

Research on the Contour Capture Algorithm of the Key Action of Dance with Computer-Aided

Jin Zhang
Xin Li

In order to solve the problem that the current key dance motion contour capture algorithm cannot effectively capture the concave part of the key dance motion contour, which leads to large contour capture error and long response time, a computer-aided key dance motion contour capture algorithm is proposed. The action sequence without background is obtained from the dance video, and the action sequence is optimized to reduce the interference of contour capture. Algorithm tracks the posture change of the object in the dance action sequence, determines the contour capture area, and completes the contour capture of the dance key action. Experimental results show that the proposed method can effectively improve the capture accuracy and shorten the response time.

Keywords: computer-aided, dance key movements, movement contour, contour capturing algorithm

Tob Regul Sci.™ 2021;7(5):1160-1169

DOI: doi.org/10.18001/TRS.7.5.34

INTRODUCTION

Dance movement is based on the dancer's body language expression, through a certain plot narrative, emotional expression or both in a unified way. In the practice of dance training with human movement as the core, the coordination of the dancer's movement, the correct degree of force way, muscle cooperation and so on not only affect the artistic expression of the dancer, the standard of movement display, but also have an important impact on the dancer's own body condition. In the dance performance, the dancer's key dance movements are the highlight of the performance, the display of the dancer's artistic skills, and the key to the performance of dance art. Therefore, the standardization of key dance movements is crucial in daily dance training¹. With the advancement of research on human motion gestures in the field of vision, both computer graphics and vision researchers can obtain dance movement data from the contour changes of dance movements and analyze the movement changes of dancers' joints in order to guide and correct the movements of dance trainers and enhance the learning effect of dancers. The

movement data of the joints on the target are captured by mechanical electric devices mentioned in the reference², and after processing the spatial displacement data collected by the sensors, the movement contours are captured by drawing. This algorithm is simple to implement and has good results for static shape capturing, but when applied to dance movement capturing, it will restrict affect the standardization of the dancer's movements and has a large limitation of use. The traditional image segmentation-based motion contour capture algorithm mentioned in the reference³ uses the rough capture of the overall motion contour of the human body by locating the person in the image and segmenting the human body parts in the image. This algorithm cannot achieve real-time processing of motion transformations and has a large capture error for some more complex dance movements. The motion contour capture algorithm mentioned in the reference⁴ uses the snake model to capture the contour of the motion image. However, because the snake model only involves the gradient information change of the pixel boundary of the target object in the image during processing, resulting in the algorithm's inability to

capture the contour of the recessed part of the captured object; secondly, the model is difficult to handle the topological changes of the image, and, because of its limited force capture range, its initial contour can only be placed near the target, which seriously affects the accuracy of capturing the contour of dance movements and the reliability of the algorithm. The above-mentioned traditional movement contour capturing algorithms cannot accurately capture the dynamic dance movement contour due to the improper processing of the target image and the limitations of the algorithm itself. In order to help dancers master dance movements accurately, improve the standard of their movements and reduce the probability of injury caused by non-standard movements, it is necessary to improve the movement contour capture algorithm by using advanced computer technology.

At present, the international common Laban dance score system will write down the dance movement sequences from one-dimensional to three-dimensional art space with a paper plane score. In addition, the force-effective movement analysis proposed in Laban dance score system, from the dance body itself, proposes the embodiment of life science, geometric law and movement development state of the dance body in movement, which constitutes the practical basis of dance movement development. Through the scientific description of human movement sequences and the enrichment of force-effect content in dance movements, it is clear that the dancer's body and dance movements are closely related to the life science of the human body, that is, they are guided by physiology and anatomy, thus helping to solve the most fundamental problem of human movement in the dancer's dance art practice. At the same time, with the Laban dance score system and human movement-related disciplines as the research background, physical dance in the era of computer-based multimedia art is essentially inseparable from human movement itself, and human movement thus becomes the core of multimedia dance research^{5,6}. The mainstream dance styles in today's world are guided by the transition from real human performance to rigid human performance, thus enabling the audience to gradually form the ability to visually interpret

