

TailorGAN: Making User-Defined Fashion Designs

Lele Chen[†] Justin Tian[†] Guo Li
 University of Rochester

lchen63@ur.rochester.edu, jtian@pas.rochester.edu, gli27@ur.rochester.edu

Cheng-Haw Wu Erh-Kan King Kuan-Ting Chen Shao-Hang Hsieh
 Viscovery

{jason.wu, jeff.king, dannie.chen, shao}@viscovery.com

Chenliang Xu
 University of Rochester
 chenliang.xu@rochester.edu
[†]equal contribution

Abstract

Attribute editing has become an important and emerging topic of computer vision. In this paper, we consider a task: given a reference garment image A and another image B with target attribute (collar/sleeve), generate a photo-realistic image which combines the texture from reference A and the new attribute from reference B . The highly convoluted attributes and the lack of paired data are the main challenges to the task. To overcome those limitations, we propose a novel self-supervised model to synthesize garment images with disentangled attributes (e.g., collar and sleeves) without paired data. Our method consists of a reconstruction learning step and an adversarial learning step. The model learns texture and location information through reconstruction learning. And, the model's capability is generalized to achieve single-attribute manipulation by adversarial learning. Meanwhile, we compose a new dataset, named GarmentSet, with annotation of landmarks of collars and sleeves on clean garment images. Extensive experiments on this dataset and real-world samples demonstrate that our method can synthesize much better results than the state-of-the-art methods in both quantitative and qualitative comparisons.

1. Introduction

Deep generative techniques [9, 15] have led to highly successful image/video generation, some focusing on style transfer [37], and others on the synthesis of desired conditions [16, 25]. In this paper, we propose a novel schema to disentangle attributes and synthesize high-quality images



Figure 1. Illustration of our task and results. Given a reference fashion garment (on the left) and a desired design attribute (in the middle), we aim to generate a new garment that seamlessly integrates the desired design attribute to the reference image. The first row shows collar-editing, and the second row shows sleeve-editing. Our results are shown on the right.

with the desired attribute while keeping other attributes unchanged. Concretely, we focus on the problem of fashion image attribute manipulation to demonstrate the capability of our method. By our method, users can switch a specific part of garments to the wanted designs (e.g., round collar to V-collar). The objective is to synthesize a photo-realistic new fashion image by combining different parts seamlessly together. Potential applications of such system range from item retrieval to professional fashion design assistant. Fig. 1 shows our task and results.

Recently, Generative Adversarial Networks (GANs) [9, 16, 25], image-to-image translation [10], and CycleGAN [37] have proved their effectiveness in generating photo-realistic images. However, image syntheses using such models usually involve highly-entangled attributes and may fail in editing target-specific attributes/objects [19, 30, 37]. Meanwhile, the large variance in the garment tex-

tures causes additional problems. It is impossible to build a dataset that is large enough to approximate the distribution of all garment texture and design combinations, which serve as paired samples to train models such as [25, 10]. Novel learning paradigm is expected to overcome this difficulty.

To solve these challenges, we propose a novel self-supervised image generative model, TailorGAN, that makes disentangled attribute manipulations. The model exploits the latent space expression of input images without paired training images. Image edge maps are used to isolate the structures from textures. By using edge maps, we achieve good results with only a small amount of data. It is also robust to a small amount of geometric misalignment between the reference and design attribute images. The model is further generalized in a GAN framework to achieve single-attribute manipulation using random fashion inputs. Besides, an attribute-aware discriminator is weaved into the model to guide the image editing. This attribute-aware discriminator helps in making high-quality single attribute editing and guides a better self-supervised learning process.

Current available fashion image datasets (*e.g.*, DeepFashion [14]) mostly consist of street photos with complex backgrounds or with user’s body parts presented. That extra visual information may hinder the performance of an image synthesis model. Thus, to simplify the training data and screen out noisy backgrounds, we introduce a new dataset, GarmentSet. The new dataset contains fashion images with no human user presented, and most images have single-color backgrounds.

Contributions made in this paper can be summarized as:

- We propose a new task in deep learning fashion studies. Instead of generating images with text guidance, virtual try-on, or texture transferring, our task is to make new fashion designs with disentangled user-appointed attributes.
- We exploit a novel training schema consists of reconstruction learning and adversarial learning. The self-supervised reconstruction learning guides the network to learn shape, location information from the edge map, and texture information from the RGB image. The unpaired adversarial learning gives the network generalizability to synthesize images with new disentangled attributes. Besides, we propose a novel attribute-aware discriminator, which helps high-quality attribute editing by isolating the design structures from the garment textures.
- A new dataset, GarmentSet, is introduced to serve our attribute editing task. Unlike existing fashion datasets in which most images illustrate the user’s face or body parts with complicated backgrounds, GarmentSet filters out the most redundant information and directly serves the fashion design purpose.

The rest of this paper is organized as follows: A brief review of related work is presented in Sec. 2. The details of the proposed method are described in Sec. 3. We introduce our new dataset in Sec. 4. Experimental details and results are presented in Sec. 5. In Sec. 5.6, we further conduct ablation studies to investigate the performances of our model explicitly. And finally, Sec. 6 concludes the paper with discussions of limitations. Our code and data are available at: <https://www.cs.rochester.edu/~cxu22/r/tailorgan/>.

