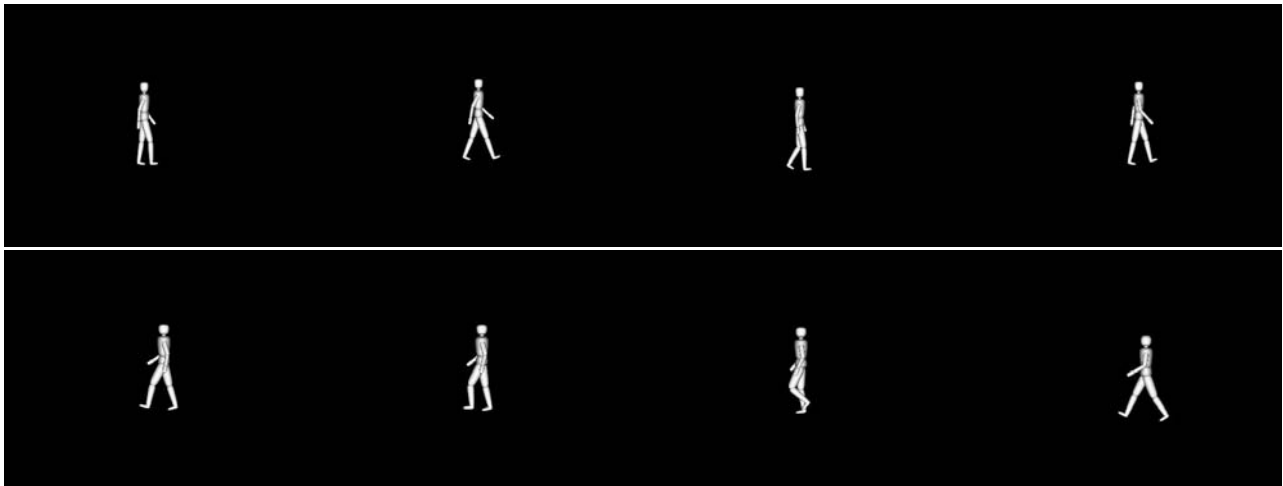


Figure 10. 3D models corresponding to the frames presented in Fig. 9 (top) and (bottom) respectively.

As aforementioned, the computed factors κ_i are used to scale the hip angles. The difference in walking between people implies that all the motion curves should be modified by using an appropriate scale factor for each one. In order to estimate these factors an error measurement (registration quality index: *RQI*) is introduced. The proposed *RQI* measures the quality of the matching between the projected 3D model and the corresponding human silhouette. It is defined as: $RQI = \text{overlappedArea} / \text{totalArea}$, where total area consists of the surface of the projected 3D model plus the surface of the walking human figure less the overlapped area, while the overlapped area is defined by the overlap of these two surfaces. Firstly, the algorithm computes the knee scale factor that maximizes the *RQI* values. In every iteration, an average *RQI* is computed for all the sequence. In order to speed up the process the number of frames was subsampled. Afterwards, the elbow and shoulder scale factors are estimated similarly. They are computed simultaneously using an efficient search method.

4 Experimental Results

The proposed technique has been tested with video sequences used in [17] and [10], together with our own video sequences. Despite that the current approach has been developed to handle sequences with a pedestrian walking over a planar surface, in a plane orthogonal to the camera direction, the technique has been also tested with an oblique walking direction (see Fig. 11) showing encouraging results. The video sequence used