art and achieve an understanding of the emotions expressed by the dancer. In the process of dance performance, the transition from real human movement to steel body performance lays the foundation for computer-assisted dance art practice. There is a developmental continuity between computer-assisted, computerized dance practice that presents a virtual human body with geometric structure as an object, and is an important object for today's study of the dance human body and dance practice trends. Based on the above analysis content, this paper will study the algorithm of capturing the key movement contour of dance under computer-assisted, and design relevant test content to verify the feasibility of this algorithm.

COMPUTER-ASSISTED ALGORITHM FOR CAPTURING KEY MOVEMENT CONTOURS IN DANCE

Acquisition of Background-free Dancer Movement Sequences

The object of the dance key action contour capture algorithm is a video image of a dance action demonstration or a high frame image of a static capture, and the background in the image will affect the recognition and contour capture of the dance key action. Therefore, the contour capturing algorithm firstly needs to acquire the dancers' action sequences without background. For the background of the project under study, the motion target segmentation method with faster segmentation speed is selected to facilitate real-time motion information acquisition.

The basic idea is to subtract the pixel value of the current frame image from the corresponding pixel value of the background image stored in advance or obtained in real time, and if the difference is greater than a certain threshold, the pixel is judged to belong to the motion target, otherwise the pixel is judged to belong to the scene background⁷⁻⁹. The motion foreground mask obtained after closed-value segmentation and binarization directly gives the position, size, shape and other information of the motion target. The disadvantage is that it is easily affected by changes in external conditions such as light and weather. Therefore the obtained motion foreground mask needs to be post-processed to get the ideal detection effect. The post-processing

of the motion foreground mask is mainly to eliminate noise and uninteresting motion targets, and when the scene background color is similar to the motion target color, the detected motion foreground will have voids inside, which needs to be analyzed and processed in the post-processing stage to reduce the missed detection.

The main purpose of background modeling, also known as background estimation, is to transform the motion target detection of video frame images into a binary classification problem based on the current estimated background, classifying all pixels into two categories: background and motion foreground, and then post-processing the classification results to obtain the final detection results. Processing of dance key motion images based on the principle of Gaussian background modeling

For each pixel point in the dance key action image, the change of its value in the sequence image can be regarded as a random process that continuously generates pixel values, which can be expressed as

$$\{x_1, x_2, \dots, x_t\} = \{I(x_0, y_0, i) | 1 \leq i \leq t\} \quad (1)$$

In formula (1), I is the pixel value; x_0 and y_0 are the horizontal and vertical coordinates of the pixel point, respectively; I is the frame number of the sequence image. In the hybrid Gaussian background model, it is considered that the color information between pixels is not correlated with each other, and the processing of each pixel point is independent of each other. If not specified, the description of the background model in this paper is for the same pixel point.

The key of background subtraction method is the background image description model, which is the basis of background subtraction method to segment moving foreground. The single Gaussian background model is suitable for the single mode background. The color value distribution of each pixel is represented by a single Gaussian distribution $\eta(X_b, \mu_b, E_t)$, in which the subscript t represents the time. μ_b represents the mean value of the Gaussian distribution at time t ; E_t is the covariance of the Gaussian distribution. Let the color value of the current pixel be I_t , $d_t = I_t - \mu_t$. If the value of $d_t^T E_t^{-1} d_t$ is greater

than a certain threshold, this point is judged as the moving foreground point; otherwise, this point is considered as the background point matching with the Gaussian distribution.