2. Related Work

Generative Adversarial Network (GAN) [9] is one of the most popular deep generative models and has shown impressive results in image synthesis studies, like image editing [23, 29] and fine-grained objects generation [25, 13]. Researchers utilize different conditions to generate images with desired properties. Existing works have explored various conditions, from category labels [22], audio [6, 5], text [25], skeleton [20, 11, 36, 24] to attributes [27]. There are a few studies that investigate the task of image translations using cGAN [16, 25, 4]. In the context of fashion-related applications, researchers apply cGAN in automated garment textures filling [32], texture transferring [12], and virtual try-on [38, 31] by replacing dress on a person with a new one. More related work is sequential attention GAN proposed by Cheng *et al.* [7]. Their model uses text as the guidance and continuously changes the fashion designs based on user’s requests, but the attribute changes are highly entangled. In contrast to this work, we propose a new training algorithm that combines self-supervised reconstruction learning with adversarial learning to make disentangled attribute manipulations with user-appointed images.

Self-supervised generation is recently introduced as a novel and effective way to train generative models without paired training data. Unpaired image-to-image translation framework such as CycleGAN [37] removes pixel-level supervision. In CycleGAN, the unpaired image to image translation is achieved by enforcing a bi-directional translation between two domains with an adversarial penalty on the translated image in the target domain. The CycleGAN variants [35, 17] are moving towards the direction of unsupervised learning approaches. However, CycleGAN-family models also create unexpected or even unwanted results, which are shown in the experiment section. One reason for such a phenomenon is due to the lack of straightforward knowledge of the target translation domain in the circularity training process. Inherent attributes of the source samples may be changed in a translation process. To avoid such unwanted changes, we keep an image reconstruction penalty in our image editing task.

Attracted by the huge profit potentials in the fashion industries, deep learning methods have been conducted on

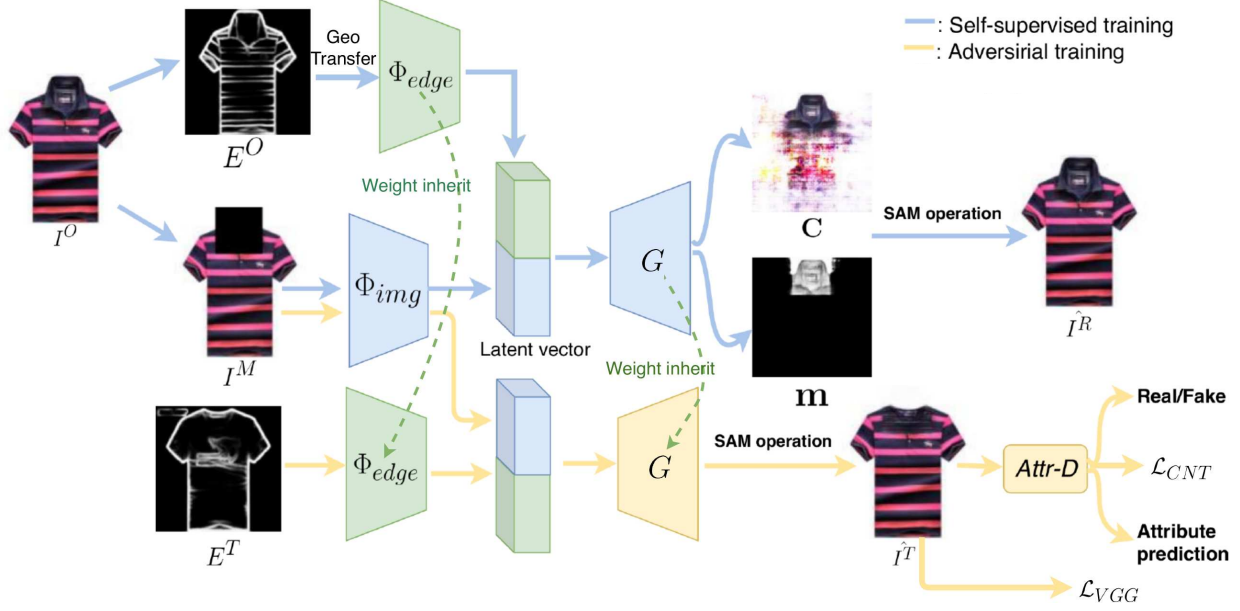


Figure 2. The network architecture of TailorGAN. The Geo-Transfer indicates random rotations, translations, and scale shifts applied on the edge map. This operation can make the trained model robust to different edge maps. The weight inherit means the weights of Φ_{edge} and G in adversarial step are inherited from the reconstruction step. In the reconstruction step (the upper part), we extract the edge feature (e.g. location and shape of the collar) by Φ_{edge} and extract image feature (e.g. overall texture and design of the garment) by Φ_{img} . Then, we merge the two vectors in the latent space and pass the resultant vector through image generator G to output mask and attention. In the adversarial training step (the lower part), we use Φ_{edge} to extract features from the target-attribute edge map E^T , which has a new collar. The attribute discriminator (Attr-D) outputs a real/fake score and classifies the collar type of the newly-generated image $\hat{I}^T$ yield by G .

fashion analysis and fashion image synthesis. Most existing researches focus on fashion trend prediction [2], clothing recognition with landmarks [14], clothing matching with fashion items in street photos [18] and fashion recommendation system [3, 33]. Different from those research lines, we focus on the fashion image synthesis task with user-appointed attribute manipulations.

