

Review

When Autonomous Systems Meet Accuracy and Transferability through AI: A Survey

Chongzhen Zhang,^{1,5} Jianrui Wang,^{1,5} Gary G. Yen,² Chaoqiang Zhao,¹ Qiyu Sun,¹ Yang Tang,^{1,*} Feng Qian,¹ and Jürgen Kurths^{3,4}

¹Key Laboratory of Advanced Control and Optimization for Chemical Process, Ministry of Education, East China University of Science and Technology, Shanghai 200237, China

²School of Electrical and Computer Engineering, Oklahoma State University, Stillwater, OK 74075, USA

³Potsdam Institute for Climate Impact Research, 14473 Potsdam, Germany

⁴Institute of Physics, Humboldt University of Berlin, 12489 Berlin, Germany

⁵These authors contributed equally

*Correspondence: yangtang@ecust.edu.cn or tangtany@gmail.com

<https://doi.org/10.1016/j.patter.2020.100050>

THE BIGGER PICTURE Accuracy and transferability are critical to the perception and decision-making tasks of autonomous systems. The focus of several learning-based perception and decision-making methods has gradually evolved from accuracy to transferability. This survey summarizes the perception and decision-making tasks of autonomous systems from the perspectives of accuracy and transferability. We introduce transfer learning and some preliminaries of adversarial learning, reinforcement learning, and meta-learning. Then, we review several perception and decision tasks of autonomous systems from the perspectives of accuracy or transferability or both. Last but not least, we discuss several challenges and future works for using adversarial learning, reinforcement learning, and meta-learning in autonomous systems.

With widespread applications of artificial intelligence (AI), the capabilities of the perception, understanding, decision-making, and control for autonomous systems have improved significantly in recent years. When autonomous systems consider the performance of accuracy and transferability, several AI methods, such as adversarial learning, reinforcement learning (RL), and meta-learning, show their powerful performance. Here, we review the learning-based approaches in autonomous systems from the perspectives of accuracy and transferability. Accuracy means that a well-trained model shows good results during the testing phase, in which the testing set shares a same task or a data distribution with the training set. Transferability means that when a well-trained model is transferred to other testing domains, the accuracy is still good. Firstly, we introduce some basic concepts of transfer learning and then present some preliminaries of adversarial learning, RL, and meta-learning. Secondly, we focus on reviewing the accuracy or transferability or both of these approaches to show the advantages of adversarial learning, such as generative adversarial networks, in typical computer vision tasks in autonomous systems, including image style transfer, image super-resolution, image deblurring/dehazing/rain removal, semantic segmentation, depth estimation, pedestrian detection, and person re-identification. We furthermore review the performance of RL and meta-learning from the aspects of accuracy or transferability or both of them in autonomous systems, involving pedestrian tracking, robot navigation, and robotic manipulation. Finally, we discuss several challenges and future topics for the use of adversarial learning, RL, and meta-learning in autonomous systems.

Introduction

Artificial intelligence (AI) has been widely used in art, government, healthcare, games, and economics due to its powerful learning ability. Especially after the representative AI algorithm AlphaGo defeated the world champion in Go games,¹ people have been paying more attention to AI. Understanding the behavior of AI agents is very important in promoting its technology.² With the rise of deep learning (DL) algorithms, the upgrading of hardware, and the availability of big data, AI technology has been making huge progress in recent years.³ Autonomous systems powered by AI, including unmanned vehicles, robotic

manipulators, and drones have been widely used in various industries and daily lives, such as intelligent transportation,⁴ intelligent logistics,⁵ and service robots.⁶ Due to the limitations of current computer perception and decision-making technologies in terms of accuracy and transferability, autonomous systems still have much room for improvement in complex and intelligent tasks via technological development. Due to the ability of DL to capture high-dimensional data features,³ DL-based algorithms are widely used in the perception and decision-making tasks of autonomous systems. There are a number of typical tasks related to perception and decision-making for autonomous



systems, such as image super-resolution (SR),^{7,8} image deblurring/dehazing/rain removal,^{9–11} semantic segmentation,^{12,13} depth estimation,^{14,15} pedestrian detection,¹⁶ person re-identification (re-ID),¹⁷ pedestrian tracking,¹⁸ robot navigation,^{19,20} and robotic manipulation.^{21,22} However, most DL-based models have good accuracy and poor transferability, i.e., they are usually effective in the testing dataset with the same data distribution or task. When a well-trained model is transferred to other datasets or real-world tasks, the accuracy usually declines drastically, which means that the transferability is poor; thus, the transferability has to be taken into account for practical applications.²³ This issue results in the fact that the current vision perception and decision-making methods cannot be used directly in actual autonomous systems. Transfer learning improves the transferability of models between different domains, i.e., a well-trained model can achieve good accuracy when applied to other testing domains.