The update of the single Gaussian distribution background model refers to the update of the parameters of the Gaussian function describing the background of the scene, and introducing a learning rate α to denote the update rate of the parameters, the parameters of the Gaussian distribution of the pixel points are updated according to the following formula¹⁰.

$$\begin{aligned} \mu_{t+1} &= (1-\alpha)\mu_t + \alpha I_t \\ E_{t+1} &= (1-\alpha)E_t + \alpha d_t d_t^T \end{aligned} \quad (2)$$

In formula (2), μ_t is the gray value of the pixel in the current background image, and it is also the mean value of the Gaussian distribution. I_t is the grayscale value of the current frame pixel; μ_{t+1} is the grayscale value of the background image after parameter update. Single gaussian background model can deal with simple scene is small and slowly changes, when a complex scene background changed or mutate, or background pixel values for peak distribution (such as small repeat movement), the change of the background pixels quickly, not by a relatively stable unimodal distribution gradually transition to another unimodal distribution, at this time should be used mixed gaussian background model. The basic idea of the mixed Gaussian model is as follows: for each pixel, K states are defined to represent the color presented, and K value is generally between 3 and 5. The larger the K value, the stronger the ability to deal with fluctuations, and the longer the corresponding processing time. Each of the K states is represented by a height and magnitude function. Some of these states represent the pixel value of the background, and the rest represent the pixel value of the moving foreground. If the color value of each pixel is expressed by the variable X, its probability density function can be expressed by K three-dimensional Gaussian functions as follows¹¹:

$$f(X_t = x) = \sum_{i=1}^K \omega_{i,t} \eta(x, \mu_{i,t}, E_{i,t}) \quad (3)$$

In formula (3), $\eta(x, \mu_{i,t}, E_{i,t})$ is the i Gaussian distribution at time t ; $\omega_{i,t}$ is the weight of the i Gaussian distribution at time t ,

and the sum of all the weights of the Gaussian distribution is 1. After obtaining the motion sequences of the dancers without background by using the mixed Gaussian background model, the key motion and pose sequences of the dance were optimized.

Optimization of Attitude Sequence

The obtained key movement sequence without background dance contains the internal details and movement information of the key movements of the dance, as well as the body movement information of the dancer when making the corresponding movements, while other details in the movement sequence will affect the subsequent contour capture. Therefore, by removing some useless details of the edge contour, the sequence of key movements and postures of the dance can be optimized. The Fourier descriptor can be used to represent a two-dimensional closed shape regardless of its starting point, position, etc.

If $x[m]$ and $y[m]$ are pixel coordinates of key movements of the dance, then all these points can be expressed as: $z[m] = x[m] + jy[m]$, and Fourier descriptor can be defined as ¹²:

$$Z[k] = DFT[z[m]] = \frac{1}{N} \sum_{m=0}^{N-1} z[m] e^{-j\frac{2\pi mk}{N}} \quad (4)$$

In the above formula, the range of k should be from 0 to $n-1$, where N is all the contour points on the closed shape. In the Fourier transform, the spectrum contains positive and negative frequencies in the middle of the DC component, so you have to include them in the inverse Fourier

transform as well. After the Fourier descriptor is transformed, the key dance sequences are optimized according to the inverse operation of Fourier descriptor. After analyzing the key movement sequences of the dance, the object posture changes in the sequence are tracked.

Dance Action Sequence Object Posture Change Tracking

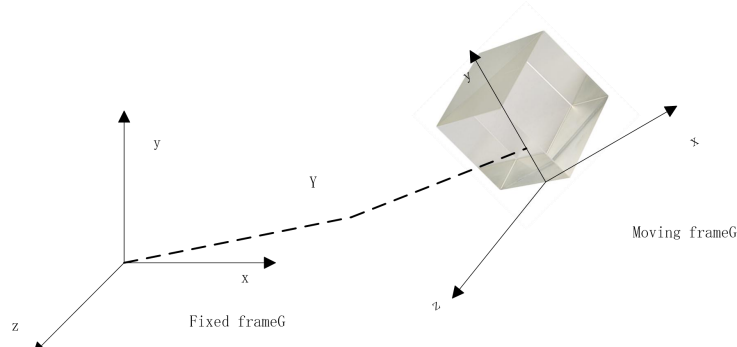
Description of human body rigidity

Motions that describe the key movements of a dancer's dance can be divided into two types: rigid movements and non-rigid movements. The rigid motion of the model is that the distances between vertices remain the same at all times. In other words, a rigidly moving object does not deform. Non-rigid movement, on the contrary, when a dancer makes a dance movement, the body can be viewed as a three-dimensional model, the model can move in three-dimensional space and can deform, the surface of the skin is elastic and can twist. It is impossible for 3D mannequin model to produce skin deformation like real people, so the motion of the model is considered as rigid motion in pose estimation.

