

Geometric Image Synthesis

Hassan Abu Alhaija¹, Siva Karthik Mustikovela¹,
Andreas Geiger^{2,3}, Carsten Rother¹

¹Visual Learning Lab, Heidelberg University

²Autonomous Vision Group, MPI for Intelligent Systems Tübingen

³University of Tübingen

Abstract. The task of generating natural images from 3D scenes has been a long standing goal in computer graphics. On the other hand, recent developments in deep neural networks allow for trainable models that can produce natural-looking images with little or no knowledge about the scene structure. While the generated images often consist of realistic looking local patterns, the overall structure of the generated images is often inconsistent. In this work we propose a trainable, geometry-aware image generation method that leverages various types of scene information, including geometry and segmentation, to create realistic looking natural images that match the desired scene structure. Our geometrically-consistent image synthesis method is a deep neural network, called Geometry to Image Synthesis (GIS) framework, which retains the advantages of a trainable method, e.g., differentiability and adaptiveness, but, at the same time, makes a step towards the generalizability, control and quality output of modern graphics rendering engines. We utilize the GIS framework to insert vehicles in outdoor driving scenes, as well as to generate novel views of objects from the Linemod dataset. We qualitatively show that our network is able to generalize beyond the training set to novel scene geometries, object shapes and segmentations. Furthermore, we quantitatively show that the GIS framework can be used to synthesize large amounts of training data which proves beneficial for training instance segmentation models.

1 Introduction

Methods for generating natural images from noise or sparse input have gained significant interest in recent years with the developments in Generative Deep Neural Networks. Additionally, Generative Adversarial Networks (GANs)[3], allowed for trainable models that can produce natural-looking images with little or no prior knowledge input just by learning to imitate the distribution in a target image set. While the generated images often consist of realistic looking local patterns, the overall structure of the images can be inconsistent. Using more sparse cues, like edge maps or semantic segmentation[4], introduces some local control over the output but doesn't address the global structure. Further, recent works have addressed the problem of global consistency by generating the image at different scales [2] or using two separate global and local networks [1]

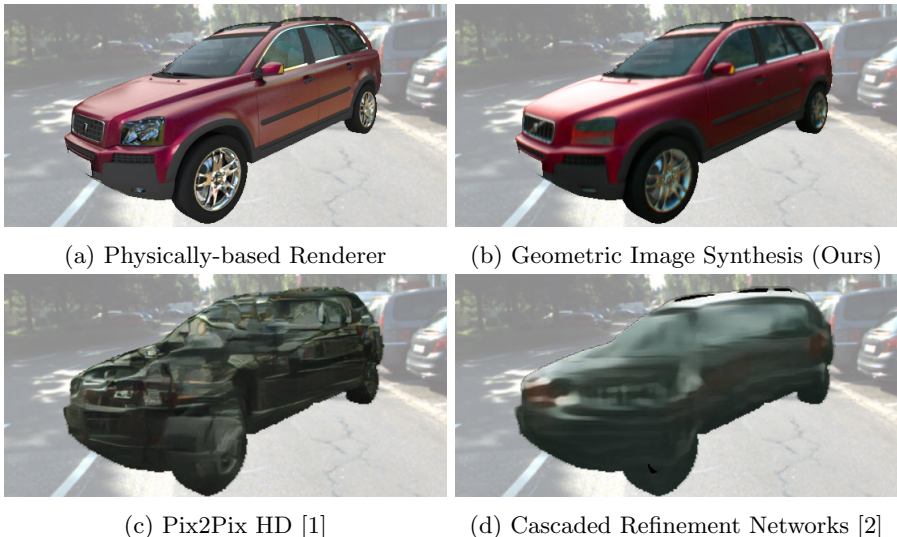


Fig. 1: (a) The result of a state-of-the-art Physically-based Renderer (“Cycle” Renderer). (b) Our GIS framework can realistically synthesize objects in a given image of a road scene, using a deep neural network. (c,d) Results of two other deep neural network based rendering methods [2][1] of road scenes. While in both cases local image patches looks plausible the whole image does not look realistic.

(Fig 1d and 1c). These solutions, nevertheless, address the global 2D structure of the image but not the 3D structure of the scene. This is evident when trying to generate an object in a different pose than those present most commonly in the training dataset (see figure 1d, 1c). While image generation from semantic segmentation can produce visually impressive images, it is not clear whether it can produce new training data for other vision tasks. This could be attributed to two factors: (i) The sparse input makes the image generation problem largely under-constrained leading to inconsistent image structure; (ii) The lack of control parameters over the image generation process (e.g. Pose and color of objects) makes it hard to define the desired attributes of the output image.

On the other hand, generating natural images from known 3D geometry, texture and material properties through rendering engines has been widely used to generate training data for various computer vision tasks. While physically-based rendering engines aim at accurately reproducing the physical properties of light-material interactions, most available rendering engines use a set of carefully designed approximations, in order to reduce the computational complexity and produce results that are visually appealing to humans. Rendered images accurately match the input scene structure but differ in local appearance from real images due the disparity between the real capturing process and the approximations in the software rendering pipeline. Previous works [5] pointed to the performance gap between synthetic and real data when used for training a task

like semantic or instance segmentation. The other limitation of rendering engines is that they require accurate and complete information about the objects and the scene, namely, detailed 3D geometry, texture and material properties, lighting information, environment maps, and so on. This usually requires laborious manual work by experts to set up the 3D scene.