Recently, since adversarial learning, such as generative adversarial networks (GANs), has shown promising results in image generation, a number of GANs-based methods have been proposed and have achieved breakthroughs in the aforementioned computer vision tasks.^{24–27} In the field of AI, GANs have become increasingly important due to their powerful generation and domain adaptation capabilities.²⁸ GANs have attracted increasing attention since they were proposed by Goodfellow et al.²⁹ in 2014. GAN is a generative model that introduces adversarial learning between the generator and the discriminator, in which the generator creates data to deceive the discriminator while the discriminator distinguishes whether its input comes from real data or generated ones. The generator and discriminator are iteratively optimized in the game, and finally reach the Nash equilibrium.³⁰ In particular, when considering a well-trained model for different datasets or real scenes, GANs can be used for domain-transfer tasks by virtue of their ability to capture high-frequency features to generate sharp images.³¹ Although some learning-based models mainly focus on the aspect of accuracy,^{7,12,14} GANs have demonstrated satisfactory results for various complex image fields in autonomous systems and other related fields, such as text-to-image generation,^{32,33} image style transfer,^{24,34} SR,²⁶ image deblurring,²⁷ image rain removal,^{35,36} object detection,^{37,38} semantic segmentation,^{24,39,40} pedestrian detection,⁴¹ person re-ID,⁴² and video generation.⁴³

Meanwhile, as a powerful tool for decision-making and control, reinforcement learning (RL) has been extensively studied in recent years because it is suitable for decision-making tasks in complex environments.^{44,45} However, when the input data are high-dimensional such as images, sounds, and videos, it is difficult to solve the problem only with RL. With the help of deep neural networks (DNNs), deep RL (DRL), which combines the high-dimensional perceptual ability of DL with the decision-making ability of RL, has achieved promising results recently in various fields of application, such as obstacle avoidance,^{46,47} robot navigation,^{48,49} robotic manipulation,^{50,51} video target tracking,^{18,52} game playing,^{53,54} and drug testing.^{55,56} However, DRL tends to require a large number of trials and needs to specify a reward function to define a certain task.⁵⁷ The former is time-consuming and the latter is significantly difficult when training from scratch. To tackle these problems, the idea of

“learn to learn,” called meta-learning, has emerged.⁵⁸ Compared with DRL, meta-learning makes the learning methods more transferable and efficient by utilizing previous experience to guide the learning of new tasks across domains. Therefore, meta-learning methods perform well especially in environments lacking data, such as image recognition,⁵⁹ classification,⁶⁰ robot navigation,⁶¹ and robotic arm control.⁶²

With the development of DL, learning-based perception and decision-making algorithms for autonomous systems have become a hot research topic. Among the reviews of autonomous systems, Tang et al.⁶³ introduced the applications of learning-based methods in perception and decision-making for autonomous systems. Gui et al.²⁸ gave a detailed overview of various GANs methods from the perspectives of algorithms, theories, and applications. Arulkumaran et al.⁶⁴ detailed the core algorithms of DRL and the advantages of RL for visual understanding tasks. Unlike previous surveys, we focus on reviewing learning-based approaches in the perception and decision-making tasks of autonomous systems from the perspectives of accuracy or transferability, or both.

The organization of this review is arranged as follows. The next section introduces transfer learning and one of its related machine-learning techniques, domain adaptation, and presents the basic concepts of adversarial learning, RL, and meta-learning. Following this, we survey some recent developments by exploring various learning-based approaches in autonomous systems, taking into account accuracy or transferability or both of these concepts. We then summarize some trends and challenges for autonomous systems, followed by our conclusions. The abbreviations used in this review are listed in [Table 1](#).

Preliminaries

Learning-based methods are used in various perception and decision-making tasks of autonomous systems, such as image style transfer, image SR, image deblurring/dehazing/rain removal, semantic segmentation, depth estimation, pedestrian detection, person re-ID, pedestrian tracking, robot navigation, and robotic manipulation. However, most traditional learning-based methods usually achieve good accuracy on the testing set with the same distribution or the same task. In recent years, with the research on the transferability of models, several typical learning-based methods have been widely used, such as adversarial learning and meta-learning.

Overview of the Section

Focusing on the transferability of models, transfer learning is proposed and first introduced in this section, which aims to make a well-trained model have good transferability, i.e., the well-trained model can be transferred to other testing sets and still have a good accuracy. We then introduce several typical learning-based methods concentrating on improving the accuracy or transferability, or both, including adversarial learning, RL, and meta-learning. In the perception tasks of autonomous systems, adversarial learning, such as GANs, has capabilities of good accuracy or transferability or both. In the decision-making tasks of autonomous systems, RL and meta-learning are often used to improve the accuracy or transferability of the system.

Transfer Learning

Transfer learning is a research topic aiming to investigate the improvement of learners from one target domain trained with

Table 1. Summary of Abbreviations in This Review

Abbreviation	Full Name
AC	actor-critic
AI	artificial intelligence
cGANs	conditional generative adversarial networks
CNNs	convolutional neural networks
CycleGAN	cycle-consistent adversarial network
DL	deep learning
DNNs	deep neural networks
DQN	deep Q network
DRL	deep reinforcement learning
GANs	generative adversarial networks
GAIL	generative adversarial imitation learning
HR	high-resolution
IRL	inverse reinforcement learning
LR	low-resolution
LSTM	long short-term memory
MAML	model-agnostic meta-learning
re-ID	re-identification
RL	reinforcement learning
SR	super-resolution
TCN	temporal convolution network

more easily obtained data from source domains.⁶⁵ In other words, the domains, tasks, and distributions used in training and testing could be different. Therefore, transfer learning saves a great deal of time and cost in labeling data when encountering various scenarios of machine-learning applications. According to different situations between domains, source tasks, and target tasks, transfer learning can be categorized into three sub-settings: inductive transfer learning, transductive transfer learning, and unsupervised transfer learning.⁶⁶ The definitions and differences between these transfer learning settings are presented in detail in Table 2.