The movement of the human body is reflected by the movement of each part of the human body in each frame with the position of the previous frame. In three dimensional space, the position of the human body corresponding to each frame is always moving, which is called "moving frame", denoted by Y ; The other reference frame is called a "fixed frame" and is denoted as G . The rigid motion of Y relative to G depicts the relative position of Y as Figure 1.

Figure 1.

Changes of moving frames during rigid movement of human body (represented by reference to fixed frames)



The coordinate system of all image frames in this paper is right-handed, and each frame is represented by x , y , and z and satisfies constraint $z = x \times y$. The rigid motion Y can be decomposed into a rotational motion R and a translational motion T . In Euclidean space, both rotation and translation can be described by a matrix. The vertex on the 3D mannequin is multiplied by this matrix to get the new coordinates for that vertex. A rotation in three dimensions can be a rotation about the x , Y , or z axes, each of which can be represented by its own matrix. The transformation of the model vertex requires the calculation of its new position in homogeneous coordinates. The position of the vertex is still in homogeneous coordinates¹³⁻¹⁵. If the vertex to nonhomogeneous coordinates, need to remove the last value of vertex coordinates, each vertex three coordinates, as a result of the model vertex motion is a linear transformation, keep the fourth element value is always 1 in this case, only the last element in the vertex removal namely it under nonhomogeneous coordinates.

Tracking of moving objects in dance action sequence

In the course of movement, not only the position of the moving object changes constantly, but the shape change of the three-dimensional human body model can be ignored basically. In the extraction result of the moving target in the previous frame, we cannot directly obtain the position of the moving target in the current frame, which leads to the impossibility of matching the moving target with a fixed template. However, in a very short time, the displacement generated by a moving object is always limited, and it is impossible to have a large displacement distance. A maximum distance is defined in advance, which specifies the upper limit of the motion displacement of all moving pixel points in the dance motion sequence in the image plane. According to the above analysis, in the dance action sequence, the movement of the dancer can be regarded as rigid movement. Therefore, the energy term of the following equation is defined¹⁶:

$$F_{fg}(I_l^t) = \min_{\|y-l\| \leq \delta} \sum_{\|x-y\| \leq \varepsilon} f(I_l^t, I_x^{t-1}) + \tau \quad (5)$$

In formula (5), I_l^t is the observed value of

the current pixel l of the input image of the current frame; I_x^{t-1} is the observed value of pixel x of the input image in the previous frame; $\tau > 0$ is a fixed constant.

When calculating the energy of any pixel point marked as fg in the current frame, a small neighborhood rigid template is used to match within the range of positions that may appear in the previous frame around the current point¹⁷. The closest color difference value is used to mark the pixel as the foreground energy. In other words, if a local range that is very similar to the foreground color distribution of the previous frame can be found around the pixel, its energy is very small, that is, the probability that the pixel belongs to the foreground is relatively large and vice versa. If the foreground of the previous frame is not found around the pixel, then we still have to assume that the energy of the foreground is a constant value.