In this work we propose a geometry-aware image generation method that leverages various types of scene information, like geometry and segmentation, to create realistic images which match the desired scene structure and properties. The network is trained with two objectives, the first is a supervised loss where the goal is to learn a mapping from the multi-channel input to an RGB image that matches the input structure. The second is an adversarial loss that learns to compare generated and real images enforcing the generator results to look similar to real data. We explore different input modalities like normals, depth, semantic and material segmentation and compare their usefulness. Using this rich input we are able to show visually a clear improvement over existing state of the art image generation approaches, e.g. [2] [4] [1].

The goal of our approach is not only to generate visually realistic images, but also to explore whether the images generated can be useful for training other networks for various computer vision tasks. The advantage of using a trainable model for image synthesis instead of a software rendering engine is two-fold. Firstly, it can produce realistically looking images from geometry and segmentation while learning, from training data, to implicitly predict the remaining rendering parameters (e.g. material properties and lighting conditions). Secondly, the trainable model has the advantage of producing images that are fine-tuned to specific characteristics of the training dataset. For instance, it can capture the specific noise distribution and color shifts in the data.

In order to demonstrate the abilities of our GIS framework, we perform two types of experiments. In the first, we utilize an augmented reality dataset where synthetic vehicles were realistically rendered into a scene using “Cycles renderer” from Blender [6]. Using the normals, depth and material labels as input and the rendered images as the training target for the supervised loss. In this way, our network is able to generate realistic looking images (see fig. 1) similar to the target data from [6] (see fig. 1a). In fact, we train the GIS network to give 9 output-images, where we observed that each image captures a different lighting condition (e.g. direct sunshine, clear sky, or cloudy), all present in the training data. We have used our framework to produce a new training dataset of 4000 augmented images, where it generalizes to poses and objects it has never seen before. This dataset is used to refine a state-of-the-art instance segmentation network, here Mask R-CNN [7]. This improves the performance of Mask R-CNN [8] over the original augmented data [9]. In the second experiment, we demonstrate how our network can be trained directly using real images only. For that we utilize the Linemod dataset [10] that includes images of several objects and their 3D scanned models in addition to the corresponding 6D pose of the objects in each image. We show that using GIS-Net trained on this data

we are able to generate large amount of training data that helps improve the performance of instance segmentation using Mask R-CNN.

To summarize, our **contributions** are as follows :

- We introduce a trainable deep neural architecture, called the GIS framework, that is able to generate geometry-consistent images from limited input information, like normals and material segmentation, While the remaining aspects of the image, e.g. lighting conditions, are learned from training data.
- We qualitatively show that our framework generalizes to novel scene geometries, objects and segmentation, for both synthetic and real data.
- We quantitatively show that our network can synthesize training data that improves the performance of a state-of-the-art instance segmentation approach, here Mask R-CNN ([7]). To the best of our knowledge, this is the first time that synthesized training data from a Neural Network is used to advance a state-of-the-art instance segmentation approach.

2 Related work

Synthetic Datasets. The success of supervised deep learning models has fueled the demand for large annotated datasets. An alternative to tedious manual annotation is provided by the creation of synthetic content, either via manual 3D scene modeling [11,12] or using some stochastic scene generation process [13,14,15,16]. Mayer et al. [17,18] demonstrate that simple synthetic datasets with “flying 3D things” can be used for training stereo and optical flow models. Ros et al. [11] proposed the SYNTHIA dataset with pixel-level semantic segmentation of urban scenes. In contrast, Gaidon et al. [19] propose “Virtual KITTI”, a synthetic dataset reproducing in detail the popular KITTI dataset [20]. Richter et al. [21,22] and Johnson-Roberson et al. [23] have been the first to demonstrate that content from commercial video games can be accessed for collecting semantic segmentation, optical flow and object detection ground truth by “playing”.

An alternative to synthesizing the entire image content is to render only specific objects into natural images. The simplest approach is to cut object instances from one image and paste them onto random background images [24] using appropriate blending or GAN-based refinement [25]. More variability can be obtained when rendering entire 3D CAD models into the image. Several works consider the augmentation of images with virtual [26,27] or scanned humans [28,29]. In contrast, Alhaija et al. [6,9] consider the problem of augmenting scenes from the KITTI dataset with virtual vehicles for improving object detectors and instance segmentation algorithms. In particular, they have shown that a well performing instance segmentation method, here MNC [30], can be considerably improved by intelligently generating additional training data.

While great progress has been made in rendering photo-realistic scenes, creating the required content and modeling all physical processes (e.g., interaction of light) correctly is a non-trivial and time-consuming task. In contrast to classical

rendering, we propose a generative feed-forward model which maps an intermediate representation of the scene to the desired output. The geometry and appearance cue of this intermediate representation are easily obtained using fast standard OpenGL rendering. Our network learns how to integrate this information with the scene and produces a variety of plausible outputs representing different illumination and environmental conditions.