3. Methodology

This section presents the implementation details of our method. There are two steps in the model training: (1) self-supervised reconstruction learning and (2) generalized attribute manipulations using adversarial learning. The first step helps the model fill correct texture for the user-appointed design pattern. The second step helps the model generate high-quality images with desired new attributes. Although we use collar translating example in Fig. 2, we emphasize here that our model can be applied to other design patterns as shown in experiments.

3.1. Learning to Manipulate Designs

Self-Supervised Learning Step. The motivation of formulating a self-supervised model is that it is almost impossible to collect paired training images for a fully supervised model. Using the collar editing task as an example,

for each image in a fully supervised training process, one needs to collect paired images for each collar type. In these paired images, only the collar parts are different while the other attributes, like body decoration, clothing texture, etc., must stay unchanged and match with other paired images. Such data is usually unavailable. Also, the dataset size increases exponentially when multiple attribute annotations are needed for each image.

Therefore, we employ an encoder-decoder structure for the self-supervised reconstruction training step. Given a masked garment image I^M (mask out the collar/sleeve part from the original garment image I^O) and edge map E^O , our reconstruction step reconstructs the original garment image (collar region). Intuitively, individual fashion designs may correlate with clothing texture and other design factors. For instance, light pink color is rarely used on men’s garments, and leather is usually used to make jackets, etc. In our task, we want our model focus only on the design structure editing rather than the clothing texture translating, which is inherited from the reference garment image. Specifically, the self-supervised learning step is defined as:

$$\hat{I}^R = SAM(G(\Phi_{img}(I^M) \oplus \Phi_{edge}(E^O)), I^M) , \quad (1)$$

where Φ_{img} and Φ_{edge} are image encoder and edge encoder, respectively. Φ_{img} and Φ_{edge} consist of several 2D convolution layers and residual blocks. G is the image de-

coder/generator, which consists of several 2D transpose-convolution layers. $\oplus$ is channel-wise concatenation. After encoding, we feed the concatenated latent vector to G to output attention mask $\mathbf{m}$ and new pixel $\mathbf{C}$. The SAM block (see Eq. 3) outputs the reconstructed image $\hat{I}^R$ based on $\mathbf{m}$, $\mathbf{C}$ and I^M . This learning step learns how to reconstruct I^O according to the texture feature from I^M and shape feature from E^O .

Different from other methods [6, 5], we only use pixel-wise loss here to supervise the reconstruction step. Specifically, the loss function for this training step is defined as:

$$\mathcal{L}_R(\hat{I}^R, I^O) = \|I^O - \hat{I}^R\|_1. \quad (2)$$

From the reconstruction, the network learns to fill the garment texture and locate the desired design to output the original unmasked image I^O . We empirically find that this learning step is critical to the full model performance. During the training, we apply random rotations, translations, and scale shifts to the input edge maps. Thus, the model is trained to handle potential geometric misalignment and can locate the desired fashion pattern at the correct position.

Self-Attention Mask Operation (SAM). During the reconstruction step, the model should learn to change only the target part and keep the rest of an image untouched. Thus, we introduce a self-attention mask to the generator. The self-attention mechanism can guide the model to focus on the target region. Concretely, the decoder/generator G produces two outputs: single-channel self-attention mask $\mathbf{m}$ and new pixel image $\mathbf{C}$. The final output combines the masked color image with the input cropped image I^M :

$$\hat{I}_{i,j}^R = \mathbf{m}_{i,j} \times \mathbf{C}_{i,j} + (1 - \mathbf{m}_{i,j}) \times I_{i,j}^M, \quad (3)$$

where $\mathbf{m}_{i,j}$, $\mathbf{C}_{i,j}$ and $\hat{I}_{i,j}^R$ are the pixels at i^{th} row and j^{th} column in the self-attention mask, the new pixel image, and the final output image. The self-attention mask layer and the color layer share the bottom transpose convolutional blocks in the decoder. The output of the last transpose convolutional layer is fed into two activation layers: a Sigmoid layer with a single-channel output (self-attention mask) and a hyperbolic tangent layer with three-channel output (the new pixel image). The self-attention mask guides the network to focus on the attribute-related region while training the network in a fully self-supervised manner.

3.2. Generalized Single Attribute Manipulations

We have introduced the self-supervised reconstruction learning in Sec. 3.1, which can reconstruct the original image. However, our task is to synthesize new images by manipulating the attributes. The model trained with the reconstruction step could not yield good results since it is not generalized to synthesize a new image with other new attributes (e.g., new collar type, and new sleeve type). Meanwhile, the synthesized image with the reconstruction step is

blurry, which makes the results unrealistic. To tackle those problems, we have another adversarial learning step, which inherits the encoder-decoder network from the reconstruction step and applies a novel attribute-aware discriminator.

Perceptual Loss. To improve the image quality and preserve the overall texture of the original garment image I^O , we use VGG perceptual loss:

$$\mathcal{L}_{VGG}(\hat{I}^T, I^O) = \|\phi(I^O) - \phi(\hat{I}^T)\|_1, \quad (4)$$

where ϕ is a VGG-19 [28] feature extractor pre-trained on ImageNet [8].