The extracted moving area often contains the sun mark and its cast shadow, which will bring great difficulty for the follow-up target classification if not removed. In order to avoid dividing the shadow into moving objects or parts of moving objects, shadow removal should be performed before target classification. A shadow cancellation term is introduced in the energy function and integrated into the previous energy function. When one component or several components of the three channels of the current pixel are greatly different from the background, the probability that the pixel belongs to the shadow is relatively small. When the difference between the three channels is small, the current pixel is likely to be the background in the shadow. Introduce the shadow elimination item as follows:

$$S(I, l) = \begin{cases} S_{bg}(I, l) = 0, l = I \\ S_{fg}(I, l) = \rho(1 - \Delta H)(1 - \Delta S)(1 - \Delta V), l \neq I \end{cases} \quad (6)$$

In formula (6), ρ is the parameter of the shadow elimination item; ΔH , ΔS and ΔV are the color components in HSV color space respectively. The effect of the shadow removal item can be summarized as follows: the brightness and saturation are slightly lower than the background model, and areas with roughly the same chroma are more likely to be judged as the background under the shadow. According to the

above content, after determining the movement area of the dancer in the dance key movement sequence, the outline of the dance key movement can be captured by processing this area.

Realize the Key Action Outline Capture of Dance

Realization of contour capture of dance key movements, Generally speaking, a human model for motion capture usually contains three parts: points, lines, and bodies (surfaces). The points represent the joints of the human body, the lines connect two adjacent joints to represent the bones of the human body, and the secondary surface or mesh model masks the geometric model of muscles or skin on the bones. The points and lines form the human skeleton, which is essential to the human body model and defines the geometric properties and motion attributes of the human body model. Usually the human skeleton is a tree joint chain structure, with the human sacrum as the root node, the terminal joint as the leaf node, and other joints as the child nodes. The human pose is represented by the translation and rotation parameters of the root node and the rotation parameters of the other joints¹⁸. A typical human body model contains 15 to 22 joints. The more the number of joints in the human skeleton, the more refined the human action is represented, but the number of dimensions of the independent variables to be computed increases accordingly, making it more difficult to capture the action contours. According to the principle of human kinematics, the degrees of freedom of human joints are not fixed in 3 dimensions, like elbow and knee joints can be simplified to 1 degree of freedom. In order to facilitate the calculation and reduce the search space in the human posture space, the root node is defined as 6 degrees of freedom (3 rotation, 3 translation), the shoulder and hip joints are three degrees of freedom, and the remaining joints are simplified to 1 or 2 degrees of freedom.

According to the Laban movement spectrum principle describes the changes in the contour of the 3D mannequin of the dancer when making the corresponding key dance movements. The Laban movement spectrum consists of two parts:

directional symbols and spectral surface¹⁹. The movement symbols mainly contain nine directional symbols and three horizontal symbols, which can be combined into 27 basic directional symbols using directional and horizontal symbols. Laban believes that human action is mainly composed of five main components: body, space, force effect, shape and relationship.

The human body is a complex joint chain structure, and the motion of the current limb is jointly determined by all the parent joints in its joint chain, especially the end of the human body, which usually right 3 to 4 joints to determine its motion. Therefore, in order to obtain the motion velocity of each joint of the human body, the Euler angles of the 3D human model are converted to the position coordinates of the joints in the world coordinate system. The specific calculation is that, for a node, let (b, p, h) be the rotation angle (radian value) of the node around the z, x, and y axes, respectively, and (x_0, y_0, z_0) be the position of its sub-joint point relative to that joint, the rotation matrix is as follows:

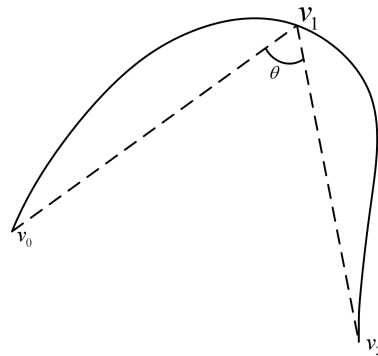
$$B = \begin{bmatrix} \cos b & \sin b & 0 \\ -\sin b & \cos b & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

$$P = \begin{bmatrix} 1 & 0 & 0 \\ 0 & \cos p & \sin p \\ 0 & -\sin p & \cos p \end{bmatrix} \quad (7)$$

$$P = \begin{bmatrix} \cos h & 0 & -\sin h \\ 0 & 1 & 0 \\ \sin h & 0 & \cos h \end{bmatrix}$$