Domain Adaptation. There are many machine-learning techniques that are connected to transfer learning,⁶⁶ for example, domain adaptation,⁶⁷ related to transductive transfer learning, and multi-task learning⁶⁸ and self-taught learning,⁶⁹ related to inductive transfer learning. Here, we focus on domain adaptation, whereby the source and target domains share the same feature spaces while the feature distributions are different but related. The difference between domain adaptation and transductive transfer learning is that domain adaptation leverages labeled data in the source domain to learn a classifier for the target domain, where the target domain is either fully unlabeled (unsupervised domain adaptation) or has few labeled samples (semi-supervised domain adaptation).⁷⁰ Domain adaptation is

promising for the transferability of perception tasks of autonomous systems because it is efficient to reduce the domain shift among different datasets arising from synthetic and real images,⁷¹ different weather conditions,⁷² different lighting conditions,⁷³ or different seasons,⁷⁴ among others. Domain adaptation for visual applications includes shallow and deep methods.³¹ There is some research studying shallow domain-adaptive methods, which mainly include homogeneous domain adaptation and heterogeneous domain adaptation, according to whether the source data and target data have the same representation.^{67,75,76} Readers who want to learn more about shallow domain adaptation methods are referred to the studies by Csurka³¹ and Patel et al.,⁷⁷ and the references therein. In this review, we mainly focus on deep domain adaptation methods, including traditional DL^{67,78,79} and adversarial learning.^{80–82}

Adversarial Learning

Early adversarial learning modeled the learner and the adversary as a competitive two-player game.⁸³ Subsequently, adversarial learning games have expanded into different forms, such as a Bayesian game,⁸⁴ a sequential game,⁸⁵ a bilevel optimization problem,⁸⁶ and so forth.⁸⁷ With the popularity of DNNs, Goodfellow et al.²⁹ used adversarial learning to generate tasks, i.e., GANs. This model is widely used in various fields of autonomous systems.

Generative Adversarial Networks. As a powerful learning-based method for computer vision tasks, adversarial learning not only improves the accuracy but also helps improve the transferability of the model by reducing the differences between the training and testing domain distributions.⁸⁰ GANs are architectures that use adversarial learning methods for generative tasks.²⁸ The framework includes two models, a generator G and a discriminator D , as shown in Figure 1. G captures the prior noise distribution $p_z(z)$ to generate fake data $G(z)$, and D outputs a single scalar to characterize whether the sample comes from training data x or generated data $G(z)$. G and D play against each other, promote each other, and finally reach the Nash equilibrium.³⁰ G and D focus on a two-player minimax game with the value function $V(G, D)$:

$$\min_G \max_D V(G, D) = \mathbb{E}_{x \sim p_{\text{data}}(x)} [\log D(x)] + \mathbb{E}_{z \sim p_z(z)} [\log(1 - D(G(z)))], \quad (\text{Equation 1})$$

where $V(G, D)$ is a binary cross-entropy function, which aims to let D classify real or fake samples. In Equation 1, D tries to maximize its output, G tries to minimize its output, and the game ends at a saddle point.³⁰

Conditional Generative Adversarial Networks. In the original generative model, since the prior comes from the noise distribution $p_z(z)$, the mode of the generated data cannot be controlled.³⁰ Mirza and Osindero⁸⁹ then proposed conditional GANs (cGANs), in which some extra information y is fed to the generator and discriminator in the model such that the data generation process can be guided, as shown in Figure 1. Note that y can be class labels or any other kind of auxiliary information. Compared with Equation 1, the objective function of cGANs is as follows:

Table 2. Definitions and Differences between Three Transfer Learning Settings

Transfer Learning Settings	Source and Target Domains	Source and Target Tasks	Source Domain Labels	Target Domain Labels
Inductive transfer learning	the same/different but related	different but related	available/unavailable	available
Transductive transfer learning	different but related	the same	available	unavailable
Unsupervised transfer learning	the same/different but related	different but related	unavailable	unavailable

Copyright 2009, IEEE. Reprinted, with permission, from Pan and Yang.⁶⁶

$$\min_G \max_D V(G, D) = \mathbb{E}_{x \sim p_{\text{data}}(x)} [\log D(x|y)] + \mathbb{E}_{z \sim p_z(z)} [\log(1 - D(G(z|y)))] \quad (\text{Equation 2})$$

Cycle-Consistent Adversarial Networks. Unlike models tailored for specific tasks, such as GANs and cGANs, cycle-consistent adversarial networks (CycleGANs) use a unified framework for various image tasks, which make the framework simple and effective.²⁴ Zhu et al.²⁴ proposed CycleGAN to learn image translation between the source domain X and the target domain Y with unpaired training examples $\{x_i\}_{i=1}^N \in X$ and $\{y_j\}_{j=1}^M \in Y$, in which N, M represent the total number of samples in the source and target domains, as shown in Figure 1. The framework includes two generators $G: X \rightarrow Y$ and $F: Y \rightarrow X$, and two discriminators D_X and D_Y , where D_X distinguishes between images x and translated images $F(y)$; similarly, D_Y distinguishes between images y and translated images $G(x)$. The output of the mapping G is $\hat{y} = G(x)$, and the output of the mapping F is $\hat{x} = F(y)$. Zhu et al. express the adversarial loss for the generator $G: X \rightarrow Y$ and the discriminator D_Y as follows:

$$\mathcal{L}_{GAN}(G, D_Y, X, Y) = \mathbb{E}_{y \sim p_{\text{data}}(y)} [\log D_Y(y)] + \mathbb{E}_{x \sim p_{\text{data}}(x)} [\log(1 - D_Y(G(x)))] \quad (\text{Equation 3})$$

They similarly define the adversarial loss for the generator $F: Y \rightarrow X$ and the discriminator D_X as $\mathcal{L}_{GAN}(F, D_X, Y, X)$. Based on the adversarial loss, they proposed a cycle-consistency loss to encourage $F(G(x)) \approx x$ and $G(F(y)) \approx y$. The cycle-consistency loss is expressed as:

$$\mathcal{L}_{\text{cyc}}(G, F) = \mathbb{E}_{x \sim p_{\text{data}}(x)} [F(G(x)) - x] + \mathbb{E}_{y \sim p_{\text{data}}(y)} [G(F(y)) - y] \quad (\text{Equation 4})$$

The full objective of CycleGAN is

$$\min_{G, F} \max_{D_X, D_Y} \mathcal{L}(G, F, D_X, D_Y) = \mathcal{L}_{GAN}(G, D_Y, X, Y) + \mathcal{L}_{GAN}(F, D_X, Y, X) + \lambda \mathcal{L}_{\text{cyc}}(G, F), \quad (\text{Equation 5})$$

where λ is a hyperparameter used to control the relative importance of the adversarial loss and the cycle-consistency loss.