Attribute-Aware Losses. To enforce the model to output image with correct attributes, we propose a discriminator with two different regressions: a binary (real/fake) label and an attribute prediction vector. The attribute prediction is optimized by cross-entropy loss, which is defined as:

$$\begin{aligned} \mathcal{L}_{ATT}(I, \vec{V}) = & -\vec{V} * \log(f(I)) \\ & - (1 - \vec{V}) * \log(1 - f(I)), \end{aligned} \quad (5)$$

where $\vec{V}$ denotes the one-hot vector of the target attribute type and f is an attribute classifier that outputs the class label vector of image I . During the adversarial training, when we forward I^O and its corresponding $\vec{V}^O$ to the discriminator, the discriminator's parameters are updated to learn how to classify the attributes; when we forward $\hat{I}^T$ and the desired target $\vec{V}^T$ to the discriminator, we do not update the discriminator's parameters, but update the generator's parameters.

By predicting the attribute type, ideally, the discriminator should focus on the replaced design region. Therefore, we introduce an additional feature-level distance measure between $\hat{I}^T$ and I^T to further improve the image quality at the perception-level:

$$\mathcal{L}_{CNT}(\hat{I}^T, I^T) = \|D_{conv}(I^T) - D_{conv}(\hat{I}^T)\|_1, \quad (6)$$

where D_{conv} denotes the feature output from the convolutional layers of the discriminator.

By combining the losses above, the full loss to train the generator can be formulated as:

$$\begin{aligned} \mathcal{L}_G = & (1 - D(G(E^T, I^M)))^2 + \mathcal{L}_{CNT}(\hat{I}^T, I^T) \\ & + \lambda_1 * \mathcal{L}_{ATT}(\hat{I}^T, \vec{V}^T) + \lambda_2 * \mathcal{L}_{VGG}(\hat{I}^T, I^O), \end{aligned} \quad (7)$$

where λ_1 and λ_2 are hyper-parameters to balance the loss terms. We set $\lambda_1 = 0.1$, $\lambda_2 = 2.5$ in all experiments. In our empirical study, we observe that the model is sensitive to λ_2 and we need to choose different λ_2 if we use different VGG layer features to compute the perceptual loss. Similarly, the full loss to train the discriminator is:

$$\begin{aligned} \mathcal{L}_D = & \frac{1}{2} [D(G(E^T, I^M))^2 + (D(I^O) - 1)^2] \\ & + \mathcal{L}_{ATT}(I^O, \vec{V}^O). \end{aligned} \quad (8)$$



Figure 3. Samples of each collar type and sleeve type from GarmentSet dataset. The pie charts demonstrate the collar type and sleeve type distributions. There are in total 12 collar types and 2 sleeve types.

Algorithm 1 Training steps

Input: learning rate α , batch size B

Output: generator parameter θ_G , discriminator parameter θ_D

for number of iterations **do**

Sample $\{I_i^O\}_{i=1}^B$

Extract edge maps and masked-out images $\{E_i^O, I_i^M\}_{i=1}^B$

$\{\hat{I}_i^R\}_{i=1}^B \leftarrow G_\theta(\{E_i^O, I_i^M\}_{i=1}^B)$

Update the generator G in the reconstruction step:

$\theta_G \leftarrow \text{Adam}\{\frac{1}{B} \sum_{i=1}^B \nabla_{\theta_G} \mathcal{L}_R(\hat{I}_i^R, I_i^O), \alpha\}$

for number of iterations **do**

Sample $\{I_i^O, \tilde{V}_i^O, I_i^T, \tilde{V}_i^T\}_{i=1}^B$

Extract edge maps and masked-out images $\{E_i^T, I_i^M\}_{i=1}^B$

$\{\hat{I}_i^T\}_{i=1}^B \leftarrow G_\theta(\{E_i^T, I_i^M\}_{i=1}^B)$

Update the discriminator D in the adversarial step:

$\theta_D \leftarrow \text{Adam}\{\frac{1}{B} \sum_{i=1}^B \nabla_{\theta_D} \mathcal{L}_D(\hat{I}_i^T, I_i^O, \tilde{V}_i^O), \alpha\}$

Update the generator G in the adversarial step:

$\theta_G \leftarrow \text{Adam}\{\frac{1}{B} \sum_{i=1}^B \nabla_{\theta_G} \mathcal{L}_G(\hat{I}_i^T, I_i^T, \tilde{V}_i^T, I_i^O), \alpha\}$

3.3. Training Algorithm

Combining the self-supervised reconstruction learning step (Sec. 3.1) with the adversarial learning step (Sec. 3.2), our full model is trained in two separate training loops. Specifically, we formulate the training algorithm in Algorithm 1. When training the generator G in the reconstruction step without discriminator D, the generator G receives a masked image I^M and its original edge map E^O as input, and it outputs the reconstructed image $\hat{I}^R$. We try to minimize the reconstruction loss $\mathcal{L}_R$ to enforce the network to learn to generate correct texture based on design sketch and fit to other parts. When training the discriminator D in the adversarial step, the generator G receives masked image I^M and the edge map E^T of a new attribute type as input, and it outputs the edited image $\hat{I}^T$. We try to minimize the generator loss $\mathcal{L}_G$, since there is no paired ground truth in this step. It enforces the network to learn to synthesize images by manipulating the target attribute type E^T . Then we update parameters in generator G in the adversarial step by minimizing the discriminator loss $\mathcal{L}_D$. By optimizing the

loss in reconstruction and adversarial steps, the TailorGAN can learn realistic texture and geometry information to yield high-quality images with user-appointed attribute type.

