

Pose Forecasting in Industrial Human-Robot Collaboration

Alessio Sampieri¹, Guido Maria D’Amely di Melendugno¹, Andrea Avogaro²,
Federico Cunico², Francesco Setti², Geri Skenderi², Marco Cristani², and Fabio
Galasso¹

¹ Sapienza University of Rome
{sampieri, damely, galasso}@di.uniroma1.it

² University of Verona
{andrea.avogaro, federico.cunico, francesco.setti, geri.skenderi,
marco.cristani}@univr.it

Abstract. Pushing back the frontiers of collaborative robots in industrial environments, we propose a new Separable-Sparse Graph Convolutional Network (SeS-GCN) for pose forecasting. For the first time, SeS-GCN bottlenecks the interaction of the spatial, temporal and channel-wise dimensions in GCNs, and it learns sparse adjacency matrices by a teacher-student framework. Compared to the state-of-the-art, it only uses 1.72% of the parameters and it is ~ 4 times faster, while still performing comparably in forecasting accuracy on Human3.6M at 1 second in the future, which enables cobots to be aware of human operators. As a second contribution, we present a new benchmark of Cobots and Humans in Industrial Collaboration (CHICO). CHICO includes multi-view videos, 3D poses and trajectories of 20 human operators and cobots, engaging in 7 realistic industrial actions. Additionally, it reports 226 genuine collisions, taking place during the human-cobot interaction. We test SeS-GCN on CHICO for two important perception tasks in robotics: human pose forecasting, where it reaches an average error of 85.3 mm (MPJPE) at 1 sec in the future with a run time of 2.3 msec, and collision detection, by comparing the forecasted human motion with the known cobot motion, obtaining an F1-score of 0.64.

Keywords: Human Pose Forecasting, Graph Convolutional Networks, Human-Robot Collaboration in Industry.

1 Introduction

Collaborative robots (cobots) and modern Human Robot Collaboration (HRC) depart from the traditional separation of functions of industrial robots [36], because of the shared workspace [32]. Additionally cobots and humans perform actions concurrently and they will therefore physically engage in contact. While there is a clear advantage in increased productivity [62] (improved by as much as 85% [63]) due to the minimization of idle times, there are challenges in the

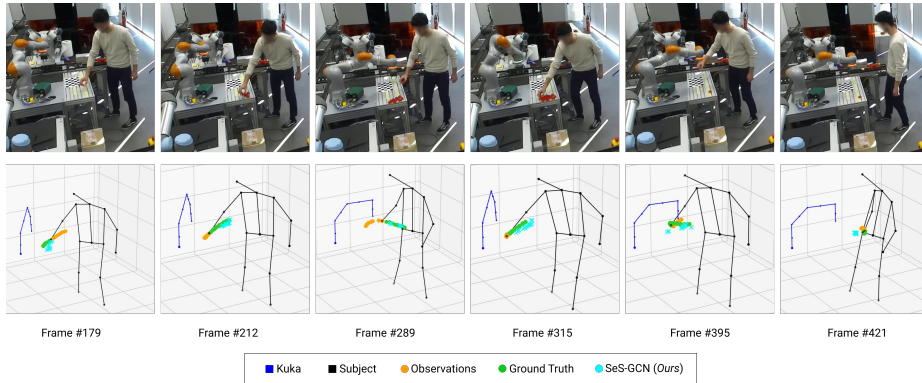


Fig. 1. A collision example from our CHICO dataset. On the top row some frames of the *Lightweight pick and place* action captured by one of the three cameras. On the bottom row, operator + robot skeletons. The forecasting takes an observation sequence (in yellow, here pictured for the right wrist only), and performs a prediction (cyan) which is compared with the ground truth (green). On frame 395 it is easy to see the robot hitting the operator, which is retracting, as it is evident in frame 421. See how the predictions by SeS-GCN follow closely the GT, except during the collision. At collision time, due to the impact, the abrupt change of the arm motion produces uncertain predictions, as it shows from the very irregular predicted trajectory.

workplace safety [24]: it is not about whether there will be contact, but rather about understanding its consequences [53].

The pioneering work of Shah et al. [63] has already shown that, in order to seamlessly and efficiently interact with human co-workers, cobots need to abide by two collaborative principles: (1) Making decisions on-the-fly, and (2) Considering the consequences of their decision on their teammates. The first calls for promptly and accurately detecting the human motion in the workspace. The second principle implies that cobots need to anticipate pose trajectories of their human co-workers and predict future collisions.

Motivated by these problems, the first contribution of our work is a novel Separable-Sparse Graph Convolutional Neural Network (SeS-GCN) for human pose forecasting. Pose forecasting requires understanding of the complex spatio-temporal joint dynamics of the human body and recent trends have highlighted the promises of modelling body kinematics within a single GCN framework [15, 17, 42, 44, 50, 69, 74]. We have designed SeS-GCN with performance and efficiency in mind, by bringing together, for the first time, three main modelling principles: depthwise-separable graph convolutions [37], space-time separable graph adjacency matrices [65], and sparse graph adjacency matrices [64]. In SeS-GCN, *separable* stands for limiting the interplay of joints with others (space), at different frames (time) and per channel (depth-wise). Within the GCN, different channels, frames and joints still interact by means of multi-hop messages. For the first time, sparsity is achieved by a teacher-student framework. The reduced interaction and sparsity results in comparable or less parameters than all GCN-based baselines [37, 65, 64], while improving performance by at least 2.7%.

Compared to the state-of-the-art (SoA) [49], SeS-GCN is lightweight, only using 1.72% of its parameters, it is ~ 4 times faster, while remaining competitive with just 1.5% larger error on Human 3.6M [31] when predicting 1 sec in the future. The model is described in detail in Sec. 3, experiments and ablation studies are illustrated in Sec. 5.

We also introduce the very first benchmark of Cobots and Humans in Industrial COllaboration (CHICO, an excerpt in Fig. 1). CHICO includes multi-view videos, 3D poses and trajectories of the joints of 20 human operators, in close collaboration with a robotic arm *KUKA LBR iiwa* within a shared workspace. The dataset features 7 realistic industrial actions, taken at a real industrial assembly line with a marker-less setup. The goal of CHICO is to endow cobots with perceptive awareness to enable human-cobot collaboration with contact. Towards this frontier, CHICO proposes to benchmark two key tasks: human pose forecasting and collision detection. Cobots currently detect collisions by mechanical-only events (transmission of contact wrenches, control torques, sensitive skins). This ensures safety but it harms the human-cobot interaction, because collisions break the motion of both, which reduces productivity, and may be annoying to the human operator. CHICO features 240 1-minute video recordings, from which two separate sets of test sequences are selected: one for estimating the accuracy in pose forecasting, so cobots may be aware of the next future (1.0 sec); and one with 226 genuine collisions, so cobots may foresee them and possibly re-plan. The dataset is detailed in Sec. 4, experiments are illustrated in Sec. 6.

When tested on CHICO, the proposed SeS-GCN outperforms all baselines and reaches an error of 85.3 mm (MPJPE) at 1.00 sec, with a negligible run time of 2.3 msec (as reported in Table 5). Additionally, the forecast human motion is used to detect human-cobot collision, by checking whether the predicted trajectory of the human body intersects that of the cobot. This is also encouraging, as SeS-GCN allows to reach an F1-score of 0.64. Both aspects contribute to a cobot awareness of the future, which is instrumental for HRC in industrial applications.