As a powerful generative model, many variants of GANs were presented by modifying loss functions or network architectures and were used for various computer vision tasks. In this review, we mainly focus on the problem of scene transfer and task transfer in autonomous systems using GANs, including image style transfer, image SR, image denoising/dehazing/rain removal, semantic segmentation, depth estimation, pedestrian detection, and person re-ID.

Reinforcement Learning

Reinforcement learning is the problem faced by an agent that learns behavior through trial-and-error interactions in a dynamic environment.⁹⁰ In the RL framework, an agent interacts with the environment to choose the action in the state of a given environment in order to maximize its long-term reward.⁹¹ RL algorithms can be classified into two kinds, model-based and model-free algorithms.⁹² Model-based RL is to learn a transition model that allows the environment to be simulated without directly interacting with the environment.⁶⁴ Model-based methods include guided policy search (GPS)⁵⁰ and model-based value expansion.⁹³ However, model-free RL uses the experience of states and environments directly to generate actions.⁹⁴ Model-free methods include deep Q network (DQN),⁹⁵ deep deterministic policy gradient (DDPG) method,⁹⁶ dynamic policy programming (DPP) method,⁹⁷ and asynchronous advantage actor-critic (A3C) method.⁶¹ Model-free algorithms can learn complex tasks but tend to be inefficient in sampling, while model-based algorithms are more efficient in sampling but usually have difficulty in scaling to complicated tasks.⁴⁸ With further research on the application of RL methods, several problems occur in that model-based algorithms are no longer applicable to more complex tasks, while model-free algorithms need more training data. Moreover, when the given environment changes or training data are insufficient, the chances are that RL methods need to train the model starting from the scratch, which is inefficient and inaccurate. Therefore, RL methods are limited when generalizing to different tasks and domains.⁴⁸ In this review, we mainly focus on several modifications on RL methods, such as amending the network structure^{98,99} and optimizing the way of training,^{100,101} to enable the model to learn the new tasks accurately in the same domain or transferably across domains.

Meta-Learning

Meta-learning, or “learning to learn” and “learning how to learn,” uses previous knowledge and experience to guide the learning of new tasks to equip the model with the ability to learn across domains.¹⁰² The goal of meta-learning is to train a model that can quickly adapt to a new task using only a few data points and training iterations.⁶⁰ Similar to transfer learning, meta-learning improves the learner’s generalization ability in a multi-task setting. However, unlike transfer learning, meta-learning focuses on the sampling of both data and tasks. Therefore, meta-learning models are trained by being exposed to a large number of tasks, which qualifies them to learn new tasks from few data settings. The meta-learning methods can be divided into three categories: recurrent models, metric learning, and learning optimizers.¹⁰³

Recurrent models are trained by various methods, such as long short-term memory (LSTM)¹⁰⁴ and temporal convolution network (TCN),¹⁰⁵ to acquire the dataset sequentially and then process new inputs from the task. LSTM¹⁰⁴ processes data sequentially and figures out its own learning strategy from