4. GarmentSet dataset

This part serves as a brief introduction to GarmentSet dataset. Current available datasets like DeepFashion [14] and FashionGen [26] mostly consist of images including the user’s face or body parts and street photos with noisy backgrounds. The redundant information raises unwanted hardness in the training process. To filter out such redundant information in the images, we build a new dataset: GarmentSet. In this dataset, we have 9636 images with collar part annotations and 8616 images with shoulder and sleeve annotations. Both classification types and landmarks are annotated. Although in our studies, the landmark locations are only used in the image pre-processing steps, they still can be useful in future researches like fashion item retrieve, clothing recognition, etc.

In Fig. 3, we present sampled pictures for each collar type, each sleeve type, and the overall data distributions in the dataset. Round collar, V-collar, and lapel images together contribute over eighty percent of the total collar-annotation images. The sleeve dataset only contains two types: short and long sleeves. Although not used in training, the dataset also contains attribute landmark locations, including collars, shoulders, and sleeve ends.

5. Experiments

5.1. Data Pre-processing

In this paper, we randomly sample 80% data for training and 20% for testing. Meanwhile, in Sec. 5.3, we keep one collar type out in the training set to demonstrate the robustness of our model. In the data pre-processing step, we generate masked-out images I^M and edge maps. We

Type	Type 1 $\Rightarrow$ Type 2			Type 2 $\Rightarrow$ Type 1			Type 1 $\Rightarrow$ Type 6			Type 6 $\Rightarrow$ Type 1			Type 2 $\Rightarrow$ Type 6			Type 6 $\Rightarrow$ Type 2		
	C.E.	SSIM	PSNR	C.E.	SSIM	PSNR	C.E.	SSIM	PSNR	C.E.	SSIM	PSNR	C.E.	SSIM	PSNR	C.E.	SSIM	PSNR
CycleGAN	12.48	0.77	18.72	1.74	0.64	13.97	5.21	0.74	23.40	2.97	0.78	18.77	6.02	0.89	18.89	10.02	0.77	17.63
Pix2pix	20.01	0.87	22.75	12.73	0.89	21.42	21.62	0.88	17.63	15.79	0.89	21.88	15.12	0.77	23.15	19.67	0.88	23.04
Ours	9.17	0.89	23.53	3.70	0.89	23.75	6.29	0.90	23.92	3.25	0.90	22.47	5.62	0.89	23.08	8.13	0.90	22.24

Table 1. Measurements for all three models based on target translating types on the testing set. C.E. stands for the average classification error score for each paired type translation. The bold numbers in each column are the best scores.



Figure 4. Testing results of collar editing. The images are clustered based on the target collar types. The generated results of simple collar patterns (round and V-collar) are, in general, better than complicated ones (lapel). We train CycleGAN, pix2pix models separately for each specific transformation.

use the left/right collar landmarks to make a bounding box around the collar region. The holistically-nested edge detection (HED) model [34] takes charge of making edge maps, which is pre-trained on the BSDS dataset [21].

We notice that the image quality of the edited results depends on the input edge map resolution. Since the HED model used is pre-trained on a general image dataset, it struggles in catching structural details and may also generate unwanted noise maps. The HED model produces low-resolution edge maps for a target image with complicated, detailed structures.

5.2. Comparative Studies

We compare our model with CycleGAN [37] and pix2pix [10]. To compare with [37], we have trained three different CycleGAN models: collar type 1 (round collar) $\Leftrightarrow$ type 2 (V-collar), collar type 1 (round collar) $\Leftrightarrow$ type 6 (lapel) and collar type 2 $\Leftrightarrow$ type 6. However, our model is trained with all different types together. The results are shown in Tab. 1. Furthermore, in Tab. 2, we compare our model with pix2pix using random collar types. Since one CycleGAN model cannot handle translation between any two collar types, we drop it in this random translating comparison test.

Qualitative results. In Fig. 4, we present testing results of three models. As one can see in the sampled images,

	C.E.	SSIM	PSNR
Our model	5.78	0.8921	23.36
pix2pix	16.12	0.8783	22.83

Table 2. Testing results of TailorGAN and pix2pix trained on all collar type inputs.

C. E.	type 1 out	type 2 out	type 6 out
full model	3.48	7.41	5.27
one-out model	4.74	8.59	5.34

Table 3. The one-out model V.S. the fully trained model.

CycleGAN does not preserve garment textures. The trained pix2pix model performs poorly in most examples. For the collar part, the pix2pix outputs only show color bulks with no structural patterns.

Quantitative results. For the quantitative comparisons, in measuring model performances, we use three metrics: classification errors (C.E.), structure similarity index (SSIM), and peak signal to noise ratio (PSNR). The classification error is measured with a classifier pretrained on GarmentSet dataset. The classification error measures the distance from the target collar designs. The SSIM and the PSNR scores are derived from the differences between the original image and the edited image. From the numerical results presented in Tab. 1 and Tab. 2, our model outperforms both CycleGAN and pix2pix in making high-quality images with higher classification accuracy. We attribute this to the poor texture/structure preserving of the CycleGAN/pix2pix model.

5.3. Synthesizing Unseen Collar Type

To test our model’s capability of processing collar types that are missing in the dataset, we take one collar type out in the training stage and test the model’s performance on this unseen collar type. In this test, we take collar type 1, 2, and 6 out. We also test the leave-one-out model with a fully trained model that meets all collar types in its training process. The classification errors for each pair of models are calculated for each taken-out collar type in the testing set. The qualitative results are shown in Fig. 5

Based on the qualitative analysis and quantitative comparisons (see Tab. 3), our model shows a strong generalization ability in synthesizing unseen collar types.