2 Related Work

Human pose forecasting. Human pose forecasting is a recent field which has some intersection with human action anticipation in computer vision [42] and HRC [18]. Previous studies exploited Temporal Convolutional Networks (TCNs) [2,22,41,59] and Recurrent Neural Networks (RNNs) [20,23,33]. Both architectures are naturally suited to model the temporal dimension. Recent works have expanded the range of available methods by using Variational Auto-Encoders [8], specific and model-agnostic layers that implicitly model the spatial structure of the human skeleton [1], or Transformer Networks [9].

Pose forecasting using Graph Convolutional Networks (GCN). Most recent research uses GCNs [17,44,49,65,74]. In [49], the authors have mixed GCN for modelling the joint-joint interaction with Transformer Networks for the temporal patterns. Others [44,65,74] have adopted GCNs to model the space-time

Table 1. Comparison between the state-of-the-art datasets and the proposed CHICO; *unk* stands for “unknown”.

	Quantitative Details							Rec. Scene	Actions Type		Tasks			Markerless
	# Classes	# Subj.	Avg Rec. Time	# Joints	FPS	Aspect Ratio	# Sensors		Industr.	HRC	Action Recog.	Pose Forec.	Coll. Det.	
Human3.6M [31]	15	11	100.49 s	32	25	normalized	15	mo-cap studio				✓		
AMASS [48]	11265	344	12.89 s	variable	variable	original	variable	mo-cap studio				✓		
3DPW [51]	47	7	28.33 s	18	60	original	18	outdoor locations				✓		
ExPI [25]	16	4	<i>unk</i>	18	25	original	88	mo-cap studio				✓		
CHI3D [19]	8	6	<i>unk</i>	<i>unk</i>	<i>unk</i>	original	14	mo-cap studio				✓		
InHARD [16]	14	16	< 8 s	17	120	original	20	assembly line	✓	✓	✓			
CHICO (ours)	7	20	55 s	15	25	original	3	assembly line	✓	✓		✓	✓	✓

body kinematics, devising, in the case of [17], hierarchical architectures to model coarse-to-fine body dynamics.

We identify three main research directions for improving efficiency in GCNs: **i.** space-time separable GCNs [65], which factorizes the spatial joint-joint and temporal patterns of the adjacency matrix; **ii.** depth-wise separable graph convolutions [30], which has been explored by [3] in the spectral domain; and **iii.** sparse GCNs [64], which iteratively prunes the terms of the adjacency matrix of a GCN. Notably, all three techniques also yield better performance than the plain GCN. Here, for the first, we bring together these three aspects into an end-to-end space-time-depthwise-separable and sparse GCN. The three techniques are complementary to improve both efficiency and performance, but their integration requires some structural changes (*e.g.*, adopting teacher-student architectures for sparsifying), as we describe in Sec. 3.

Human Robot Collaboration (HRC).

HRC is the study of collaborative processes where human and robot agents work together to achieve shared goals [4,11]. Computer vision studies on HRC are mostly related to pose estimation [10,21,40] to locate the articulated human body in the scene.

In [12,34,57], methodologies for robot motion planning and collision avoidance are proposed; their study perspective is opposite to ours, since we focus on the human operator. In this regard, the works of [5,14,35,45] model the operators’ whereabouts through detection algorithms which approximate human shapes using simple bounding boxes. Approaches that predict the human motion during collaborative tasks are in [66,73] using RNNs and in [68] using Gaussian processes. Other work [38] models the upper body and the human right hand (which they call the Human End Effector) by considering the robot-human handover phase. As motion prediction engine, DCT-RNN-GCN [49] is considered, against which we compare in the experiments.

Datasets for pose forecasting. Human pose forecasting datasets cover a wide spectrum of scenarios, see Table 1 for a comparative analysis. Human3.6M [31] considers everyday actions such as conversing, eating, greeting and smoking.

Data were acquired using a 3D marker-based motion capture system, composed of 10 high-speed infrared cameras. AMASS [48] is a collection of 15 datasets where daily actions were captured by an optical marker-based motion capture. Human3.6M and AMASS are standard benchmarks for human pose forecasting, with some overlap in the type of actions they deal with. The 3DPW dataset [51] focuses on outdoor actions, captured with a moving camera and 17 Inertial Measurement Units (IMU), embedded on a special suit for motion capturing [60]. The recent ExPI dataset [25] contains 16 different dance actions performed by professional dancers, for a total of 115 sequences, and it is aimed at motion prediction. ExPI has been acquired with 68 synchronised and calibrated color cameras and a motion capture system with 20 mocap cameras. Finally, the CHI3D dataset [19] reports 3D data taken from MOCAP systems to study human interactions.

None of these datasets answer our research needs, i.e., a benchmark taken by a sparing, energy-efficient markerless system, focused on the industrial HRC scenario, where forecasting may be really useful for anticipating collisions between the humans and robots. In fact, the only dataset relating to industrial applications is InHARD [16]. Therein, humans are asked to perform an assembly task while wearing inertial sensors on each limb. The dataset is designed for human action recognition, and it involves 16 individuals performing 13 different actions each, for a total of 4800 action samples over more than 2 million frames. Despite showcasing a collaborative robot, in this dataset the robot is mostly static, making it unsuitable for collision forecasting.

3 Methodology

We build an accurate, memory efficient and fast GCN by bridging three diverse research directions: **i.** Space-time separable adjacency matrices; **ii.** Depth-wise separable graph convolutions; **iii.** Sparse adjacency matrices. This results in an all-Separable and Sparse GCN encoder for the human body kinematics, which we dub SeS-GCN, from which the future frames are forecast by a Temporal Convolutional Network (TCN).

3.1 Background

Problem Formalization. Pose forecasting is formulated as observing the 3D coordinates $\mathbf{x}_{v,t}$ of V joints across T frames and predicting their location in the K future frames. For convenience of notation, we gather the coordinates from all joints at frame t into the matrix $X_t = [\mathbf{x}_{v,t}]_{v=1}^V \in \mathbb{R}^{3 \times V}$. Then we define the tensors $\mathcal{X}_{in} = [X_1, X_2, \dots, X_T]$ and $\mathcal{X}_{out} = [X_{T+1}, X_{T+2}, \dots, X_{T+K}]$ that contain all observed input and target frames, respectively.

We consider a graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ to encode the body kinematics, with all joints at all observed frames as the node set $\mathcal{V} = \{\mathbf{v}_{i,t}\}_{i=1, t=1}^{V,T}$, and edges $(\mathbf{v}_{i,t}, \mathbf{v}_{j,s}) \in \mathcal{E}$ that connect joints i, j at frames t, s .

Graph Convolutional Networks (GCN). A GCN is a layered architecture:

$$\mathcal{X}^{(l+1)} = \sigma \left(A^{(l)} \mathcal{X}^{(l)} W^{(l)} \right) \quad (1)$$

The input to a GCN layer l is the tensor $\mathcal{X}^{(l)} \in \mathbb{R}^{C^{(l)} \times V \times T}$ which maintains the correspondence to the V body joints and the T observed frames, but increases the depth of features to $C^{(l)}$ channels. $\mathcal{X}^{(1)} = \mathcal{X}_{in}$ is the input tensor at the first layer, with $C^{(1)} = 3$ channels given by the 3D coordinates. $A^{(l)} \in \mathbb{R}^{VT \times VT}$ is the adjacency matrix relating pairs of VT joints from all frames. Following most recent literature [17,49,64,65], $A^{(l)}$ is learnt. $W^{(l)} \in \mathbb{R}^{C^{(l)} \times 1 \times 1}$ are the learnable weights of the graph convolutions. σ is a the non-linear PReLU activation function.