The human skeleton is a joint chain structure. Except the root node, all the other nodes are described as rotating motion around their parent node. After obtaining the position of the joint, the end of the joint chain, i.e. the velocity of the limb with large motion range such as hand and foot, was selected as the basis for segmentation. So let's just think about the motion of one joint. Set the motion positions v_0 , v_1 and v_2 of the front and rear frames, and the Angle relationship between the motion direction of the current frame is as Figure 2:

Figure 2.
Angle of motion direction



Then the formula for calculating the Angle between the motion direction of the frame is as follows:

$$\theta = \arccos \frac{(v_0 - v_1)(v_2 - v_1)}{|v_0 - v_1||v_2 - v_1|} \quad (8)$$

Motion velocity is obtained by differencing between frames. Motion capture data acquisition frame rate is usually greater than 60fps, some devices even up to 1000fps. high frame rate data acquisition adjacent frames difference is small, increasing the noise of using speed judgment. Using uniform sampling, the frame rates are all reduced to 10fps, and then the motion direction angle and motion velocity of adjacent frames are calculated. The angle threshold and velocity threshold are set^{20,21}. When it is less than the threshold value, it is used as the splitting point of the action, and the corresponding current frame is the key frame.

The motion capture data is recorded according to the hierarchical structure of the human skeleton. Usually the motion of the root node describes the motion of the center of gravity of the human body, which describes the motion of the human body as a whole, while the other joints are the joint chain motion starting from the root node. Therefore, we use a hierarchical clustering approach to first divide the motion of the root node into clusters, and then cluster them in order according to the hierarchy of the skeleton. Since the direction symbol of Laban motion spectrum is the direction under the root node coordinate system, we calculate the position offset of all joints in the root node coordinate system relative to the parent node in this paper. After clustering the key frame data, each limb forms 27 classes, and each class is regarded as a state of the current limb, thus forming a hierarchical state node

according to the hierarchical structure of the human skeleton, which represents the starting and ending points of limb movements. When capturing the key dance movement contours, calculating the changes in the starting and ending points of limb movements in consecutive frames requires matching the corresponding points from frame to frame. Each histogram is the shape context of contour C1. Since the contour points p_i and q_i are in two consecutive key frames C1 and C2, respectively, the cost function to match these two points can be expressed as

$$C_{ij} = C(p_i, q_i) = \frac{1}{2} \sum_{k=1}^K \frac{[h_i(k) - h_j(k)]^2}{h_i(k) + h_j(k)} \quad (9)$$

In formula (9), $h_i(k)$, $h_j(k)$ represent the normalized histogram of contour points p_i and q_i . According to the principle of thin plate spline, all contour points are connected to get the outline of key dance movements in continuous frames. Through the above steps, the key movement outline of the dance was captured. So far, the research on the contour capture algorithm of key movements of dance with computer aid is completed.

RESEARCH ON ALGORITHM PERFORMANCE TEST

In order to reduce dancers' physical injuries caused by non-standard movements in the daily training process, and to better help them master the relevant dance key movements. In the above paper, the computer-aided dance key movement contour capture algorithm is studied to address the problems of the traditional movement contour capture algorithm. In this section, the performance of the algorithm will be tested by conducting relevant experimental studies.

Experiment Content

To verify the effectiveness of the algorithm, a dance video collected from the Internet is used for testing. The resolution of the dance video image is 1080*640, in order to avoid too many observed objects in the video, which affects the algorithm processing efficiency and the experimental process. The dance videos selected for this test are all single dances. The experiments were conducted in the form of comparison and validation, and the action contour capture algorithms mentioned in the reference ³ and reference ⁴ were selected as comparison group 1 and comparison group 2, and the key action contour capture algorithm for dance under computer-aided research in this paper was used as the experimental group. The comparison indexes selected for the experiments

were the capture error of the contour capture algorithm and the processing response time of the algorithm. By analyzing the experimental data, the performance advantages and disadvantages of the three contour capture algorithms are evaluated.