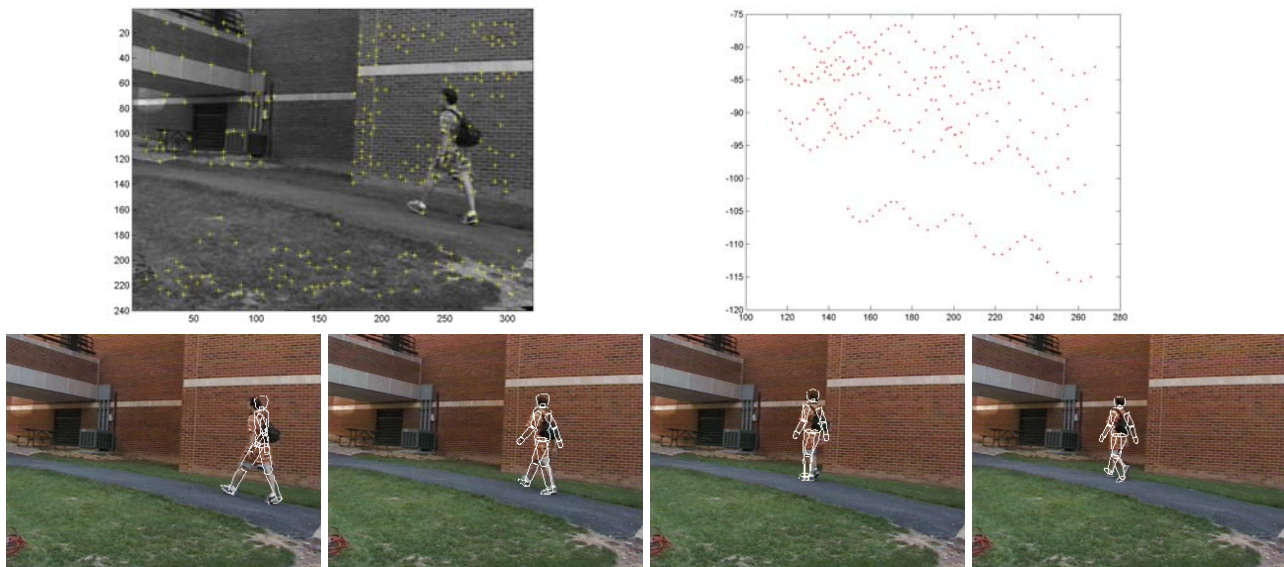


Figure 11. (top-left) Feature points of the first frame. (top-right) Feature points' trajectories. (bottom) Some frames illustrating the final result (segmented input has been provided by [10]).

as an illustration throughout this work consists of 85 frames of 480×720 pixels each, which have been segmented using the technique presented in [11]. Some of the computed 3D walking models are presented in Fig. 8(bottom), while the original frames together with the projected boundaries are presented in Fig. 8(top).

Fig. 9(top) presents a few frames of a video sequence defined by 103 frames (240×320 pixels each), while Fig. 9(bottom) corresponds to a video sequence defined by 70 frames (240×320 pixels each). Although the speed and walking style is considerably different, the proposed technique can handle both situations. The corresponding 3D models are presented in Fig. 10 (top) and (bottom) respectively.

Finally, the proposed algorithm was also tested on a video sequence, consisting of 70 frames of 240×320 pixels each, containing a diagonal walking displacement (Fig. 11). The segmented input frames have been provided by the authors of [10]. Although the trajectory was not on a plane orthogonal to the camera direction, feature point information was enough to capture the pedestrian's attitude.

5 Conclusions and Future Work

A new approach towards human motion modeling has been presented. It exploits prior knowledge regarding a person's movement as well as human body kinematics constraints. At this paper only walking has been modeled. Although constraints about walking direction and planar surfaces have been imposed, we expect to extend this technique in order to include frontal and oblique walking directions [18]. A preliminary result has been presented in Fig. 11.

Modeling other kinds of human body cyclic movements (such as running or going up/down stairs) using this technique constitutes a possible extension and will be studied. In addition, the use of a similar approach to model the displacement of other articulated beings (animals in general [19]) will be studied. Animal motion (i.e. cyclic movement) can be understood as an open articulated structure, however, when more than one extremity is in contact with the floor, that structure becomes a closed kinematics chain with a reduced set of DOFs. Therefore a motion model could be computed by exploiting these particular features.

Further work will also include the tuning of not only motion model's parameters but also geometric model's parameters in order to find a better fitting. In this way, external objects attached to the body (like a handbag or backpack) could be added to the body and considered as a part of it.