3.2 Separable & Sparse Graph Convolutional Networks (SeS-GCN)

We build SeS-GCN by integrating the three mentioned modelling dimensions: **i.** separating spatial and temporal interaction terms in the adjacency matrix of a GCN; **ii.** separating the graph convolutions depth-wise; **iii.** sparsifying the adjacency matrices of the GCN.

Separating space-time. STS-GCN [65] has factored the adjacency matrix $A^{(l)}$ of the GCN, at each layer l , into the product of two terms $A_s^{(l)} \in \mathbb{R}^{V \times V \times T}$ and $A_t^{(l)} \in \mathbb{R}^{T \times T \times V}$, respectively responsible for the temporal-temporal and joint-joint relations. The GCN formulation becomes:

$$\mathcal{X}^{(l+1)} = \sigma \left(A_s^{(l)} A_t^{(l)} \mathcal{X}^{(l)} W^{(l)} \right) \quad (2)$$

Eq. (2) bottlenecks the interplay of joints across different frames, implicitly placing more emphasis on the interaction of joints on the same frame ($A_s^{(l)}$) and on the temporal pattern of each joint ($A_t^{(l)}$). This reduces the memory-footprint of a GCN by approx. 4x while improving its performance (cf. Sec. 5.1). Note that this differs from alternating spatial and temporal modules, as it is done in [71] and [7], respectively for trajectory forecasting and action recognition.

Separating depth-wise. Inspired by depth-wise convolutions [13,30], the approach in [37] has introduced depth-wise graph convolutions for image classification, followed by [3] which resorted to a spectral formulation of depth-wise graph convolutions for graph classification. Here we consider depth-wise graph convolutions for pose forecasting. The depth-wise formulation bottlenecks the interplay of space and time (operated by the adjacency matrix $A^{(l)}$) with the channels of the graph convolution $W^{(l)}$. The resulting all-separable model which we dub STS-DW-GCN is formulated as such:

$$\mathcal{H}^{(l)} = \gamma \left(A_s^{(l)} A_t^{(l)} \mathcal{X}^{(l)} W_{DW}^{(l)} \right) \quad (3a)$$

$$\mathcal{X}^{(l+1)} = \sigma \left(\mathcal{H}^{(l)} W_{MLP}^{(l)} \right) \quad (3b)$$

Adding the depth-wise graph convolution splits the GCN of layer l into two terms. The first, Eq. (3a), focuses on space-time interaction and limits the channel cross-talk by the use of $W_{DW}^{(l)} \in \mathbb{R}^{\frac{C^{(l)}}{\alpha} \times 1 \times 1}$, with $1 \leq \alpha \leq C^{(l)}$ setting

the number of convolutional groups ($\alpha = C^{(l)}$ is the plain single-group depth-wise convolution). The second, Eq. (3b), models the intra-channel communication just. This may be understood as a plain (MLP) 1D-convolution with $W_{\text{MLP}}^{(l)} \in \mathbb{R}^{C^{(l)} \times 1 \times 1}$ which re-maps features from $C^{(l)}$ to $C^{(l+1)}$. γ is the ReLU6 non-linear activation function. Overall, this does not significantly reduce the number of parameters, but it deepens the GCN without over-smoothing [58], which improves performance (see Sec. 5.1 for details).

Sparsifying the GCN. Sparsification has been used to improve the efficiency (memory and, in some cases, runtime) of neural networks since the seminal pruning work of [39]. [64] has sparsified GCNs for trajectory forecasting. This consists in learning masks $\mathcal{M}$ which selectively erase certain parameters in the adjacency matrix of the GCN. Here we integrate sparsification with the all-separable GCN design, which yields our proposed SeS-GCN for human pose forecasting:

$$\mathcal{H}^{(l)} = \gamma \left((\mathcal{M}_s^{(l)} \odot A_s^{(l)}) (\mathcal{M}_t^{(l)} \odot A_t^{(l)}) \mathcal{X}^{(l)} W_{DW}^{(l)} \right) \quad (4a)$$

$$\mathcal{X}^{(l+1)} = \sigma \left(\mathcal{H}^{(l)} W_{\text{MLP}}^{(l)} \right) \quad (4b)$$

$\odot$ is the element-wise product and $\mathcal{M}_{\{s,t\}}^{(l)}$ are binary masks. Both at training and inference, [64] generates masks, it uses those to zero certain coefficients of the adjacency matrix A , and it adopts the resulting GCN for trajectory forecasting. By contrast, we adopt a teacher-student framework during training. The teacher learns the masks, and the student only considers the spared coefficients in A . At inference, our proposed SeS-GCN only consists of the student, which simply adopts the learnt sparse A_s and A_t . Compared to [64], the approach of SeS-GCN is more robust at training, it yields fewer model parameters at inference ($\sim 30\%$ less for both A_s and A_t), and it reaches a better performance, as it is detailed in Sec. 5.1.

3.3 Decoder Forecasting

Given the space-time representation, as encoded by the SeS-GCN, the future frames are then decoded by using a temporal convolutional network (TCN) [2, 22, 41, 65]. The TCN remaps the temporal dimension to match the sought output number of predicted frames. This part of the model is not considered for improvement because it is already efficient and it performs satisfactorily.

4 The CHICO dataset

In this section the CHICO dataset is detailed by describing the acquisition scenario and devices, the cobot and the performed actions. We release RGB videos, skeletons and calibration parameters [†].

[†]Code and dataset are available at: <https://github.com/AlessioSam/CHICO-PoseForecasting>.

The scenario. We are in a smart-factory environment, with a single human operator standing in front of a $0.9\text{ m} \times 0.6\text{ m}$ workbench and a cobot at its end (see Fig.1). The human operator has some free space to turn towards some equipment and carry out certain assembly, loading and unloading actions [54]. In particular, light plastic pieces and heavy tiles, a hammer, abrasive sponges are available. The detailed setups for each action are reported graphically in the additional material. A total of 20 human operators have been hired for this study. They attended a course on how to operate with the cobot and signed an informed consent form prior to the recordings.

The collaborative robot. A 7 degrees-of-freedom Kuka LBR iiwa 14 R820 collaborates with the human operator during the data acquisition process. Weighing in at 29.5 kg and with the ability to handle a payload up to 14 kg, it is widely used in modern production lines. More details on the cobot can be found in the supplementary material.