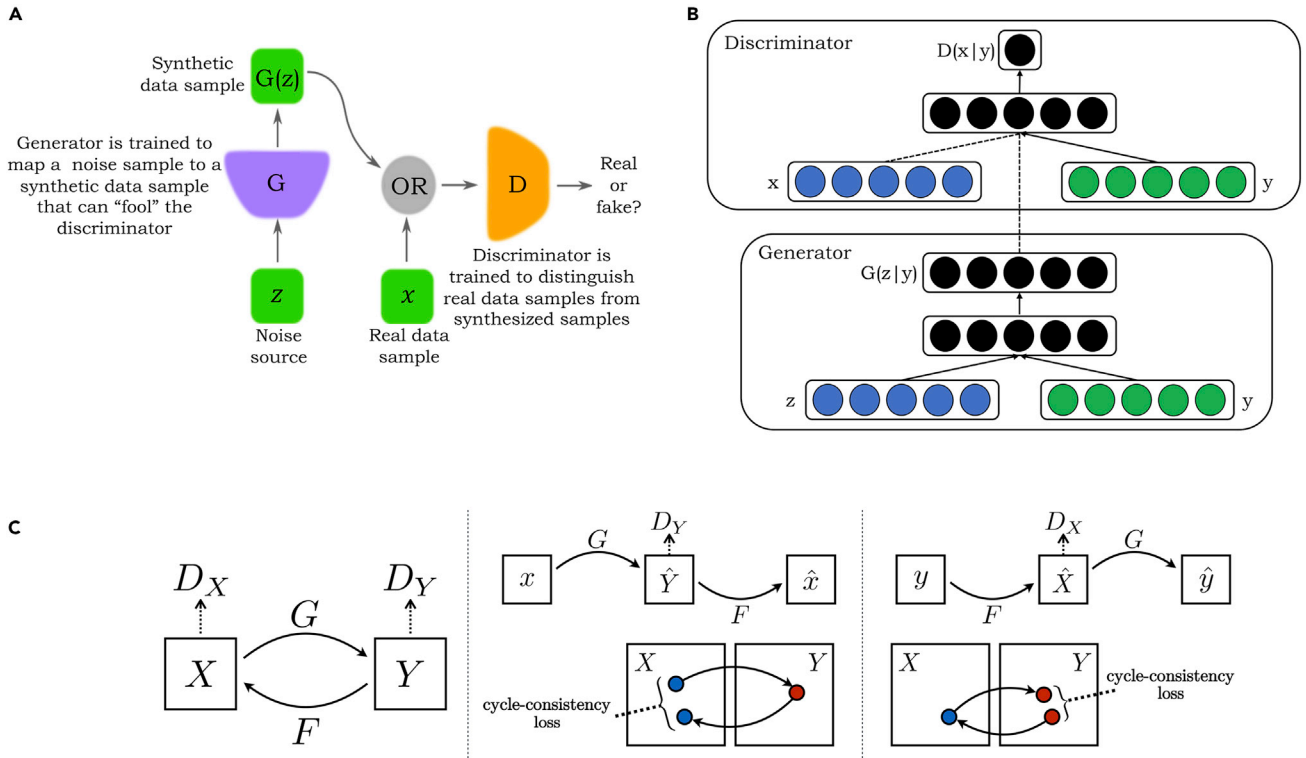


Figure 1. Generative Adversarial Networks and Several Typical Variants

(A) Generative adversarial networks. Copyright (2018) IEEE. Reprinted, with permission, from Creswell et al.⁸⁸
 (B) Conditional generative adversarial networks. From Mirza and Osindero.⁸⁹
 (C) Cycle-consistent adversarial networks. Copyright (2017) IEEE. Reprinted, with permission, from Zhu et al.²⁴

scratch. Moreover, TCN¹⁰⁵ uses convolution structures to capture long-range temporal patterns, whose framework is simpler and more accurate than LSTM.

Metric learning is a way to calculate the similarity between two targets from different tasks. For a specific task, the input target is classified into a target category with large similarity judging from a metric distance function.¹⁰⁶ It has been widely used for few-shot learning,¹⁰⁷ during which the data belong to a large number of categories; some categories are unknown at the stage of training and the training samples of each category are particularly small.¹⁰⁸ These characteristics are consistent with the characteristics of meta-learning. There are four sorts of typical networks proposed for metric learning: Siamese network,¹⁰⁹ prototypical network,¹⁰³ matching network,¹¹⁰ and relation network.¹¹¹

Learning is optimized, i.e., one meta-learner learns how to update the learner so that the learner can learn the task efficiently.¹¹² This method has been extensively studied to obtain better optimization results of neural networks. Combined with RL¹¹³ or imitation learning,⁶² meta-learning is able to learn new policies accurately or adapt to new tasks effectively. Model-agnostic meta-learning (MAML)⁶⁰ is a representative and popular meta-learning optimization method, which uses stochastic gradient descent (SGD)¹¹⁴ to update. It adapts quickly to new tasks due to no assumptions being made about the form of the model and no extra parameters being introduced for meta-learning. MAML includes a base-model learner and a meta-

learner. Each base-model learner learns a specific task and the meta-learner learns the average performance θ of multiple specific tasks as the initialization parameters of one new task.⁶⁰ From Figure 2, the model is represented by a parameterized function f_θ with the parameter θ . When adapting to a new task T_i that is drawn from a distribution over tasks $p(T)$, the model's parameter is updated to θ'_i . θ'_i is computed by one or more gradient descent updates on task T_i . Moreover, $\mathcal{L}_{T_i}$ represents the loss function for task T_i and the step size α is regarded as a hyperparameter. For example, we consider one gradient update on task T_i ,

$$\theta'_i = \theta - \alpha \nabla_{\theta} \mathcal{L}_{T_i}(f_{\theta}). \quad (\text{Equation 6})$$

The model parameters are trained by optimizing for the performance of a parameterized function $f_{\theta'_i}$ with parameter θ'_i , corresponding to the following problem:

$$\min_{\theta} \sum_{T_i \sim p(T)} \mathcal{L}_{T_i}(f_{\theta'_i}) = \sum_{T_i \sim p(T)} \mathcal{L}_{T_i}(f_{\theta - \alpha \nabla_{\theta} \mathcal{L}_{T_i}(f_{\theta})}). \quad (\text{Equation 7})$$

When extending MAML to the imitation learning setting, the model's input, $\mathbf{o}_t$, is the agent's observation sampled at time t , whereas the output $\mathbf{a}_t$ is the agent's action taken at time t . The demonstration trajectory can be represented as $\tau := \{\mathbf{o}_1, \mathbf{a}_1, \dots, \mathbf{o}_T, \mathbf{a}_T\}$, using a mean squared error loss as a function of policy parameters φ as follows:

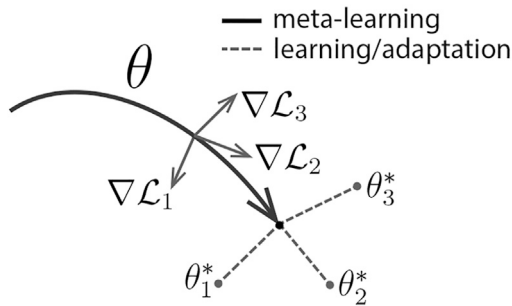


Figure 2. Diagram of the MAML Algorithm, which Optimizes for a Representation θ That Can Quickly Adapt to New Tasks
From Finn et al.⁶⁰

$$\mathcal{L}_{T_i}(\mathbf{f}_\varphi) = \sum_{\tau^{(j)} \sim T_i} \sum_t \mathbf{f}_\varphi(\mathbf{o}_t^{(j)}) - \mathbf{a}_t^{(j)2}. \quad (\text{Equation 8})$$