Experimental Results and Analysis

The dance videos used in the experiments are all single performances from various dance events at home and abroad, and all videos are free from overly elaborate backgrounds. Three sets of contour capture algorithms were used to capture the key movements of the excerpted frames in different videos. The test data of the algorithms corresponding to the capture error and processing response time are shown in Table 1.

Table 1.
Algorithm capture error and response time

Dance number	Comparison group 1		Comparison group 2		Experimental group	
	Capture error /mm	Response time /ms	Capture error /mm	Response time /ms	Capture error /mm	Response time /ms
1	70.4	293.2	46.2	110.8	12.8	75.7
2	64.1	293.2	74.8	101.8	12.7	83.4
3	65.1	291.8	49.7	103.3	12.6	73.9
4	64.5	292.6	43.3	107.7	13.5	72.3
5	98.1	292.8	46.2	109.5	12.2	59.4
6	66.7	289.7	50.7	112.8	12.1	76.8
7	86.8	292.5	42.6	102.6	13.0	65.5
8	61.4	289.7	44.7	106.9	13.7	59.2
9	66.6	291.7	69.5	101.1	13.4	75.5
10	69.1	291.9	42.4	108.9	13.6	81.6

Analysis of the Table 1 shows that the contour capture error of the experimental group algorithm is much smaller than the contour capture error of the two comparison groups' algorithms for different single dance performance videos with key movements contour capture. There are significant fluctuations in the contour capturing errors of the two comparison groups' algorithms for different dance videos. It indicates that the comparison group algorithms are less reliable when performing contour capture compared to the experimental group algorithms. The response time of the experimental group algorithm is significantly shorter than the other two algorithms when performing dance key action capture. Seeking the average value of the index data in the above table, it can be seen from the average value that the experimental group algorithm improves the capture accuracy by at

least 71.66% and shortens the response time by at least about 32.11%. The above two results show that the experimental group algorithm is better than the two algorithms mentioned in the reference for capturing the contours of dance movements. In conclusion, the computer-aided dance key movement contour capture algorithm studied in this paper has better reliability and accuracy and can be practically applied to dance movement contour capture.

CONCLUSION

With the rapid development of information technology, people can use images or videos as information carriers to study the change information of dance movement contours, which can help identify the subtle contour changes in dance movements, thus helping dancers to reduce injuries beyond the training process. Regarding

the problems of traditional contour capture algorithms, this paper investigates the computer-aided contour capture algorithm for key dance movements. The reliability and effectiveness of the contour capture algorithm are verified by using a single person dance video as the research object. Since the single dance video is used as the research object for all the experiments, in the future research, we should try to increase the capture objects in the video image for corresponding in-depth research to improve the performance of the algorithm.