The acquisition setup. The acquisition system is based on three RGB HD cameras providing three different viewpoints of the same workplace: two frontal-lateral, one rear view. The frame rate is 25Hz. The videos were first checked for erroneous or spurious frames, then we used Voxelpose [67] to extract 3D human pose for each frame. Extrinsic parameters of each camera are estimated w.r.t. the robot’s reference frame by means of a calibration chessboard of $1 \times 1\text{ m}$, and temporal alignment is guaranteed by synchronization of all the components with an Internet Time Server. In our environment, Voxelpose estimates a joint positioning accuracy in terms of Mean Per Joint Position Error (MPJPE) of 24.99mm using three cameras, which is enough for our purposes, as an ideal compromise between portability of the system and accuracy. We confirm these numbers in two ways: the first is by checking that human-cobot collisions were detected with 100% F1 score (we have a collision when the minimum distance between the human limbs and the robotic links is below a predefined threshold). Secondly, we show that the new CHICO dataset does not suffer from a trivial zero velocity solution [52], *i.e.* results achieved by a zero velocity model underperform the current SoA in equal proportion as for the large-scale established Human3.6M.

Actions. The 7 types of actions of CHICO are inspired from ordinary work sessions in an HRC disassembly line as described in the review work of [29]. Each action is repeated over a time interval of ~ 1 minute on average. Each action is associated to a goal that the human operator has to achieve by a given time limit, which requires them to move with a certain velocity. Each action consists of repeated interactions with the robot (*e.g.*, robot place, human picks) which, due to the limited space, lead to some *unconstrained collisions** which we label accordingly. Globally, from the 7 actions $\times$ 20 operators, we collect 226 different collisions. On the additional material, an excerpt of the videos with collisions are available. In the following, each action is shortly described.

*Unconstrained collisions is a term coming from [26], indicating a situation in which only the robot and human are directly involved into the collision.

- **Lightweight pick and place** (*Light PEP*). The human operator is required to move small objects of approximately 50 grams from a loading bay to a delivery location within a given time slot. The bay and the delivery location are at the opposite sides of the workbench. Meanwhile, the robot loads on of this bay so that the human operator has to pass close to the robotic arm. In many cases the distance between the limbs and the robotic arm is few centimeters.
- **Heavyweight pick and place** (*Heavy PEP*). The setup of this action is the same as before, but the objects to be moved are floor tiles weighing 0.75 kg. This means that the actions have to be carried out with two hands.
- **Surface polishing** (*Polishing*). This action was inspired by [47], where the human operator polishes the border of a 40 by 60cm tile with some abrasive sponge, and the robot mimics a visual quality inspection.
- **Precision pick and place** (*Prec. PEP*). The robot places four plastic pieces in the four corners of a 30×30cm table in the center of the workbench, and the human has to remove them and put on a bay, before the robot repeats the same unloading.
- **Random pick and place** (*Rnd. PEP*). Same as the previous action, except for the plastic pieces which were continuously placed by the robot randomly on the central 30×30cm table, and the human operator has to remove them.
- **High shelf lifting** (*High lift*). The goal was to pick light plastic pieces (50 grams each) on a sideways bay filled by the robot, putting them on a shelf located at 1.70m, at the opposite side of the workbench. Due to the geometry of the workspace, the arms of the human operator were required to pass above or below the moving robotic arm. In this way, close distances between the human arm and forearm and the robotic links were realized.
- **Hammering** (*Hammer*). The operator hits with a hammer a metallic tide held by the robot. In this case, the interest was to check how much the collision detection is robust to an action where the human arm is colliding close to the robotic arm (that is, on the metallic tile) without properly colliding *with the robotic arm*.

5 Experiments on Human3.6M

We benchmark the proposed SeS-GCN model on the large and established Human3.6M [31]. In Sec. 5.1, we analyze the design choices corresponding to the models discussed in Sec. 3, then we compare with the state-of-the-art in Sec. 5.2.

Human3.6M [31] is an established dataset for pose forecasting, consisting of 15 daily life actions (e.g. Walking, Eating, Sitting Down). From the original skeleton of 32 joints, 22 are sampled as the task, representing the body kinematics. A total of 3.6 million poses are captured at 25 fps. In line with the literature [49,52,17], subjects 1, 6, 7, 8, 9 are used for for training, subject 11 for validation, and subject 5 for testing.

Metric. The prediction error is quantified via the MPJPE error metric [31,50], which considers the displacement of the predicted 3D coordinates w.r.t. the

ground truth, in millimeters, at a specific future frame t :

$$L_{\text{MPJPE}} = \frac{1}{V} \sum_{v=1}^V \|\hat{\mathbf{x}}_{vt} - \mathbf{x}_{vt}\|_2. \quad (5)$$

5.1 Modelling choices of SeS-GCN

We review and quantify the impact of the modelling choices of SeS-GCN:

Efficient GCN baselines. In Table 2, we first validate the three difference modelling approaches to efficient GCNs, namely space-time separable STS-GCN [65], depth-wise separable graph convolutions DW-GCN [37], and Sparse-GCN [64]. STS-GCN yields the lowest MPJPE error of 117.0 mm at a 1 sec forecasting horizon (2.4% better than DW-GCN, 4.8% better than Sparse-GCN) with the fewest parameters, 57.6k (ca. x4 less). We build therefore on this approach.

Deeper GCNs. It is a long standing belief that Deep Neural Networks (DNN) owe their performance to depth [27,46,70,72]. However, deeper models require more parameters and have a longer processing time. Additionally, deeper GCNs may suffer from over-smoothing [58]. Seeking both better accuracy and efficiency, we consider three pathways for improvement: (1) add GCN layers; (2) add MLP layers between layers of GCNs; (3) adopt depth-wise graph convolutions, which also add MLP layers between GCN ones (cf. Sec. 3.2).

As shown in Table 2, there is a slight improvement in performance with 5 STS-GCN layers (MPJPE of 115.9 mm), but deeper models underperform. Adding MLP layers between the GCN ones (depth of 5+5) also decreases performance (MPJPE of 125.2). By contrast, adding depth by depth-wise separable graph convolutions (STS-DW-GCN of depth 5+5) reduces the error to 114.8 mm. This may be explained by the virtues of the increased depth in combination with the limiting cross-talk of joint-time-channels, which existing literature confirms [13,37,65]. We note that space-time and depth-wise channel separability

Table 2. MPJPE error (millimeters) for long-range predictions (25 frames) on Human3.6M [31] and numbers of parameters. Best figures overall are reported in bold, while underlined figures represent the best in each block. The proposed model has comparable or less parameters than the GCN-based baselines [30,64,65] and it outperforms the best of them [65] by 2.6%.

	Depth	MPJPE	Parameters (K)	DW-Separable	ST-Separable	Sparse	w/ MLP layers	Teacher-Student
GCN	4	123.2	222.7					
DW-GCN [37]	4+4	119.8	223.2	✓			✓	
STS-GCN [†] [65]	4	<u>117.0</u>	57.6		✓			
Sparse-GCN [64]	4	122.7	257.9			✓		
STS-GCN	5	115.9	<u>68.6</u>		✓			
STS-GCN	6	116.1	79.9		✓			
STS-GCN w/ MLP	5+5	125.2	101.4		✓		✓	
STS-DW-GCN	5+5	<u>114.8</u>	70.0	✓	✓		✓	
STS-DW-Sparse-GCN	5+5	115.7	122.4	✓	✓	✓	✓	
SeS-GCN (proposed)	5+5	<u>113.9</u>	<u>58.6</u>	✓	✓	✓	✓	✓

is complementary. Altogether, this performance is beyond the STS-GCN performance (114.8 Vs. 117.0 mm), at a slight increase of the parameter count (70k Vs. 57.6k).