Figure 5. Testing results of the unseen target collar type. The left side indicates which type we are generating (remove this type in training set). The 1th, 4th columns are the original reference images. The 2th, 5th columns are images with target collar types. The 3th, 6th columns are generated images with target collar types with texture/style of reference image.

5.4. Sleeve Generation

In the previous discussions, we applied our model in collar part editing. GarmentSet dataset also contains sleeve landmarks and types information. Thus, we test our model's capability in editing sleeves and present testing results in Fig. 6. We use the same training scheme. Instead of collars, the sleeve parts are masked out. Due to simpler edge structures and better resolutions in the edge maps, the edited sleeve images have better image qualities and are close to the real images. We discuss more details of sleeve editing results in the user evaluation section.

5.5. User Evaluations and Item Retrieves

To evaluate the performance in a human perceptive level, we conduct thoughtful user studies in this section. Human subjects evaluation (see Fig. 7) is conducted to investigate the image quality and the attribute (collar) similarity of our generated results compared with [37, 10]. Here, we present the average scores for each model based on twenty users' evaluations. The maximum score is ten. As shown in Fig. 7, our model receives the best scores in both image quality and similarity evaluations.

We also collected users' feedback to the sleeve changing results, and the feedback shows that users can hardly distinguish real/fake between our generated images and the

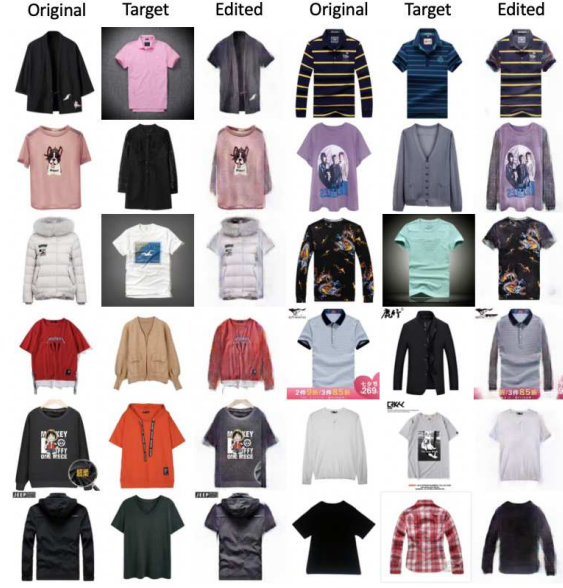


Figure 6. Testing results of editing sleeves using TailorGAN. The 1th, 4th columns are the original reference images. The 2th, 5th columns are images with target sleeve types. The 3th, 6th columns are generated images with target sleeve types with texture/style of reference image.

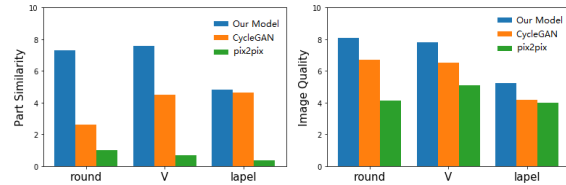


Figure 7. Average user evaluation scores based on image quality and target part similarity. The part similarity is based on users' scores on the edited part structure similarity between the edited image and the target image.

real images. Sleeve tests only evaluate image quality. In each test case, there are two pictures: (1) the original picture and (2) the edited picture. Users decide scores ranging from zero to ten to both pictures based on the image quality. Translating from short to long sleeves is, in general, a harder task due to auto-texture filling and background changing. Users may find that it is harder to distinguish the real images from the edited ones for short sleeve garments. Those observations are reflected in the evaluation scores in Fig. 8.

To prove that TailorGAN can be useful in image item retrieves, we upload the edited images to a searching-based website [1] and show top search results in Fig. 9.

5.6. Ablation Studies

In this section, we want to clarify our choices of two separate training steps are crucial to generating photo-realistic

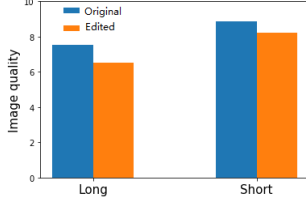


Figure 8. User feedback to sleeve editing results.



Figure 9. Top 5 matching items from [1]. The first column (green box) are generated images and the rest of the column (yellow) are retrieved images based on the generated images from [1].

results. As we argued, to disentangle the texture and the design structure meanwhile keep the rest similar, the model needs to have the ability to reconstruct. Also, using two separate training steps makes sharper results. Tab. 4 shows that without using the reconstruction training step produces blurry images. Similarly, feeding RGB reference images as input instead of gray-scale edge images, the model produces blurry collars since the lack of clear geometric edge information. For qualitative analysis, We plot testing results of RGB inputs and results of only using adversarial training step and compare those changes with our full model.