REFERENCE

1. Struyvenberg M R , Groof A J D , Putten J V D , et al. A computer-assisted algorithm for narrow-band-imaging-based tissue characterization in Barrett's esophagus[J]. *Gastrointestinal Endoscopy*, 2020,93(1):89-98.
2. Jean-Philippe Rivière, Sarah Fdili Alaoui, Baptiste Caramiaux, et al. Mackay. Capturing Movement Decomposition to Support Learning and Teaching in Contemporary Dance[J]. *Proceedings of the ACM on Human-Computer Interaction*,2019,3(CSCW): 1-22.
3. Khonina S N , Porfirev A P . Generation of multi-contour plane curves using vortex beams[J]. *Optik - International Journal for Light and Electron Optics*, 2021, 229(9):166299.
4. Nguyen N Q , Vo D M , Lee S W . Contour-aware Polyp Segmentation in Colonoscopy Images using Detailed Upsampling Encoder-Decoder Networks[J]. *IEEE Access*, 2020, PP(99):1-1.
5. Wen Bo Z , Xin P , Yao Y , et al. Expert Consensus for the Treatment Algorithm for Navigation-assisted Reconstruction of Maxillofacial Deformities[J]. *The Chinese journal of dental research : the official journal of the Scientific Section of the Chinese Stomatological Association (CSA)*, 2020, 23(1):33-42.
6. Weichel F , Eisenmann U , Richter S , et al. A computer-assisted optimization approach for orthognathic surgery planning[J]. *Current Directions in Biomedical Engineering*, 2019, 5(1):41-44.
7. Zheng J , Lin D , Gao Z , et al. Deep Learning Assisted Efficient AdaBoost Algorithm for Breast Cancer Detection and Early Diagnosis[J]. *IEEE Access*, 2020, 8:96946-96954.
8. Bogach N , Boitsova E , Chernonog S , et al. Speech Processing for Language Learning: A Practical Approach to Computer-Assisted Pronunciation Teaching[J]. *Electronics*, 2021, 10(3):235.
9. Thomas N , Blanc V . Break it then Build Again: An Arts Based Duoethnographic Pilot Reconstructing Music Therapy and Dance-Movement Therapy Histories[J]. *The Arts in Psychotherapy*, 2021(2):101765.
10. Heather Spooner, Jenny B. Lee, Diane G. Langston, et al. Using distance technology to deliver the creative arts therapies to veterans: Case studies in art, dance/movement and music therapy[J]. *The Arts in Psychotherapy*,2019,62: 12-18.
11. Dieterich-Hartwell R , Goodill S , Koch S . Dance/Movement Therapy with Resettled Refugees: A Guideline and Framework Based on Empirical Data[J]. *The Arts in Psychotherapy*, 2020, 69:101664.
12. Crooks A , Mensinga J . Body, Relationship, Space: Dance Movement Therapy as an Intervention in Embodied Social Work With Parents and Their Children[J]. *Australian Social Work*, 2021(4):1-9.
13. Pylvninen P , Hyvnen K , Muotka J . The Profiles of Body Image Associate With Changes in Depression Among Participants in Dance Movement Therapy Group[J]. *Frontiers in Psychology*, 2020, 11:564788.
14. Bukhari Sarah A, Proussaefs Periklis, AlHelal Abdulaziz, et al. Use of Implant-Supported Custom Milled Impression Copings to Capture Soft-Tissue Contours and Incisal Guidance[J]. *Journal of prosthodontics : official journal of the American College of Prosthodontists*,2019,28(4): 473-479.
15. Kruisbrink, P. Paleo Cageao, H.P. Morvan, et al. Operating under jet splashing conditions can increase the capture efficiency of scoops[J]. *International Journal of Heat and Fluid Flow*,2019,76: 296-308.
16. Stumpel Lambert J, Wadhwani Chandur. Development and capture of soft tissue contours at time of implant placement[J]. *The Journal of prosthetic dentistry*,2017,117(6): 709-713.
17. Farha S A , Binder T R , Bronte C R , et al. Evidence of spawning by lake trout *Salvelinus namaycush* on substrates at the base of large boulders in northern Lake Huron[J]. *Journal of Great Lakes Research*, 2020, 46(6):1674-1688.
18. Liu C , Zhao R , Xie W , et al. Pathological lung segmentation based on random forest combined with deep model and multi-scale superpixels[J]. *Neural Processing Letters*, 2020, 52(2):1631-1649.
19. Dang T , Pham V , Nguyen H N , et al. AgasedViz: visualizing groundwater availability of Ogallala Aquifer, USA[J]. *Environmental Earth Sciences*, 2020, 79(5):1631-1649.
20. Persello C , Tolpekin V , Bergado J R , et al. Delineation of agricultural fields in smallholder farms from satellite images using fully convolutional networks and combinatorial grouping[J]. *Remote Sensing of Environment*, 2019, 231:111253.
21. A R M , A T Y , A M S , et al. A random forest algorithm-based approach to capture latent decision variables and their cutoff values[J]. *Journal of Biomedical Informatics*, 2020, 110:103548.