Conditional Adversarial Learning. Recently, generative adversarial networks (GANs) [3] have proven powerful tools for image generation. Isola et al. [4] formulate the image-to-image translation problem by conditioning GANs on images from another domain and combining an adversarial with a reconstruction loss. They are one of the first to demonstrate image synthesis from semantic labels alone. Yang et al. [31] introduce an additional diversity loss to generate more diverse outputs. Wang et al. [32] propose a multi-scale conditional GAN architecture for generating images of up to 2 Megapixel resolution. Wang et al. [33] use a GAN to synthesize surface normals and another GAN to generate an image from the resulting normal map. GANs’ major advantage is that they don’t require source and target images to be matching but rather enforce the generator to produce images that match the target data distribution. We exploit this by adding an Adversarial loss in our GIS framework such that the generated images are realistic. Besides, we explore a richer set of input modalities compared to just raw images [34,33,35] or semantic segmentations [4,31,1] for generating higher-quality outputs. We demonstrate that our model compares favorably to the High-Resolution Image Synthesis model of Wang et al. [1] (see fig 1b).

Feed-Forward Image Synthesis. Dosovitskiy et al. [36] consider an alternative formulation to GANs using feedforward synthesis with a regression loss. Their work demonstrates that an adversarial loss is not necessary to generate accurate images of 3D models given a model ID and a viewpoint. In the same spirit, Chen et al. [2] consider the problem of synthesizing photographic images conditioned on semantic layouts using a purely feedforward approach. They demonstrate detailed reconstructions at resolutions up to 2 Megapixels, improving considerable upon the results of Isola et al. [4].

Our work also uses a feedforward formulation for the image synthesis problem. Unlike [2], however, our focus is on synthesizing controllable, high quality images. Thus, we consider 3D geometry and segmentation (semantic or material) as input, provided by a simple OpenGL rendering unit. Our technique is not only able to synthesize high quality objects. Instead, we also demonstrate that it allows for seamlessly augmenting natural images with synthetic objects while respecting photometric constraints (e.g., shadows).

3 Method

A general image generation process can be defined as a mapping $\mathcal{F} : \{G, A, E\} \rightarrow I$ from scene description $\{G, A, E\}$ to an RGB image I . The scene description consists of three parts, (i) the geometry parameters G which include the poses and shapes of objects, (ii) the appearance parameters A which describe the

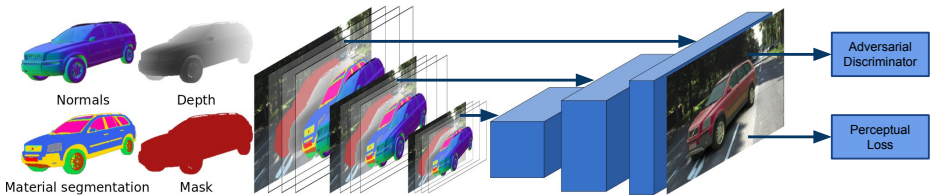


Fig. 2: **Overview of our approach.** We propose Geometric Image Synthesis framework, feed-forward architecture for synthesizing RGB images based on geometric and semantic cues. Here we show the case where a car is augmented onto an empty road. Compared to existing image synthesis approaches, our model benefits from a rich set of input modalities, while learning realistic mappings which generalize to novel geometries and segmentations, and integrate the objects seamlessly into provided image content.

objects’ material, texture and transparency and finally (iii) the environment parameters E which describe global conditions of the scene that affect all objects such as lighting, camera parameters, and the environment. In this work in contrast, our goal is to train a mapping $\mathcal{R} : \{G, S\} \rightarrow I$ that can produce natural images from a given geometry G and material segmentation S only, without the knowledge of exact appearance A or environment parameters E . Similar to semantic segmentation, the material segmentation labels each pixel with a specific material label (e.g. metal, glass etc.) from a pre-defined set of materials without providing any properties or parameters of the material. The task of the network is to learn the unknown parameters from the training data directly and apply them to generate images from new input geometry.

The target image I , which is used for training, can either be a real image of a known scene geometry or a rendered synthetic image obtained through a high-quality software renderer. While learning image generation directly from real images is desirable, it is often difficult to obtain geometry and material labels which are pixel-accurately aligned with real-world images. For this reason, it is possible to exploit synthetically rendered data using a state-of-the-art physically based renderer as supervised target while using an adversarial loss with real images to acquire realistic looking images. Using realistically rendered images also gives us fine-grained control over the data, which we exploit to conduct various experiments for analyzing our model. Additionally, we demonstrate how our method can be trained directly using real images for the supervised loss when an exact 6D pose of the objects in the image is known.

3.1 Input Representation

Geometry plays a major role in defining the appearance of an object in an image since it defines its shape in addition to its shading through interaction with light. Providing the geometry as an input changes the learning objective from learning

to create objects to learning a correlation between geometry and appearance. This makes the network more generalizable to new geometries as we show in later experiments. To use geometry in a deep neural network, it is important to find a compact representation of the object’s 3D surface. While meshes are one of the most common representations for 3D objects, they are problematic in the context of convolutional neural networks due to their irregular 3D structure. Another popular representation of 3D objects are voxels. Voxel-based representations can be handled using 3D convolutions [37] but suffer from two shortcomings: high computational requirements and comparably low resolution.