Sparsifying GCNs and the proposed SeS-GCN. Finally, we target to improve efficiency by model compression. Trends have reduced the size of models by reducing the parameter precision [61], by pruning and sparsifying some of the parameters [56], or by constructing teacher-student frameworks, whereby a smaller student model is paired with a larger teacher to reach its same performance [28,43]. Note that the last technique is the current go-to choice in deploying very large networks such as Transformers [6].

We start off by compressing the model with sparse adjacency matrices by the approach of Sparse-GCN [64]. They iteratively optimize the learnt parameters and the masks to select some (the selection occurs by a network branch, also at inference, cf. 3.2). As illustrated in Table 2, the approach of [64] does not make a viable direction (STS-DW-Sparse-GCN), since the error increases to 115.7 mm and the parameter count to 122.4k.

Reminiscent of teacher-student models, in the proposed SeS-GCN we first train a teacher STS-DW-GCN, then use its learnt parameters to sparsify the affinity matrices of a student STS-DW-GCN, which is then trained from scratch. SeS-GCN achieves a competitive parameter count and the lowest MPJPE error of 113.9 mm, being comparable with the current SoA [49] and using only 1.72% of its parameters (58.6k Vs. 3.4M).

5.2 Comparison with the state-of-the-art (SoA)

In Table 5.2, we evaluate the proposed SeS-GCN against three most recent techniques, over a short time horizon (10 frames, 400 msec) and a long time horizon (25 frames, 1000 msec). The first, DCT-RNN-GCN [49], the current SoA, uses DCT encoding, motion attention and RNNs and, differently from other models, demands more frames as input (50 vs. 10). The other two, MSR-GCN [17] and STS-GCN [65] adopt GCN-only frameworks, the former adopts a multi-scale approach, the latter acts a separation between spatial and temporal encoding.

Both on Short- and long-term predictions, at the 400 and 1000 msec horizons, the proposed SeS-GCN outperforms other techniques [65,49] and it is within a 1.5% error w.r.t. the current SoA [49], while only using 1.72% parameters and being ~ 4 times faster than [49].

6 Experiments on CHICO

We benchmark on CHICO the SoA and the proposed SeS-GCN model. The two HRC tasks of human pose forecasting and collision detection are discussed in Secs. 6.1 and 6.2 respectively.

[†]Results for STS-GCN differ from [65], due to revision by the authors, cf. <https://github.com/FraLuca/STSGCN>.

Table 3. MPJPE error in mm for short-term (400 msec, 10 frames) and long-term (1000 msec, 25 frames) predictions of 3D joint positions on Human3.6M. The proposed model achieves competitive performance with the SoA [49], while adopting 1.72% of its parameters and running ~ 4 times faster, cf. Table 5. Results are discussed in Sec. 5.2.

Time Horizon (msec)	Walking		Eating		Smoking		Discussion		Directions		Greeting		Phoning		Posing	
	400	1000	400	1000	400	1000	400	1000	400	1000	400	1000	400	1000	400	1000
DCT-RNN-GCN [49]	39.8	58.1	36.2	75.5	36.4	69.5	65.4	119.8	56.5	106.5	78.1	138.8	49.2	105.0	75.8	178.2
MSR-GCN [17]	45.2	63.0	40.4	77.1	38.1	71.6	69.7	117.5	53.8	100.5	93.3	147.2	51.2	104.3	85.0	174.3
STS-GCN [†] [65]	51.0	70.2	43.3	82.6	42.3	76.1	71.9	118.9	63.2	109.6	86.4	136.1	53.8	108.3	84.7	178.4
SeS-GCN (proposed)	48.8	67.3	41.7	78.1	40.8	73.7	70.6	116.7	60.3	106.9	83.8	137.2	52.6	106.7	82.6	173.5

Time Horizon (msec)	Purchases		Sitting		Sitting Down		Taking Photo		Waiting		Walking Dog		Walking Together		Average	
	400	1000	400	1000	400	1000	400	1000	400	1000	400	1000	400	1000	400	1000
DCT-RNN-GCN [49]	73.9	134.2	56.0	115.9	72.0	143.6	51.5	115.9	54.9	108.2	86.3	146.9	41.9	64.9	58.3	112.1
MSR-GCN [17]	79.6	139.1	57.8	120.0	76.8	155.4	56.3	121.8	59.2	106.2	93.3	148.2	43.8	65.9	62.9	114.1
STS-GCN [†] [65]	83.1	141.0	60.8	121.4	79.4	148.4	59.4	126.3	62.0	113.6	97.3	151.5	49.1	72.5	65.8	117.0
SeS-GCN (proposed)	82.2	139.1	59.9	117.5	78.1	146.0	57.7	121.2	58.5	107.5	94.0	147.7	48.3	70.8	64.0	113.9

6.1 Pose forecasting benchmark

Here we describe the evaluation protocol proposed for CHICO and report comparative evaluation of pose forecasting techniques.

Evaluation protocol. We create the train/validation/test split by assigning 2 subjects to the validation (subjects 0 and 4), 4 to the test set (subjects 2, 3, 18 and 19), and the remaining 14 to the training set. For short-range prediction experiments, abiding the setup of Human3.6M [31], we consider 10 frames as observation time and 10 or 25 frames as forecasting horizon. Differently from all reported techniques, DCT-RNN-GCN requires 50 input frames.

We adopt the same Mean Per Joint Position Error (MPJPE)[31] as Human3.6M, in Eq. (5), which also defines the training loss for the evaluated techniques.

None of the motion sequences for pose forecasting contains collisions. In fact, the objective is to train and test the “correct” collaborative human behavior, and not the human retractions and the pauses due to the collisions^{**}.

Table 4. MPJPE error in mm for short-term (400 msec, 10 frames) and long-term (1000 msec, 25 frames) prediction of 3D joint positions on CHICO dataset. The average error is 7.9% lower than the other models in the short-term and 2.4% lower in the long-term prediction. See Sec. 6.1 for a discussion.

Time Horizon (msec)	Hammer		High Lift		Prec. P&P		Rnd. P&P		Polishing		Heavy P&P		Light P&P		Average	
	400	1000	400	1000	400	1000	400	1000	400	1000	400	1000	400	1000	400	1000
DCT-RNN-GCN [49]	41.1	39.0	69.4	128.8	50.6	83.3	52.7	88.2	42.1	76.0	64.1	121.5	62.1	104.2	54.6	91.6
MSR-GCN [17]	41.6	39.7	67.8	130.2	50.2	81.3	53.4	90.3	41.1	73.2	62.7	118.2	61.5	101.9	54.1	90.7
STS-GCN [65]	46.6	52.1	64.2	116.4	48.3	79.5	52.0	87.9	42.1	73.9	60.6	106.5	57.2	95.2	53.0	87.4
SeS-GCN (proposed)	40.9	49.3	62.1	116.3	46.0	77.4	48.4	84.8	38.8	72.4	56.1	104.4	56.2	92.2	48.8	85.3

Comparative evaluation. In Table 4, we compare pose forecasting techniques from the SoA and the proposed SeS-GCN. On the short-term predictions the

^{**}After the collisions, the robot stops for 1 seconds, during which the human operator usually stands still, waiting for the robot to resume operations.

best performance is that of SeS-GCN, reaching an MPJPE error of 48.8 mm, which is 7.9% better than the second best STS-GCN [65].