6 References

- [1] Sappa, N. Aifanti, N. Grammalidis and S. Malassiotis, "Advances in Vision-Based Human Body Modeling", Chapter book in: *3D Modeling and Animation: Synthesis and Analysis Techniques for the Human Body*, N. Sarris and M.G. Strintzis (Eds.), Idea-Group Inc., 2004.
- [2] K. Rohr, "Human Movement Analysis Based on Explicit Motion Models", Chapter 8 in *Motion-Based Recognition*, M. Shah and R. Jain (Eds.), Kluwer Academic Publisher, Dordrecht Boston 1997, pp. 171-198.
- [3] Q. Delamarre and O. Faugeras, "3D Articulated Models and Multi-View Tracking with Physical Forces", Special Issue on Modelling People, *Computer Vision and Image Understanding*, Vol. 81, 328-357, 2001.
- [4] H. Ning, T. Tan, L. Wang and W. Hu, "Kinematics-Based Tracking of Human Walking in Monocular Video Sequences", *Image and Vision Computing*, Vol. 22, 2004, 429-441.
- [5] L. Wang, T. Tan, W. Hu and H. Ning, "Automatic Gait Recognition Based on Statistical Shape Analysis", *IEEE Trans. on Image Processing*, Vol. 12 (9), September 2003, 1-13.
- [6] D. Gavrila and L. Davis, "3-D Model-Based Tracking of Humans in Action: a Multi-View Approach", *IEEE Int. Conf. on Computer Vision and Pattern Recognition*, San Francisco, USA, 1996.
- [7] Barron and I. Kakadiaris, "Estimating Anthropometry and Pose from a Single Camera", *IEEE Int. Conf. on Computer Vision and Pattern Recognition*, Hilton Head Island, USA, 2000.
- [8] Metaxas and D. Terzopoulos, "Shape and Nonrigid Motion Estimation through Physics-Based Synthesis", *IEEE Trans. on Pattern Analysis and Machine Intelligence*, Vol. 10 (6), June 1993, 580-591.
- [9] Y. Song, L. Goncalves and P. Perona, "Unsupervised Learning of Human Motion", *IEEE Trans. on Pattern Analysis and Machine Intelligence*, Vol. 25 (7), July 2003, 1-14.
- [10] S. Jabri, Z. Duric, H. Wechsler and A. Rosenfeld, "Detection and Location of People in Video Images Using Adaptive Fusion of Color and Edge Information", *15th. Int. Conf. on Pattern Recognition*, Barcelona, Spain, Sep. 2000.
- [11] Kim and J. Hwang, "Fast and Automatic Video Object Segmentation and Tracking for Content-Based Applications", *IEEE Trans. on Circuits and Systems for Video Technology*, Vol. 12 (2), February 2002, 122-129.
- [12] M. Yamamoto, A. Sato, S. Kawada, T. Kondo and Y. Osaki, "Incremental Tracking of Human Actions from Multiple Views", *IEEE Int. Conf. on Computer Vision and Pattern Recognition*, Santa Barbara, CA, USA, 1998.
- [13] Cohen, G. Medioni and H. Gu, "Inference of 3D Human Body Posture from Multiple Cameras for Vision-Based User Interface", *World Multiconference on Systemic, Cybernetics and Informatics*, USA, 2001.
- [14] Solina and R. Bajcsy, "Recovery of Parametric Models from Range Images: The Case for Superquadrics with Global Deformations", *IEEE Trans. on Pattern Analysis and Machine Intelligence*, Vol. 12, No. 2, February 1990.
- [15] Sappa, N. Aifanti, S. Malassiotis and M. Strintzis, "Monocular 3D Human Body Reconstruction Towards Depth Augmentation of Television Sequences", *IEEE Int. Conf. on Image Processing*, Barcelona, Spain, Sep. 2003.
- [16] Y. Ma, S. Soatto, J. Kosecká and S. Sastry, *An Invitation to 3-D Vision: From Images to Geometric Models*, Springer-Verlag New York, 2004.
- [17] P. Phillips, S. Sarkar, I. Robledo, P. Grother and K. Bowyer, "Baseline Results for the Challenge Problem of Human ID Using Gait Analysis", *IEEE Int. Conf. on Automatic Face and Gesture Recognition*, Washington, USA, May 2002.

- [18] L. Wang, T. Tan, H. Ning and W. Hu, "Silhouette Analysis-Based Gait Recognition for Human Identification", *IEEE Trans. on Pattern Analysis and Machine Intelligence*, Vol. 25 (12), December 2003, 781-796.
- [19] P. Schneider and J. Wilhelms, "Hybrid Anatomically Based Modeling of Animals", *IEEE Computer Animation'98*, Philadelphia, USA, June 1998.