A common 2D representation of 3D shapes is their depth in the camera view. The advantage of such an image-based geometry representation is that it can be directly processed with a regular 2D convolutional neural network. Nevertheless, the object appearance doesn’t usually depend on its absolute depth, but rather on the small changes in depth between various points with respect to each other in addition to the surface orientation with respect to the light source. This can be better characterized by computing the surface normals in the camera coordinates system at each point of the visible surface.

The main advantage of learning the image generation from geometry compared to rendering is the ability to exploit high-level semantic and context cues to predict the appearance of an object. This allows it to *learn* non-geometric attributes of the object appearance such as material parameters, lighting, environment reflections and texture directly from data. Using semantic [2] or instance segmentation [1] allows the network to learn the appearance of semantically similar objects across multiple examples and scenes. This can be challenging in the case when a semantic class has a large variety in appearance, e.g. cars with different models and colors. We propose in this work to use material segmentation instead. Each pixel in the segmentation input gets a label from a pre-defined set of materials (e.g. metal, plastic, glass etc.). This doesn’t include any material properties or parameters. Rather, it groups parts made of similar materials together allowing the network to learn the material appearance model from multiple objects in different contexts, e.g. different lighting conditions. This results in more generalization power since the material appearance is often independent of the object class, pose or shape. We expect this labeling to be particularly effective when rendering objects that consist of a small number of materials but vary significantly in shape, e.g., cars.

3.2 Network Architecture

We now define our network architecture in detail. As discussed before, our goal is to learn a mapping from an intermediate representation to a natural RGB image using a deep neural network. As input layers to our network, we use the normal map, the depth map object mask and material segmentation of the object which can be easily obtained using OpenGL based rendering. Additionally, by providing the network with a background image I_{bg} , the network can learn to augment synthetic objects realistically into real images e.g. add shadows underneath a synthetic car and blending edges.

Fig. 2 illustrates the input layers to our network. Let N, D, S be the 2D images representing the normal map, depth map and semantic segmentation of the input object respectively. Let M denote the material label where each pixel is represented by a one-hot encoding vector which identifies its material id, see Fig. 2 for an illustration. We are now ready to formally represent the mapping as $\mathcal{R} : \{N, D, O, M, I_{bg}\} \rightarrow I$.

For our generator R , we follow Chen et al. [2] and use a feed-forward coarse-to-fine network architecture for image synthesis. More specifically, we leverage a cascade of convolutional layer modules C starting from a very low resolution input and growing to modules of higher resolutions (Fig. 2). Each convolutional module C_i has an input resolution of $w_i \times h_i$ and produces a feature map F_i of the same size. C_i receives the feature map F_{i-1} from the previous module, upsampled to $w_i \times h_i$, and concatenated with the input downsampled to the same resolution. The following layer, C_{i+1} , operates at twice the resolution of the previous layer ($2w_i \times 2h_i$), and receives the feature maps F_i and the input rescaled to $2w_i \times 2h_i$. Each convolutional module C_i consists of an input layer, intermediate layer and output layer, each of which is followed by a set of convolutions, layer normalization and leaky ReLU nonlinearity. The output layer of the final module is followed by a 1×1 convolution applied to the feature map and normalized to obtain the synthesized image. For the adversarial discriminator D , we use a patch-wise binary classification network operating on patches of size 4×4 similar to the PatchGAN introduced in [4]. This is specially useful when synthesizing objects into real images where the same image would contain both real and synthetic patches. We adopt a fully convolutional network architecture consisting of 5 convolutional layers each followed by a leaky ReLU with a stride of 2 for the first two layers and 1 for the rest. The discriminator’s output is a 2D binary map where each value describes the discriminator classification of a 4×4 patch of the input image into real or synthesized by the generator. To further stabilize the adversarial training, we employ the simple discriminator gradient regularization method proposed in [38]

3.3 Training

We train the generator R in our GIS framework to produce synthesized images I_s that resemble the target images I_t obtained using the “Cycles” render engines while at the same time being close in appearance to real images in order to “confuse” the adversarial discriminator. Effectively, the task of the network is to learn the process of generating images, directly from the target images, given $\{N, D, O, M, I_{bg}\}$ without information such as lighting, environment map or material properties. Those properties are estimated by the network during training and combined with the geometry and segmentation input to produce a high quality image. To achieve this, we choose to compute the perceptual loss (feature matching loss) as proposed in [39] between the generated object and the target object. The goal of the perceptual loss is to match the feature activations produced by I_t and I_s at various convolutional levels through a perception network, e.g., VGG. This helps the network to learn fine-grained image patterns

while also preserving the global object structure. We use VGG pre-trained on ImageNet as our perceptual network. Let us denote this network by V , and let V_l denote a layer of this network. The loss between I_t and I_s is given by $\mathcal{L}^P = \sum_l S_l \lambda_l \|V_l(I_t) - V_l(I_s)\|_1$ where $V_l(\cdot)$ denotes the feature activations of VGG at layer l and S_l is the binary mask of the object rescaled to the size of V_l . The GIS framework can also learn to synthesize objects on top of real images. In this case, our goal is to create augmented images by learning not just the target object appearance but also its interaction with the environment in the real image, including shadows, reflection and blending at the object's edges. Towards this goal, we add the background image to the input and train the network using an ℓ_1 loss for the background areas outside the object mask: $\mathcal{L}^B = (1 - S)\|I_t - I_s\|_1$. The network is trained by minimizing the overall loss $\mathcal{L} = w \mathcal{L}^P + (1 - w) \mathcal{L}^B$ where w is a weight inversely proportional to the number of pixels of the synthesized object(s).