On the longer-term predictions, the best performance (MPJPE error of 85.3 mm) is also detained by SeS-GCN, which is 2.4% better than the second best STS-GCN [65]. The proposed model outperforms all techniques on all actions except *Hammer*, a briefly repeating action which may differ for single hits. We argue that DCT-RNN-GCN [49] may get an advantage from using 50 input frames (all other methods use 10 frames)

For a graphical illustration, Fig. 2 shows a distribution of the error per joint calculated over all the actions, for the horizons 400 (*left*) and 1000 msec (*right*). In both cases the error gets larger as we get closer to the extrema of the kinematic skeleton, since those joints move the most. The slightly larger error at the right hands (70.03 and 125.76 mm, respectively) matches that subjects are right-handed (but some actions are operated with both hands).

For a sanity check of results, we have also evaluated the performance of a trivial zero velocity model. [52] has found that keeping the last observed positions may be a surprisingly strong (trivial) baseline. For CHICO, the zero velocity model scores an MPJPE of 110.6 at 25-frames, worse than the 85.3 mm score of SeS-GCN. This is in line with the large-scale dataset Human3.6M [31], where the performance of the trivial model is 153.3 mm.

6.2 Collision detection experiments

Evaluation protocol. We consider a collision to occur when any body limb of the subject gets too close to any part of the cobot, i.e. within a distance threshold, for at least one frame. In particular, a collision refers to the proximity between the cobot and the human in the forecast portion of the trajectory. The (Euclidean) distance threshold is set to 13 cm.

The motion of the cobot is scripted beforehand, thus known. The motion of the human subjects in the next 1000 msec needs to be forecast, starting from

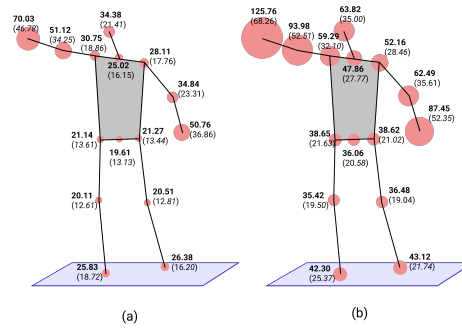


Fig. 2. Average MPJPE distribution for all actions in CHICO on different joints for (a) short-term (0.40 s) and (b) long-term (1.00 s) predictions. The radius of the blob gives the spatial error with the same scale of the skeleton.

the observation of 400 msec. The train/validation/test sets sample sequences of 10+25 frames with stride of 10.

Evaluation of collision detection. For the evaluation of collision, following [55], both the cobot arm parts and the human body limbs are approximated by cylinders. The diameters for the cobot are fixed to 8cm. Those of the body limbs are taken from a human atlas.

In Table 5, we report precision, recall and F_1 scores for the detection of collisions on the motion of 2 test subjects, which contains 21 collisions. The top performer in pose forecasting, our proposed SeS-GCN, also yields the largest F_1 score of 0.64. The lower performing MSR-GCN [17] yields poor collision detection capabilities, with an F_1 score of 0.31.

Table 5. Evaluation of collision detection performance achieved by competing pose forecasting techniques, with indication of inference run time. See discussion in Sec. 6.2.

<i>Time Horizon (msec)</i>	1000			
<i>Metrics</i>	<i>Prec</i>	<i>Recall</i>	<i>F₁</i>	<i>Inference Time (sec)</i>
DCT-RNN-GCN [49]	0.63	0.58	0.56	9.1×10^{-3}
MSR-GCN [17]	0.63	0.30	0.31	25.2×10^{-3}
STS-GCN [65]	0.68	0.61	0.63	2.3×10^{-3}
SeS-GCN (proposed)	0.84	0.54	0.64	2.3×10^{-3}

7 Conclusions

Towards the goal of forecasting the human motion during human-robot collaboration in industrial (HRC) environments, we have proposed the novel SeS-GCN model, which integrates three most recent modelling methodologies for accuracy and efficiency: space-time separable GCNs, depth-wise separable graph convolutions and sparse GCNs. Also, we have contributed a new CHICO dataset, acquired at real assembly line, the first providing a benchmark of the two fundamental HRC tasks of human pose forecasting and collision detection. Featuring an MPJPE error of 85.3 mm at 1 sec in the future with a negligible run time of 2.3 msec, SeS-GCN and CHICO unleash great potential for perception algorithms and their application in robotics.

Acknowledgements. This work was supported by the Italian MIUR through the project "Dipartimenti di Eccellenza 2018-2022", and partially funded by DsTech S.r.l.

References

1. Aksan, E., Kaufmann, M., Hilliges, O.: Structured prediction helps 3d human motion modelling. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV) (October 2019) [3](#)
2. Bai, S., Kolter, J.Z., Koltun, V.: An empirical evaluation of generic convolutional and recurrent networks for sequence modeling. ArXiv [abs/1803.01271](#) (2018) [3](#), [7](#)
3. Balcilar, M., Renton, G., Héroux, P., Gaüzère, B., Adam, S., Honeine, P.: Spectral-designed depthwise separable graph neural networks. In: Proceedings of Thirty-seventh International Conference on Machine Learning (ICML 2020)-Workshop on Graph Representation Learning and Beyond (GRL+ 2020) (2020) [4](#), [6](#)
4. Bauer, A., Wollherr, D., Buss, M.: Human-robot collaboration: a survey. International Journal of Humanoid Robotics **5**(01), 47–66 (2008) [4](#)
5. Beltran, E.P., Diwa, A.A.S., Gales, B.T.B., Perez, C.E., Saguisag, C.A.A., Serano, K.K.D.: Fuzzy logic-based risk estimation for safe collaborative robots. In: 2018 IEEE 10th International Conference on Humanoid, Nanotechnology, Information Technology, Communication and Control, Environment and Management (HNICEM). pp. 1–5 (2018) [4](#)
6. Benesova, K., Svec, A., Suppa, M.: Cost-effective deployment of bert models in serverless environment (2021) [11](#)
7. Bertasius, G., Wang, H., Torresani, L.: Is space-time attention all you need for video understanding? In: Proceedings of the International Conference on Machine Learning (ICML) (2021) [6](#)
8. Bütepage, J., Kjellström, H., Kragic, D.: Anticipating many futures: Online human motion prediction and synthesis for human-robot collaboration. ArXiv [abs/1702.08212](#) (2017) [3](#)
9. Cai, Y., Huang, L., Wang, Y., Cham, T.J., Cai, J., Yuan, J., Liu, J., Yang, X., Zhu, Y., Shen, X., Liu, D., Liu, J., Thalmann, N.M.: Learning progressive joint propagation for human motion prediction. In: The European Conference on Computer Vision (ECCV) (2020) [3](#)
10. Cao, Z., Hidalgo Martinez, G., Simon, T., Wei, S., Sheikh, Y.A.: Openpose: Real-time multi-person 2d pose estimation using part affinity fields. IEEE Transactions on Pattern Analysis and Machine Intelligence (2019) [4](#)
11. Castro, A., Silva, F., Santos, V.: Trends of human-robot collaboration in industry contexts: Handover, learning, and metrics. Sensors **21**(12), 4113 (2021) [4](#)
12. Chen, J.H., Song, K.T.: Collision-Free Motion Planning for Human-Robot Collaborative Safety under Cartesian Constraint. IEEE Int. Conf. Robot. Autom. pp. 4348–4354 (2018) [4](#)
13. Chollet, F.: Xception: Deep learning with depthwise separable convolutions. In: 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 1800–1807 (2017) [6](#), [10](#)
14. Costanzo, M., De Maria, G., Lettera, G., Natale, C.: A multimodal approach to human safety in collaborative robotic workcells. IEEE Transactions on Automation Science and Engineering **PP**, 1–15 (01 2021) [4](#)
15. Cui, Q., Sun, H., Yang, F.: Learning dynamic relationships for 3d human motion prediction. In: 2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 6518–6526 (2020) [2](#)
16. Dallel, M., Havard, V., Baudry, D., Savatier, X.: Inhard - industrial human action recognition dataset in the context of industrial collaborative robotics. In: 2020 IEEE International Conference on Human-Machine Systems (ICHMS) (2020) [4](#), [5](#)