Fig. 10 shows edited image results of our current model versus w/o reconstruction and w/o edge maps. The model w/o reconstruction does not prioritize high-frequency details and tends to average the pixel values in the editing region. On the other hand, using RGB images as inputs, the results show unexpected effects. The RGB-input includes extra structures that are not related to the target images or original images. We hypothesize that those unexpected structures are from the texture-design entanglement. We measure the classification error, SSIM, and PSNR for three variants. Since the major part of the image is left untouched in the result, the leading score may not be im-

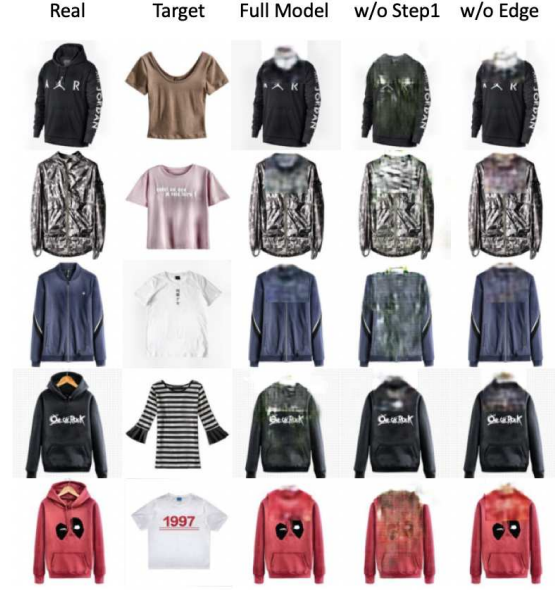


Figure 10. Results of full model, without reconstruction step (w/o reconstruction step) and using RGB pictures (w/o E^T) as inputs instead of edge maps.

	C.E.	SSIM	PSNR
Full model	5.78	0.8921	23.36
w/o reconstruction step	6.34	0.7706	19.46
w/o edge input	7.21	0.8782	21.58

Table 4. Our final model versus two variants on the testing set. w/o reconstruction step represents we use only adversarial training step rather than two steps. w/o edge input indicates that we use RGB images as input rather than grey-scale edge maps.

pressive in numbers. But, through the qualitative analysis, we can confirm that our current model can generate high-quality images with more details preserved.

6. Conclusion

In this paper, we introduce a novel task to the deep learning fashion field. We investigate the problem of doing image editing to a fashion item with user-defined attribute manipulation. Methods developed under this task can be useful in real world applications such as novel garment retrieval and assistant to designers. We propose a novel training schema that can manipulate a single attribute to an arbitrary fashion image. To serve a better model training, we collect our own GarmentSet dataset. Our model outperforms the baseline models and successfully generates photo-realistic images with desired attribute manipulation.

Acknowledgement. This work was partly supported by a research gift from Viscosity.

References

- [1] Taobao. www.TaoBao.com. Accessed: 2019-07-30.
- [2] Z. Al-Halah, R. Stiefelham, and K. Grauman. Fashion forward: Forecasting visual style in fashion. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 388–397, 2017.
- [3] H. Chen, A. C. Gallagher, and B. Girod. Describing clothing by semantic attributes. In *Computer Vision - ECCV 2012 - 12th European Conference on Computer Vision, Florence, Italy, October 7-13, 2012, Proceedings, Part III*, pages 609–623, 2012.
- [4] L. Chen, Z. Li, R. K. Maddox, Z. Duan, and C. Xu. Lip movements generation at a glance. In *Computer Vision - ECCV 2018 - 15th European Conference, Munich, Germany, September 8-14, 2018, Proceedings, Part VII*, pages 538–553, 2018.
- [5] L. Chen, R. K. Maddox, Z. Duan, and C. Xu. Hierarchical cross-modal talking face generation with dynamic pixel-wise loss. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pages 7832–7841, 2019.
- [6] L. Chen, S. Srivastava, Z. Duan, and C. Xu. Deep cross-modal audio-visual generation. In *Proceedings of the on Thematic Workshops of ACM Multimedia 2017, Mountain View, CA, USA, October 23 - 27, 2017*, pages 349–357, 2017.
- [7] Y. Cheng, Z. Gan, Y. Li, J. Liu, and J. Gao. Sequential attention gan for interactive image editing via dialogue. *arXiv preprint arXiv:1812.08352*, 2018.
- [8] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei. ImageNet: A Large-Scale Hierarchical Image Database. In *CVPR09*, 2009.
- [9] I. J. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. C. Courville, and Y. Bengio. Generative adversarial nets. In *Advances in Neural Information Processing Systems 27: Annual Conference on Neural Information Processing Systems 2014, December 8-13 2014, Montreal, Quebec, Canada*, pages 2672–2680, 2014.
- [10] P. Isola, J.-Y. Zhu, T. Zhou, and A. A. Efros. Image-to-image translation with conditional adversarial networks. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 1125–1134, 2017.
- [11] N. Jetchev and U. Bergmann. The conditional analogy gan: Swapping fashion articles on people images. In *Proceedings of the IEEE International Conference on Computer Vision*, pages 2287–2292, 2017.
- [12] S. Jiang and Y. fu. Fashion style generator. *IJCAI*, 2017.
- [13] N. Kato, H. Osone, K. Oomori, C. W. Ooi, and Y. Ochiai. Gans-based clothes design: Pattern maker is all you need to design clothing. In *Proceedings of the 10th Augmented Human International Conference 2019*, page 21. ACM, 2019.
- [14] M. H. Kiapour, X. Han, S. Lazebnik, A. C. Berg, and T. L. Berg. Where to buy it: Matching street clothing photos in online shops. In *2015 IEEE International Conference on Computer Vision, ICCV 2015, Santiago, Chile, December 7-13, 2015*, pages 3343–3351, 2015.
- [15] D. P. Kingma and M. Welling. Auto-encoding variational bayes. *arXiv preprint arXiv:1312.6114*, 2013.
- [16] C. Lassner, G. Pons-Moll, and P. V. Gehler. A generative model of people in clothing. In *IEEE International Conference on Computer Vision, ICCV 2017, Venice, Italy, October 22-29, 2017*, pages 853–862, 2017.
- [17] M. Li, H. Huang, L. Ma, W. Liu, T. Zhang, and Y. Jiang. Unsupervised image-to-image translation with stacked cycle-consistent adversarial networks. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 184–199, 2018.
- [18] S. Liu, Z. Song, G. Liu, C. Xu, H. Lu, and S. Yan. Street-to-shop: Cross-scenario clothing retrieval via parts alignment and auxiliary set. In *2012 IEEE Conference on Computer Vision and Pattern Recognition, Providence, RI, USA, June 16-21, 2012*, pages 3330–3337, 2012.
- [19] Y. Lu, Y. Tai, and C. Tang. Attribute-guided face generation using conditional cyclegan. In *Computer Vision - ECCV 2018 - 15th European Conference, Munich, Germany, September 8-14, 2018, Proceedings, Part XII*, pages 293–308, 2018.
- [20] L. Ma, X. Jia, Q. Sun, B. Schiele, T. Tuytelaars, and L. Van Gool. Pose guided person image generation. In *Advances in Neural Information Processing Systems*, pages 406–416, 2017.
- [21] D. Martin, C. Fowlkes, D. Tal, and J. Malik. A database of human segmented natural images and its application to evaluating segmentation algorithms and measuring ecological statistics. In *Proc. 8th Int'l Conf. Computer Vision*, volume 2, pages 416–423, July 2001.
- [22] A. Nguyen, A. Dosovitskiy, J. Yosinski, T. Brox, and J. Clune. Synthesizing the preferred inputs for neurons in neural networks via deep generator networks. In *Advances in Neural Information Processing Systems*, pages 3387–3395, 2016.
- [23] A. Radford, L. Metz, and S. Chintala. Unsupervised representation learning with deep convolutional generative adversarial networks. *arXiv preprint arXiv:1511.06434*, 2015.
- [24] A. Raj, P. Sangkloy, H. Chang, J. Lu, D. Ceylan, and J. Hays. Swapnet: Garment transfer in single view images. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 666–682, 2018.
- [25] S. E. Reed, Z. Akata, X. Yan, L. Logeswaran, B. Schiele, and H. Lee. Generative adversarial text to image synthesis. In *Proceedings of the 33rd International Conference on Machine Learning, ICML 2016, New York City, NY, USA, June 19-24, 2016*, pages 1060–1069, 2016.
- [26] N. Rostamzadeh, S. Hosseini, T. Boquet, W. Stokowiec, Y. Zhang, C. Jauvin, and C. Pal. Fashion-gen: The generative fashion dataset and challenge. *arXiv preprint arXiv:1806.08317*, 2018.
- [27] W. Shen and R. Liu. Learning residual images for face attribute manipulation. In *2017 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2017, Honolulu, HI, USA, July 21-26, 2017*, pages 1225–1233, 2017.
- [28] K. Simonyan and A. Zisserman. Very deep convolutional networks for large-scale image recognition. In *3rd International Conference on Learning Representations, ICLR 2015, San Diego, CA, USA, May 7-9, 2015, Conference Track Proceedings*, 2015.