3.4 Diversity

Synthesizing images from geometry and segmentation alone is an ill-posed problem. That is, for a specific set of inputs, there are infinitely plausible outputs due to different possible lighting conditions, object colors etc. Thus, we task our network to produce K diverse outputs from the last layer using multiple-choice learning [40,2]. More specifically, we compute the loss for each of these outputs, but only back propagate the gradient of the best configuration for the foreground prediction, while averaging the background predictions as none of them should deviate from the input: $\mathcal{L} = w \min_k \mathcal{L}_k^P + (1 - w) \frac{1}{K} \sum_k \mathcal{L}_k^B$. Note that only the foreground object with the smallest loss is taken into account, thus the min operator effectively acts as a multiple choice switch. This encourages the network to output a diverse set of images to spread its bets over the domain of possible images that could be produced from the current input. In all our experiments we use $K = 9$ as the number of diverse outputs, see 5 for an illustration.

4 Experiments

To demonstrate the ability of our GIS network to synthesize realistic images, we perform a set of experiments which assess the quality and generalization capacity of the method. We mainly focus on two scenarios, outdoor driving and indoor objects. Realistically synthesizing augmented objects like cars or obstacles into real-world scenes is an important feature for expanding manually annotated training datasets. Moreover, using a learnable image synthesis method like GIS can greatly increase the effectiveness of the generated data for training, since it can be specifically tuned to the task at hand. In the case of indoor objects, a learned network can be used to synthesize novel views of objects to provide extensive training data for various tasks. Our goal is to show that the GIS framework produces images on par with a state-of-the-art software rendering engine for generating training data.

4.1 Augmentation of Kitti-360 dataset

Introduced by Alhaija et al. [6], the augmented KITTI-360 dataset features 4000 augmented images obtained from 200 real-world images through careful rendering up to 5 high-quality 3D car models into each image using classical rendering techniques. The set of 28 car models have been manually created and placed to ensure high realism. Rendering was performed using the physically-based Cycles Renderer in Blender and a manually designed post-processing pipeline was applied to increase the realism of the output in terms of its low-level statistics. Additional scene information like 360 degree environment maps and camera calibration has been used to ensure realistic reflections and good integration with the real image.

Our goal for this experiment is train our GIS network using the available 3D car models and real images from KITTI-360 to create the input data, and the pre-rendered images from the augmented KITTI-360 as target images. We then leverage the trained GIS model to create a new dataset of 4000 augmented images with new car poses and combinations. Mixing the real images with rendered cars presents an additional challenge since interactions between the inserted objects and the real background, e.g. shadows and transparencies, have to be taken into account. To deal with this, we input the background image to the network in addition to the geometry and the segmentation information of the car model.

To train the GIS network for this task, we use normals, depth, material segmentation and semantic masks of the augmented cars provided by the augmented KITTI-360 pipeline. The material labels include 16 materials of different properties (e.g. plastic, chrome glass etc.) in addition to 15 car paint materials which differ only in color. We use the corresponding RGB images from augmented KITTI-360 as target images for training the parameters of GIS network, tasking the network to learn the process of synthesizing images realistically and blend them into the surrounding environment by appropriately adding reflections, shadows and transparencies, amongst others.

During inference, we obtain a new set of car model positions and orientations following the procedure in [6]. We render the mask, depth, normals and material labels from the camera viewpoint and use them as input to our GIS network. Note that during the inference phase, we do not require a sophisticated rendering pipeline, like “Cycles” renderer, since normals, depth maps and segmentations can be obtained directly using a simple OpenGL based renderer.

Qualitative evaluation: Fig. 3 shows augmented images produced by our GIS framework when trained on the augmented KITTI-360 dataset. Note that the synthesized cars exhibit realistic appearance properties like shading, shadows, reflections and specularities, despite the fact that this information is not provided to our model. The material labeling of the cars allows the model to tune the synthesis process to each material. Importantly, note that the material label is just a semantic label of material and does not contain any information with respect to the physical properties of the material. Interestingly, our model is able to learn the transparency property of the material with label “glass” from data, without providing any alpha channel or explicitly modeling transparency.



Fig. 3: Images from KITTI-360 dataset augmented with cars synthesized using our GIS framework (Real image without augmentation in upper left corner).

Dataset	IoU 50%	AP
Real KITTI-360	58.80%	31.92%
Augmented KITTI-360	66.68%	37.88%
GIS (ours)	67.74%	38.69%

Table 1: Accuracy of Mask R-CNN when trained with real, augmented or GIS generated images.

Additionally, the model is able to replicate camera effects such as blur and chromatic aberrations, which are present in the augmented KITTI-360 dataset.