17. Dang, L., Nie, Y., Long, C., Zhang, Q., Li, G.: MSR-GCN: Multi-scale residual graph convolution networks for human motion prediction. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV) (2021) [2](#), [3](#), [4](#), [6](#), [9](#), [11](#), [12](#), [14](#)
18. Duarte, N.F., Raković, M., Tasevski, J., Coco, M.I., Billard, A., Santos-Victor, J.: Action anticipation: Reading the intentions of humans and robots. *IEEE Robotics and Automation Letters* **3**(4), 4132–4139 (2018) [3](#)
19. Fieraru, M., Zanfir, M., Oneata, E., Popa, A.I., Olaru, V., Sminchisescu, C.: Three-dimensional reconstruction of human interactions. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition. pp. 7214–7223 (2020) [4](#), [5](#)
20. Fragkiadaki, K., Levine, S., Felsen, P., Malik, J.: Recurrent network models for human dynamics. In: 2015 IEEE International Conference on Computer Vision (ICCV). pp. 4346–4354 (2015) [3](#)
21. Garcia-Esteban, J.A., Piardi, L., Leitao, P., Curto, B., Moreno, V.: An interaction strategy for safe human Co-working with industrial collaborative robots. *Proc. - 2021 4th IEEE Int. Conf. Ind. Cyber-Physical Syst. ICPS 2021* pp. 585–590 (2021) [4](#)
22. Gehring, J., Auli, M., Grangier, D., Yarats, D., Dauphin, Y.N.: Convolutional sequence to sequence learning. In: The International Conference on Machine Learning (ICML) (2017) [3](#), [7](#)
23. Gopalakrishnan, A., Mali, A., Kifer, D., Giles, L., Ororbia, A.G.: A neural temporal model for human motion prediction. In: 2019 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 12108–12117 (2019) [3](#)
24. Gualtieri, L., Palomba, I., Wehrle, E.J., Vidoni, R.: The Opportunities and Challenges of SME Manufacturing Automation: Safety and Ergonomics in Human–Robot Collaboration. Springer International Publishing (2020) [2](#)
25. Guo, W., Bie, X., Alameda-Pineda, X., Moreno-Noguer, F.: Multi-person extreme motion prediction with cross-interaction attention. *arXiv preprint arXiv:2105.08825* (2021) [4](#), [5](#)
26. Haddadin, S., Albu-Schaffer, A., Frommberger, M., Rossmann, J., Hirzinger, G.: The “dlr crash report”: Towards a standard crash-testing protocol for robot safety-part i: Results. In: 2009 IEEE International Conference on Robotics and Automation. pp. 272–279. IEEE (2009) [8](#)
27. He, K., Zhang, X., Ren, S., Sun, J.: Deep residual learning for image recognition. In: 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 770–778 (2016) [10](#)
28. Hinton, G., Dean, J., Vinyals, O.: Distilling the knowledge in a neural network. In: NIPS. pp. 1–9 (2014) [11](#)
29. Hjorth, S., Chrysostomou, D.: Human–robot collaboration in industrial environments: A literature review on non-destructive disassembly. *Robotics and Computer-Integrated Manufacturing* **73**, 102–208 (2022) [8](#)
30. Howard, A.G., Zhu, M., Chen, B., Kalenichenko, D., Wang, W., Weyand, T., Andreetto, M., Adam, H.: Mobilenets: Efficient convolutional neural networks for mobile vision applications (2017) [4](#), [6](#), [10](#)
31. Ionescu, C., Papava, D., Olaru, V., Sminchisescu, C.: Human3.6m: Large scale datasets and predictive methods for 3d human sensing in natural environments. *IEEE Transactions on Pattern Analysis and Machine Intelligence* **36**(7) (2014) [3](#), [4](#), [9](#), [10](#), [12](#), [13](#)
32. ISO: ISO/TS 15066:2016. Robots and robotic devices — Collaborative robots (2021), <https://www.iso.org/obp/ui/#iso:std:iso:ts:15066:ed-1:v1:en> [1](#)

33. Jain, A., Zamir, A.R., Savarese, S., Saxena, A.: Structural-rnn: Deep learning on spatio-temporal graphs. In: 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 5308–5317 (2016) [3](#)
34. Kanazawa, A., Kinugawa, J., Kosuge, K.: Adaptive Motion Planning for a Collaborative Robot Based on Prediction Uncertainty to Enhance Human Safety and Work Efficiency. *IEEE Trans. Robot.* **35**(4), 817–832 (2019) [4](#)
35. Kang, S., Kim, M., Kim, K.: Safety Monitoring for Human Robot Collaborative Workspaces. *Int. Conf. Control. Autom. Syst.* **2019-October**(Iccas), 1192–1194 (2019) [4](#)
36. Knudsen, M., Kaivo-oja, J.: Collaborative robots: Frontiers of current literature. *Journal of Intelligent Systems: Theory and Applications* **3**, 13–20 (06 2020) [1](#)
37. Lai, G., Liu, H., Yang, Y.: Learning graph convolution filters from data manifold (2018) [2](#), [6](#), [10](#)
38. Laplaza, J., Pumarola, A., Moreno-Noguer, F., Sanfeliu, A.: Attention deep learning based model for predicting the 3d human body pose using the robot human handover phases. In: 2021 30th IEEE International Conference on Robot & Human Interactive Communication (RO-MAN). pp. 161–166. IEEE (2021) [4](#)
39. LeCun, V., Denker, J., Solla, S.: Optimal brain damage. In: *Advances in Neural Information Processing Systems* (1989) [7](#)
40. Lemmerz, K., Glogowski, P., Kleineberg, P., Hypki, A., Kuhlenkötter, B.: A Hybrid Collaborative Operation for Human-Robot Interaction Supported by Machine Learning. *Int. Conf. Hum. Syst. Interact. HSI* **2019-June**, 69–75 (2019) [4](#)
41. Li, C., Zhang, Z., Sun Lee, W., Hee Lee, G.: Convolutional sequence to sequence model for human dynamics. In: *The IEEE Conference on Computer Vision and Pattern Recognition (CVPR)* (2018) [3](#), [7](#)
42. Li, M., Chen, S., Zhao, Y., Zhang, Y., Wang, Y., Tian, Q.: Dynamic multiscale graph neural networks for 3d skeleton based human motion prediction. In: 2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 211–220 (2020) [2](#), [3](#)
43. Li, M., Lin, J., Ding, Y., Liu, Z., Zhu, J.Y., Han, S.: Gan compression: Efficient architectures for interactive conditional gans. In: 2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 5283–5293 (2020) [11](#)
44. Li, X., Li, D.: Gpfs: A graph-based human pose forecasting system for smart home with online learning. *ACM Trans. Sen. Netw.* **17**(3) (2021) [2](#), [3](#)
45. Lim, J., Lee, J., Lee, C., Kim, G., Cha, Y., Sim, J., Rhim, S.: Designing path of collision avoidance for mobile manipulator in worker safety monitoring system using reinforcement learning. *ISR 2021 - 2021 IEEE Int. Conf. Intell. Saf. Robot.* pp. 94–97 (2021) [4](#)
46. Liu, S., Deng, W.: Very deep convolutional neural network based image classification using small training sample size. In: 2015 3rd IAPR Asian Conference on Pattern Recognition (ACPR). pp. 730–734 (2015) [10](#)
47. Magrini, E., Ferraguti, F., Ronga, A.J., Pini, F., De Luca, A., Leali, F.: Human-robot coexistence and interaction in open industrial cells. *Robotics and Computer-Integrated Manufacturing* **61** (2020) [9](#)
48. Mahmood, N., Ghorbani, N., Troje, N.F., Pons-Moll, G., Black, M.J.: AMASS: Archive of motion capture as surface shapes. In: *International Conference on Computer Vision* (2019) [4](#), [5](#)
49. Mao, W., Liu, M., Salzmann, M.: History repeats itself: Human motion prediction via motion attention. In: *The European Conference on Computer Vision (ECCV)* (2020) [3](#), [4](#), [6](#), [9](#), [11](#), [12](#), [13](#), [14](#)