During meta-training, several demonstrations are sampled as training tasks. The demonstrations help to compute θ_i^* for each task T_i using gradient descent with Equation 6 and to compute the gradient of the meta-objective by using Equation 7 with the loss in Equation 8. During meta-testing, we consider using only a single demonstration as a new task T , updating with SGD. Therefore, the model is updated to acquire a policy for that task.¹¹⁵

The Relationship between Adversarial Learning, RL, and Meta-Learning

RL¹¹⁶ is a method to describe and solve the problem that agents learn policies to achieve the maximum returns or specific goals in the interactions with the environment. Pfau and Vinyals¹¹⁷ discussed the connection between GANs and actor-critic (AC) methods. AC is a kind of RL method that learns the policy and value function simultaneously. To be specific, the actor network chooses the proper action in a continuous action space, while the critic network implements a single-step update, which improves the learning efficiency.¹¹⁸ Pfau and Vinyals¹¹⁷ argued that GANs can be viewed as an AC approach in an environment where actors cannot influence rewards. RL and GANs are integrated for various tasks, such as real-time point cloud shape completion¹¹⁹ and image synthesis.¹²⁰

In the field of RL, using the cost function to understand the underlying behavior is called inverse reinforcement learning (IRL).¹²¹ The policy distribution in the IRL can be regarded as the data distribution of the generator in GANs, and the reward in the IRL can be regarded as the discriminator in GANs. However, IRL learns the cost function to explain expert behavior but cannot directly tell the learner how to take action, which leads to high running costs. Ho and Ermon¹²² proposed generative adversarial imitation learning (GAIL), combining GANs with imitation learning, which employs GANs to fit the states and actions distributions that define expert behavior. GAIL significantly improves the performance in large-scale and high-dimensional planning problems.¹²³

Introducing meta-learning to RL methods is called meta-RL methods,¹²⁴ which equips the model to solve new problems more efficiently by utilizing the experience from prior tasks. A meta-RL model is trained over a distribution of different but

related tasks, and during testing it is able to learn to solve a new task quickly by developing a new RL algorithm.¹²⁵ There are several meta-RL algorithms that utilize past experience to achieve good performance on new tasks. For example, MAML⁶⁰ and Reptile¹²⁶ are typical methods for updating model parameters and optimizing model weights; MAESN (model-agnostic exploration with structured noise)¹²⁷ can learn structured action noise from prior experience; evolved policy gradient¹²⁸ defines the policy gradient loss function as a temporal convolution over previous experience. Moreover, when dealing with unlabeled training data, unsupervised meta-RL methods¹²⁹ effectively acquire accelerated RL procedures without manual task design, such as collecting data and labeling data. Therefore, both supervised and unsupervised meta-RL can transfer previous task information to new tasks across domains.

Autonomous Systems Meet Accuracy and Transferability

Computer vision and robot control tasks are critical to autonomous systems. Currently there is a variety of learning-based methods for perception and decision-making tasks. As mentioned at the beginning of the previous section (Preliminaries), most traditional learning-based methods show good accuracy in the same data distribution or task but suffer from poor transferability; specifically, when considering the application of a well-trained model to different scenarios, its accuracy often decreases heavily. This is due to the obvious domain gap between different datasets. Therefore, domain adaptation between different domains is very important for autonomous systems.

Overview of the Section

In this section, we mainly focus on learning-based methods in the perception and decision-making tasks of autonomous systems, from the perspectives of accuracy or transferability or both, such as image style transfer, image SR, image denoising/dehazing/rain removal, semantic segmentation, depth estimation, other geometry information (surface normal and optical flow) prediction, pedestrian detection/re-ID/tracking, robot navigation, and robotic manipulation in autonomous systems. Although some traditional DL-based methods mainly focus on improving the accuracy of the model, in recent years methods for the above visual tasks have gradually attached importance to the transferability, using adversarial learning, RL, and meta-learning. We summarize some typical computer vision tasks and robot control tasks in autonomous systems in Tables 3 and 4, including their training manners, loss functions, learning methods, and experimental platforms. As shown in Table 3, the training manners of some computer vision tasks gradually change from supervised to unsupervised ones, and their loss functions change from accuracy to transferability between domains. Table 4 indicates that for robot control tasks, informative simulation environments and flexible practice platforms will help to accurately transfer information across domains.

Image Style Transfer

Images can be well transferred between different styles, which is conducive to the perception and decision-making algorithms of autonomous systems applicable to various scenarios. Autonomous systems inevitably face the problem of the image style transfer arising from seasonal conversion,⁷⁴ varying weather conditions,⁷² or day conversion.⁷³ In particular, it is more