Quantitative evaluation: To verify the effectiveness of data produced by our model, we train the state-of-the-art Mask R-CNN model [7] for car instance segmentation using the images produced by our network. Alongside, we also train the same model with images from the original Cycles Rendering pipeline from [6], and a baseline model using the unaugmented real images from KITTI-360. We evaluate all models on the KITTI 2015 training set. The results are presented in Table 1. We observe that the model trained on images synthesized by the GIS network significantly outperforms the one trained only on real data, and marginally outperform the highly-tuned data from [6]. This clearly indicates that our model does not only learn to imitate the training data, but also the adversarial loss can contribute in make the result appearance more realistic and, therefore, more effective in training. This indicates that the GIS can be reliably leveraged to train instance segmentation models for improving their performance.

4.2 Generalization and Ablation Study

A key feature of our GIS framework is that it learns a mapping from any geometry to natural images and is not limited to a specific set of objects or shapes. In the following sections, we present an extensive experimental study to demonstrate that our model learns a generic image formation function and does not overfit or limit itself only to objects of certain geometry and material.

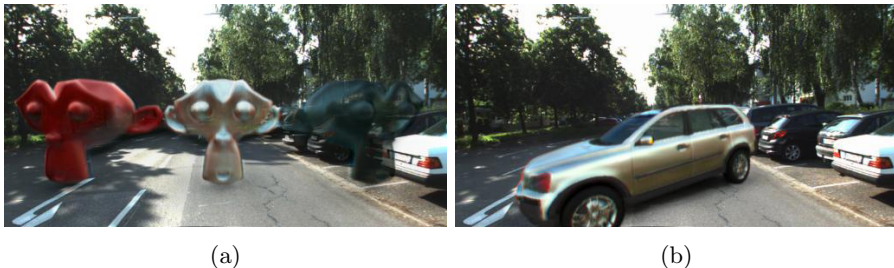


Fig. 4: (a) GIS output for a monkey model with material labels: car paint, chrome and glass. (b) GIS output for a car with material label chrome.



Fig. 5: Three diverse outputs obtained from GIS on the KITTI-360 dataset. Note how the renderings vary in lighting conditions and reflections.

To show generalization ability, we present the network with two tasks: (i) synthesize seen objects with material combinations never seen before, and (ii) synthesize learned materials on new, unseen, geometries. In Fig. 4a we show the results of our model applied to the “Monkey” model from blender with different material labels applied to it. Our results clearly demonstrate that the material properties have been learned by the network independently from the geometry. In Fig. 4b we replace the car paint with the chrome material previously seen in the training data only on the car wheel rims. The resulting image looks realistic, demonstrating that the material properties learned from one part of the model can be transferred to other, geometrically different, parts of the model by simply changing the material label. Using the diversity loss, described in Section 3.4, our GIS model produces 9 different possible images from the same input. The results in Fig. 5 show how the network can learn different lighting conditions (direct light, cloudy etc.) without providing it with any explicit information about lighting.

Ablation study: To better understand the importance of different input modalities, we perform an ablation study where we train a GIS model from scratch using all inputs excluding one at a time. We qualitatively compare the results in Fig. 6. When normals are not used for training the model, the output images become smooth and lack fine geometric details. Excluding depth maps from the input, on the other hand, leads to no noticeable difference. We hypothesize that this is due to the fact that most of the shading of the object can be modeled based on the local geometry cues that are expressed well in the normals, but little difference in appearance relates to the absolute depth of an object. In contrast, removing the material segmentation results in blurry images.

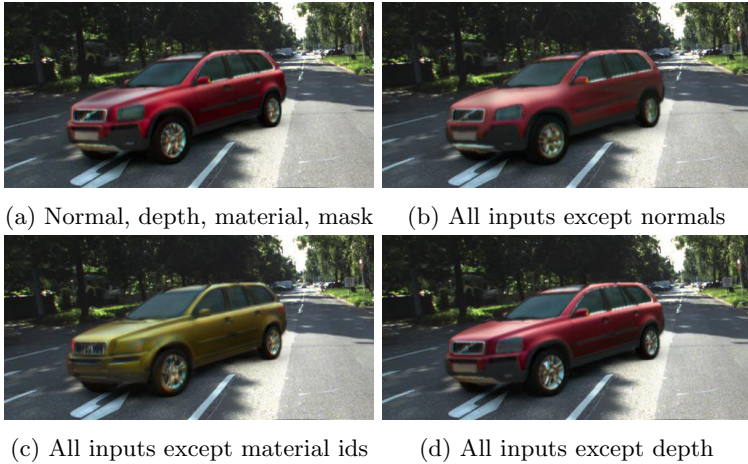


Fig. 6: Output of GIS using various types of inputs. Note that GIS with all four inputs, or all inputs except depth, synthesizes realistic images.

4.3 Novel view synthesis on Linemod dataset

The Linemod dataset was introduced by Hinterstoisser et al. [10] for evaluation of 6D object pose estimation algorithms. The training set of this data contains only synthetic images of various objects rendered using their 3D CAD models in various poses. The evaluation set contains images with multiple known objects, each of them annotated with a 6 degree-of-freedom pose with respect to a 3D CAD model of that object. Additionally, we annotated various materials present on each CAD model with a material class. Hence, using the 6D pose and the CAD model, we can obtain the normal map (N), material segmentation (S) and depth map (D) for an object.