- [29] X. Wang and A. Gupta. Generative image modeling using style and structure adversarial networks. In *European Conference on Computer Vision*, pages 318–335. Springer, 2016.
- [30] Z. Wang, X. Tang, W. Luo, and S. Gao. Face aging with identity-preserved conditional generative adversarial networks. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 7939–7947, 2018.
- [31] R. Y. X. Han, Z. Wu and L. S. Davis. Viton: an image based virtual try-on network. *CVPR*, 2018.
- [32] W. Xian, P. Sangkloy, V. Agrawal, A. Raj, J. Lu, C. Fang, F. Yu, and J. Hays. Texturegan: Controlling deep image synthesis with texture patches. In *2018 IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2018, Salt Lake City, UT, USA, June 18-22, 2018*, pages 8456–8465, 2018.
- [33] T. Xiao, T. Xia, Y. Yang, C. Huang, and X. Wang. Learning from massive noisy labeled data for image classification. In *IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2015, Boston, MA, USA, June 7-12, 2015*, pages 2691–2699, 2015.
- [34] S. Xie and Z. Tu. Holistically-nested edge detection. In *Proceedings of the IEEE international conference on computer vision*, pages 1395–1403, 2015.
- [35] Y. Yuan, S. Liu, J. Zhang, Y. Zhang, C. Dong, and L. Lin. Unsupervised image super-resolution using cycle-in-cycle generative adversarial networks. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops*, pages 701–710, 2018.
- [36] B. Zhao, X. Wu, Z.-Q. Cheng, H. Liu, Z. Jie, and J. Feng. Multi-view image generation from a single-view. In *2018 ACM Multimedia Conference on Multimedia Conference*, pages 383–391. ACM, 2018.
- [37] J. Zhu, T. Park, P. Isola, and A. A. Efros. Unpaired image-to-image translation using cycle-consistent adversarial networks. In *IEEE International Conference on Computer Vision, ICCV 2017, Venice, Italy, October 22-29, 2017*, pages 2242–2251, 2017.
- [38] S. Zhu, S. Fidler, R. Urtasun, D. Lin, and C. C. Loy. Be your own prada: Fashion synthesis with structural coherence. In *IEEE International Conference on Computer Vision, ICCV 2017, Venice, Italy, October 22-29, 2017*, pages 1689–1697, 2017.