Table 3. Summary of Methods for Computer Visual Tasks in Autonomous Systems

Year	Reference	Task	Multi-Task	GANs-Based	Supervision ^a	Loss ^b
2016	Gatys et al. ¹³⁰	style transfer			Supervised	C
2016	Johnson et al. ¹³¹	style transfer	√		Supervised	B
2017	Li et al. ¹³²	style transfer			Supervised	C
2017	Pix2Pix ³⁴	style transfer	√	√	Supervised	A, E
2017	CycleGAN ²⁴	style transfer	√	√	Unsupervised	A, D
2019	DLOW ¹³³	style transfer	√	√	Unsupervised	A, D
2019	INIT ¹³⁴	style transfer		√	Unsupervised	A, C
2014	SRCNN ⁷	super-resolution			Supervised	F
2015	SRCNN ⁸	super-resolution			Supervised	F
2016	FSRCNN ¹³⁵	super-resolution			Supervised	F
2016	Johnson et al. ¹³¹	super-resolution	√		Supervised	B
2017	SRGAN ²⁶	super-resolution		√	Supervised	A, F
2017	EnhanceNet ¹³⁶	super-resolution		√	Supervised	A, B, F
2018	ZSSR ¹³⁷	super-resolution			Unsupervised	E
2018	ESRGAN ¹³⁸	super-resolution		√	Supervised	A, B, E
2018	CinCGAN ¹³⁹	super-resolution		√	Unsupervised	A, D, F
2019	Soh et al. ¹⁴⁰	super-resolution		√	Supervised	A, C, F
2020	Gong et al. ¹⁴¹	super-resolution		√	Unsupervised	A, D, E
2018	DeblurGAN ²⁷	image deblurring		√	Supervised	A, B
2019	DeblurGAN-v2 ¹⁴²	image deblurring		√	Supervised	A, E, F
2019	Dr-Net ¹⁴³	image deblurring		√	Supervised	A, E
2018	Li et al. ¹⁴⁴	image dehazing		√	Supervised	A, B, E
2018	Cycle-Dehaze ¹⁴⁵	image dehazing		√	Unsupervised	A, D
2019	Kim et al. ¹⁴⁶	image dehazing		√	Supervised	A, D, E, F
2019	CDNet ¹⁴⁷	image dehazing		√	Unsupervised	A, D
2020	Sharma et al. ¹⁴⁸	image dehazing		√	Supervised	A, B, E, F
2018	Qian et al. ³⁵	image rain removal		√	Supervised	A, B, F
2019	Li et al. ¹⁴⁹	image rain removal		√	Supervised	A, B, F
2019	ID-CGAN ³⁶	image rain removal		√	Supervised	A, B, E
2020	AI-GAN ¹⁵⁰	image rain removal		√	Supervised	A, F
2016	Hoffman et al. ⁷¹	semantic segmentation			Unsupervised	F, G
2017	SegNet ¹³	semantic segmentation			Supervised	F
2017	Mask R-CNN ¹⁵¹	instance segmentation			Supervised	F
2017	CyCADA ⁷⁴	semantic segmentation		√	Unsupervised	A, D, F
2018	FCAN ¹⁵²	semantic segmentation		√	Unsupervised	A, F
2018	Hu et al. ¹⁵³	instance segmentation			partially supervised ^c	F
2018	Hong et al. ³⁹	semantic segmentation		√	Unsupervised	A, F, G
2019	CrDoCo ¹⁵⁴	semantic segmentation	√	√	Unsupervised	A, C, D, F
2019	CLAN ¹⁵⁵	semantic segmentation		√	Unsupervised	A, F
2019	Li et al. ¹⁵⁶	semantic segmentation		√	self-supervised	A, B, C, F
2020	Erkent et al. ¹⁵⁷	semantic segmentation		√	Unsupervised	A, F
2014	Eigen et al. ¹⁴	depth estimation			Supervised	F
2015	Eigen et al. ¹⁵	depth estimation	√		Supervised	F
2015	Liu et al. ¹⁵⁸	depth estimation			Supervised	F
2018	Atapour-Abarghouei et al. ²⁵	depth estimation		√	Supervised	A, C
2019	ASM ¹⁵⁹	depth estimation	√		Supervised	F
2019	CrDoCo ¹⁵⁴	depth estimation	√	√	Unsupervised	A, C, D, F
2019	GASDA ¹⁶⁰	depth estimation		√	Unsupervised	A, D, F
2020	ARC ¹⁶¹	depth estimation		√	Supervised	A, B, C, D, F

(Continued on next page)

Table 3. Continued

Year	Reference	Task	Multi-Task	GANs-Based	Supervision ^a	Loss ^b
2013	ConvNet ¹⁶²	pedestrian detection			Unsupervised	F
2015	TA-CNN ¹⁶	pedestrian detection			Supervised	F
2017	SAF R-CNN ¹⁶³	pedestrian detection			Supervised	F
2019	Kim et al. ⁴¹	pedestrian detection		✓	Unsupervised	A, E, F
2018	SPGAN ⁴²	person re-ID		✓	Unsupervised	A, D, F
2018	CamStyle ¹⁶⁴	person re-ID		✓	Unsupervised	A, D, F
2019	ATNet ¹⁶⁵	Person re-ID		✓	Unsupervised	A, D, F
2017	ADNet ¹⁶⁶	pedestrian tracking	✓		Supervised	F
2017	Supancic et al. ¹⁸	pedestrian tracking			supervised	— ^d
2018	Chen et al. ¹⁶⁷	pedestrian tracking			supervised	E
2019	ConvNet-LSTM ⁵²	pedestrian tracking			supervised	— ^d

^aFor models that do not explicitly state whether they are supervised in the references, this review considers models that require paired images as supervised and models that do not require paired images as unsupervised.

^bWe classify the loss function into several classes. “A” represents adversarial (GAN) loss. “B” represents perceptual loss. “C” represents reconstruction loss. “D” represents cycle consistency loss. “E” represents pixel-wise loss. “F” represents specific task loss such as depth loss and semantic loss. “G” represents domain transfer loss such as domain adversarial loss and domain classifier loss.

^cPartially supervised learning problems refer to training on the combination of strong and weak labels.¹⁵³ According to Schwencker and Trentin,¹⁶⁸ partially supervised learning includes active learning, general semi-supervised learning, semi-supervised learning with graphs, partially supervised learning in ensembles, and multiple classifier systems.