The objective of this experiment is to use real images as target training data for our network with the corresponding geometric information (N, S, D) as inputs. Unlike the KITTI-360 dataset, where the target data is acquired using a manually designed post-processing framework, the availability of real images as target data in this case allows the network to model real world images and their noise statistics. We use the objects *Ape*, *Can*, *Cat*, *Duck*, *Eggbox*, *Holepuncher* each with 1200 images with pose annotations of which we use 600 images of each object as training data. For each image and each object, we obtain the normal map, depth and material segmentation. While these are used as the input to our network, the corresponding RGB images are used as the target training data.

Once our network is trained, we can further use the 3D models to obtain (N, S, D) for a large number of new 6D poses. Using this information as input, we can synthesize RGB images of these objects in novel viewpoints which the network has not seen during training phase.

To demonstrate the efficacy of the images produced by our network, we train the Mask R-CNN instance segmentation framework. To this end, we create two kinds of datasets to train the Mask R-CNN. First, we crop the pure synthetic im-



Fig. 7: Top row contains pure synthetic images. Middle row contains real images in similar poses. Bottom row contains images synthesized by our network. The middle row images are not seen during training phase. GIS is still able to synthesize novel views of objects realistically.

ages provided in Linemod training data and augment them at random locations on NYU dataset scenes as background. Alongside, we use the object images from our network to also augment them on NYU dataset scenes. To evaluate the performance of Mask R-CNN, we test it on Linemod-Occluded dataset, proposed by Michel et al. [41] where a sequence of real images picked from Linemod dataset are annotated with instance masks of the objects present in them (note that we do not use this data while training our GTI network). We observe that the Mask R-CNN trained with only synthetic augmented images performs at an average accuracy of 21.1% for all objects. On the other hand, the Mask R-CNN trained with our synthesized images performs at an average accuracy of 68.4%. This clearly indicates that the images synthesized by our network are highly realistic and are useful for training other deep networks which cannot be achieved with pure synthetic data. This can also be seen in 7 where the first row contains pure synthetic images, the second row contains real images in similar poses and the third row contains images produced by our network. It can be clearly seen that our network synthesizes images which are very close to real images and can also be used to train other tasks.

5 Conclusion

In this work, we have proposed GIS, a deep neural network which is able to learn to synthesize realistic objects by leveraging semantic and geometric scene information. Through various experiments we have demonstrated the generalization performance of our GIS framework with respect to varying geometry, semantics and materials. Further, we have provided empirical evidence that the images synthesized by GIS are realistic enough to train the state-of-the-art instance segmentation method Mask R-CNN, and improve its accuracy on car instance segmentation with respect to a baseline model trained on non-augmented images from the same dataset. We believe that our approach opens new avenues towards ultimately reaching the goal of photo-realistic image synthesis using deep neural networks.

References

1. Wang, T.C., Liu, M.Y., Zhu, J.Y., Tao, A., Kautz, J., Catanzaro, B.: High-resolution image synthesis and semantic manipulation with conditional gans. In: CVPR. (2018)
2. Chen, Q., Koltun, V.: Photographic image synthesis with cascaded refinement networks. In: ICCV. (2017)
3. Goodfellow, I.J., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A.C., Bengio, Y.: Generative adversarial nets. In: NIPS. (2014)
4. Isola, P., Zhu, J., Zhou, T., Efros, A.A.: Image-to-image translation with conditional adversarial networks. In: CVPR. (2017)
5. Tremblay, J., Prakash, A., Acuna, D., Brophy, M., Jampani, V., Anil, C., To, T., Cameracci, E., Boochoon, S., Birchfield, S.: Training deep networks with synthetic data: Bridging the reality gap by domain randomization. arXiv preprint arXiv:1804.06516 (2018)
6. Alhaija, H.A., Mustikovela, S.K., Mescheder, L., Geiger, A., Rother, C.: Augmented reality meets deep learning for car instance segmentation in urban scenes. In: BMVC. (2017)
7. He, K., Gkioxari, G., Dollár, P., Girshick, R.B.: Mask R-CNN. arXiv.org **1703.06870** (2017)
8. He, K., Gkioxari, G., Dollr, P., Girshick, R.: Mask r-cnn. In: ICCV. (2017) 2980–2988
9. Alhaija, H., Mustikovela, S., Mescheder, L., Geiger, A., Rother, C.: Augmented reality meets computer vision: Efficient data generation for urban driving scenes. IJCV (2018)
10. Hinterstoisser, S., Lepetit, V., Ilic, S., Holzer, S., Bradski, G.R., Konolige, K., Navab, N.: Model based training, detection and pose estimation of texture-less 3d objects in heavily cluttered scenes. In: ACCV. (2012)
11. Ros, G., Sellart, L., Materzynska, J., Vazquez, D., Lopez, A.: The synthia dataset: A large collection of synthetic images for semantic segmentation of urban scenes. In: CVPR. (2016)
12. Zhang, Y., Song, S., Yumer, E., Savva, M., Lee, J., Jin, H., Funkhouser, T.A.: Physically-based rendering for indoor scene understanding using convolutional neural networks. In: CVPR. (2017)
13. McCormac, J., Handa, A., Leutenegger, S., Davison, A.J.: Scenenet RGB-D: can 5m synthetic images beat generic imagenet pre-training on indoor segmentation? In: ICCV. (2017)
14. de Souza, C.R., Gaidon, A., Cabon, Y., Peña, A.M.L.: Procedural generation of videos to train deep action recognition networks. arXiv.org **1612.00881** (2016)
15. Tsirikoglou, A., Kronander, J., Wrenninge, M., Unger, J.: Procedural modeling and physically based rendering for synthetic data generation in automotive applications. arXiv.org **1710.06270** (2017)
16. Veeravasaru, V.S.R., Rothkopf, C.A., Ramesh, V.: Model-driven simulations for deep convolutional neural networks. arXiv.org **1605.09582** (2016)
17. Mayer, N., Ilg, E., Fischer, P., Hazirbas, C., Cremers, D., Dosovitskiy, A., Brox, T.: What makes good synthetic training data for learning disparity and optical flow estimation? arXiv.org **1801.06397** (2018)
18. Mayer, N., Ilg, E., Haeusser, P., Fischer, P., Cremers, D., Dosovitskiy, A., Brox, T.: A large dataset to train convolutional networks for disparity, optical flow, and scene flow estimation. In: CVPR. (2016)