50. Mao, W., Liu, M., Salzmann, M., Li, H.: Learning trajectory dependencies for human motion prediction. In: The IEEE International Conference on Computer Vision (ICCV) (2019) [2](#), [9](#)
51. von Marcard, T., Henschel, R., Black, M., Rosenhahn, B., Pons-Moll, G.: Recovering accurate 3d human pose in the wild using imus and a moving camera. In: European Conference on Computer Vision (ECCV) (sep 2018) [4](#), [5](#)
52. Martinez, J., Black, M.J., Romero, J.: On human motion prediction using recurrent neural networks. In: The IEEE Conference on Computer Vision and Pattern Recognition (CVPR) (2017) [8](#), [9](#), [13](#)
53. Matthias, B., Reisinger, T.: Example application of ISO/TS 15066 to a collaborative assembly scenario. 47th Int. Symp. Robot. ISR 2016 **2016**, 88–92 (2016) [2](#)
54. Michalos, G., Makris, S., Tsarouchi, P., Guasch, T., Kontovrakis, D., Chryssolouris, G.: Design considerations for safe human-robot collaborative workplaces. *Procedia CIRP* **37**, 248–253 (2015) [8](#)
55. Minelli, M., Sozzi, A., De Rossi, G., Ferraguti, F., Setti, F., Muradore, R., Bonfè, M., Secchi, C.: Integrating model predictive control and dynamic waypoints generation for motion planning in surgical scenario. In: IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS). pp. 3157–3163 (2020) [14](#)
56. Molchanov, P., Tyree, S., Karras, T., Aila, T., Kautz, J.: Pruning convolutional neural networks for resource efficient inference (2017) [11](#)
57. Nascimento, H., Mujica, M., Benoussaad, M.: Collision avoidance in human-robot interaction using kinect vision system combined with robot’s model and data. *IEEE Int. Conf. Intell. Robot. Syst.* pp. 10293–10298 (2020) [4](#)
58. Oono, K., Suzuki, T.: Graph neural networks exponentially lose expressive power for node classification. In: International Conference on Learning Representations (2020) [7](#), [10](#)
59. Pavlo, D., Feichtenhofer, C., Grangier, D., Auli, M.: 3d human pose estimation in video with temporal convolutions and semi-supervised training. In: Conference on Computer Vision and Pattern Recognition (CVPR) (2019) [3](#)
60. Ramon, J.A.C., Herias, F.A.C., Torres, F.: Safe human-robot interaction based on dynamic sphere-swept line bounding volumes. *Robot. Comput. Integr. Manuf.* **27**(1), 177–185 (2011) [5](#)
61. Rastegari, M., Ordonez, V., Redmon, J., Farhadi, A.: Xnor-net: Imagenet classification using binary convolutional neural networks (2016) [11](#)
62. Rodriguez-Guerra, D., Sorrosal, G., Cabanes, I., Calleja, C.: Human-Robot Interaction Review: Challenges and Solutions for Modern Industrial Environments. *IEEE Access* **9**, 108557–108578 (2021) [1](#)
63. Shah, J., Wiken, J., Breazeal, C., Williams, B.: Improved human-robot team performance using Chaski, a human-inspired plan execution system. *HRI 2011 - Proc. 6th ACM/IEEE Int. Conf. Human-Robot Interact.* pp. 29–36 (2011) [1](#), [2](#)
64. Shi, L., Wang, L., Long, C., Zhou, S., Zhou, M., Niu, Z., Hua, G.: Sparse graph convolution network for pedestrian trajectory prediction. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (2021) [2](#), [4](#), [6](#), [7](#), [10](#), [11](#)
65. Sofianos, T., Sampieri, A., Franco, L., Galasso, F.: Space-time-separable graph convolutional network for pose forecasting. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV) (2021) [2](#), [3](#), [4](#), [6](#), [7](#), [10](#), [11](#), [12](#), [13](#), [14](#)

66. Torkar, C., Yahyanejad, S., Pichler, H., Hofbaur, M., Rinner, B.: Rnn-based human pose prediction for human-robot interaction. In: Proceedings of the ARW & OAGM Workshop 2019. pp. 76–80 (2019) [4](#)
67. Tu, H., Wang, C., Zeng, W.: Voxelpose: Towards multi-camera 3d human pose estimation in wild environment. In: European Conference on Computer Vision. pp. 197–212. Springer (2020) [8](#)
68. Vianello, L., Mouret, J.B., Dalin, E., Aubry, A., Ivaldi, S.: Human posture prediction during physical human-robot interaction. IEEE Robotics and Automation Letters (2021) [4](#)
69. Wang, C., Wang, Y., Huang, Z., Chen, Z.: Simple baseline for single human motion forecasting. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV) Workshops. pp. 2260–2265 (2021) [2](#)
70. Xie, S., Girshick, R., Dollár, P., Tu, Z., He, K.: Aggregated residual transformations for deep neural networks. In: 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 5987–5995 (2017) [10](#)
71. Yu, C., Ma, X., Ren, J., Zhao, H., Yi, S.: Spatio-temporal graph transformer networks for pedestrian trajectory prediction. In: The European Conference on Computer Vision (ECCV) (2020) [6](#)
72. Zeiler, M.D., Fergus, R.: Visualizing and understanding convolutional networks. In: European Conference on Computer Vision (ECCV) (2014) [10](#)
73. Zhang, J., Liu, H., Chang, Q., Wang, L., Gao, R.X.: Recurrent neural network for motion trajectory prediction in human-robot collaborative assembly. CIRP annals **69**(1), 9–12 (2020) [4](#)
74. Zhao, Y., Dou, Y.: Pose-forecasting aided human video prediction with graph convolutional networks. IEEE Access **8**, 147256–147264 (2020) [2](#), [3](#)