19. Gaidon, A., Wang, Q., Cabon, Y., Vig, E.: Virtual worlds as proxy for multi-object tracking analysis. In: CVPR. (2016)
20. Geiger, A., Lenz, P., Urtasun, R.: Are we ready for autonomous driving? The KITTI vision benchmark suite. In: CVPR. (2012)
21. Richter, S.R., Vineet, V., Roth, S., Koltun, V.: Playing for data: Ground truth from computer games. In: ECCV. (2016)
22. Richter, S.R., Hayder, Z., Koltun, V.: Playing for benchmarks. In: ICCV. (2017)
23. Johnson-Roberson, M., Barto, C., Mehta, R., Sridhar, S.N., Rosaen, K., Vasudevan, R.: Driving in the matrix: Can virtual worlds replace human-generated annotations for real world tasks? In: ICRA. (2017)
24. Dwibedi, D., Misra, I., Hebert, M.: Cut, paste and learn: Surprisingly easy synthesis for instance detection. In: ICCV. (2017)
25. Xu, W., Li, Y., Lu, C.: Generating instance segmentation annotation by geometry-guided GAN. arXiv.org **1801.08839** (2018)
26. Cheung, E., Wong, T.K., Bera, A., Manocha, D.: STD-PD: generating synthetic training data for pedestrian detection in unannotated videos. arXiv.org **1707.09100** (2017)
27. Hattori, H., Boddeti, V.N., Kitani, K.M., Kanade, T.: Learning scene-specific pedestrian detectors without real data. In: CVPR. (2015)
28. Chen, W., Wang, H., Li, Y., Su, H., Wang, Z., Tu, C., Lischinski, D., Cohen-Or, D., Chen, B.: Synthesizing training images for boosting human 3d pose estimation. In: 3DV. (2016)
29. Varol, G., Romero, J., Martin, X., Mahmood, N., Black, M.J., Laptev, I., Schmid, C.: Learning from synthetic humans. In: CVPR. (2017)
30. Dai, J., He, K., Sun, J.: Instance-aware semantic segmentation via multi-task network cascades. In: CVPR. (2016)
31. Yang, Z., Liu, H., Cai, D.: On the diversity of realistic image synthesis. arXiv.org **1712.07329** (2017)
32. Wang, T., Liu, M., Zhu, J., Tao, A., Kautz, J., Catanzaro, B.: High-resolution image synthesis and semantic manipulation with conditional gans. arXiv.org **1711.11585** (2017)
33. Wang, X., Gupta, A.: Generative image modeling using style and structure adversarial networks. In: ECCV. (2016)
34. Denton, E.L., Chintala, S., Szlam, A., Fergus, R.: Deep generative image models using a laplacian pyramid of adversarial networks. In: NIPS. (2015)
35. Radford, A., Metz, L., Chintala, S.: Unsupervised representation learning with deep convolutional generative adversarial networks. arXiv.org **1511.06434** (2015)
36. Dosovitskiy, A., Springenberg, J.T., Tatarchenko, M., Brox, T.: Learning to generate chairs, tables and cars with convolutional networks. PAMI **39** (2017) 692–705
37. Riegler, G., Ulusoy, A.O., Geiger, A.: Octnet: Learning deep 3d representations at high resolutions. In: CVPR. (2017)
38. Roth, K., Lucchi, A., Nowozin, S., Hofmann, T.: Stabilizing training of generative adversarial networks through regularization. In: NIPS. (2017) 2018–2028
39. Gatys, L.A., Ecker, A.S., Bethge, M.: A neural algorithm of artistic style. arXiv.org **1508.06576** (2015)
40. Guzmán-Rivera, A., Batra, D., Kohli, P.: Multiple choice learning: Learning to produce multiple structured outputs. In: NIPS. (2012)
41. Michel, F., Kirillov, A., Brachmann, E., Krull, A., Gumhold, S., Savchynskyy, B., Rother, C.: Global hypothesis generation for 6d object pose estimation. CoRR abs/1612.02287 (2016)