^dRL-based methods mainly focus on reward and action instead of loss.

challenging and interesting to consider transferring training data for night to day, rainy to sunny, or winter to summer, since most autonomous systems have a better ability to perceive under good lighting or weather conditions than some harsh environments. The task of image style transfer is to change the content of the source domain image to the target domain one while ensuring that the style is consistent with the target domain.¹³⁰ In addition, style transfer, as an interesting data augmentation strategy, can extend the range of lighting and weather changes, thus further improving the transferability of the model.¹⁹¹ In addition, using the image style transfer algorithm to achieve the transfer from the simulated environment to the real world is very useful for semantic segmentation, robot navigation, and grasping tasks, because training directly in the real world may lead to higher experimental costs due to some possible damage to hardware.¹⁹¹ Traditional methods to achieve style transfer mainly rely on non-parametric techniques to manipulate the pixels of the image (e.g., Efros and Freeman,¹⁹² Hertzmann et al.¹⁹³). Although traditional methods have achieved good results in style transfer, they are limited to using only low-level features of the image for texture transfer, but not semantic transfer.¹³⁰

Traditional DL-Based Style Transfer. Convolutional neural networks (CNNs) have been used in image style transfer, since they have achieved impressive results in numerous visual perception areas. Gatys et al.¹³⁰ first proposed to utilize CNNs (pre-trained VGG Networks) to separate content and style from natural images, and then combined the content of one image with the style of another into a new image to achieve an artistic style transfer. This work opened up a new viewpoint for style transfer using DNNs. To reduce the computational burden, Johnson et al.¹³¹ proposed to use the perceptual loss instead of the per-pixel loss for the image style transfer task. This method achieves results similar to those of Gatys et al.¹³⁰ while being three orders

of magnitude faster. Since Gatys et al.¹³⁰ used the Gram matrices to represent the artistic style of an image, the subsequent improvement works did not investigate its principles in depth. Li et al.¹³² first regarded neural style transfer as a domain adaptation problem, and theoretically showed that the second-order interaction in the Gram matrix is not necessary for style transfer, which is equivalent to a specific maximum mean discrepancy. In addition, Chen et al.¹⁹⁴ presented a stereo neural style conversion that can be used in emerging technologies such as three-dimensional (3D) movies or virtual reality. This method seems promising for improving the perception accuracy of autonomous systems in unmanned scenes because the transferred results contain more stereo information in the scene.

GANs-Based Style Transfer. Traditional CNN-based methods minimize the Euclidean distance between predicted pixels and ground-truth pixels, which may cause blurry results.³⁴ GANs can be used for image style transfer, which can produce more realistic images.³⁴ Isola et al.³⁴ used cGANs to transfer image style, and the experimental results showed that cGANs (with L1 loss) not only have satisfactory results for style transfer tasks but also can produce reasonable results for a wide variety of problems such as semantic segmentation and background removal. However, this method requires paired image samples, which is often difficult to implement in practice. By considering this issue, Zhu et al.²⁴ proposed CycleGAN to learn image translation between domains with unpaired examples, as shown in Figure 3. As mentioned in Preliminaries, the framework of CycleGAN includes two generators and two discriminators to achieve mutual translation between the source and the target domain. The main insight of CycleGAN is to preserve the key attributes between the input and the translated image by using a cycle-consistency loss. At almost the same time, DiscoGAN¹⁹⁵ and DualGAN¹⁹⁶ were presented to adopt similar cycle-consistency ideas to achieve an image transfer task across domains. To

Table 4. Summary of Traditional RL/Meta-Learning Methods for Scenario-Transfer Tasks

Year	Reference	Task	RL Method	Meta-Learning Method	Simulation Platform	Practice Platform
2016	Sadeghi et al. ¹⁶⁹	UAV navigation	F		3D CAD environment	Parrot Bebop
2017	Tai et al. ¹⁷⁰	robot navigation	I		V-REP	Turtlebot
2017	Zhang et al. ¹⁷¹	robot navigation	H		maze-like 3D environment	Robotino
2017	Polvara et al. ⁹⁹	UAV navigation	H		gazebo	Parrot AR Drone 2
2017	Zhu et al. ⁶¹	robot navigation	K	B	AI2-THOR	SCITOS
2018	Banino et al. ¹⁷²	robot navigation	K	A	multi-room 2D environment	None
2018	Faust et al. ²²	robot navigation	I		simulated building plans	differential drive robot
2019	Zhu et al. ¹⁷³	robot navigation	K	A	SUNCG	Matterport3D
2019	Niroui et al. ¹⁷⁴	robot navigation	K	A	Turtlebot Stage simulator	Turtlebot
2019	Wortsman et al. ¹⁷⁵	robot navigation		A, C, E	AI2-THOR	none
2019	Jabri et al. ¹⁷⁶	robot navigation		E	ViZDoom	none
2019	Koch et al. ¹⁷⁷	UAV navigation	I, O, N		GymFC	none
2020	Gaudet et al. ¹⁷⁸	UAV navigation		E	Mars and asteroid landing simulation	none
2015	Zhang et al. ¹⁷⁹	robotic manipulation	H		none	Baxter arm
2016	Levine et al. ⁵⁰	robotic manipulation	L		MuJoCo	PR2 robot
2017	Gu et al. ¹⁸⁰	robotic manipulation	M		MuJoCo	7-DoF arm
2017	Finn et al. ¹¹⁵	robotic manipulation		C, D	MuJoCo	7-DoF PR2 arm
2018	Haarnoja et al. ¹⁸¹	robotic manipulation	G		MuJoCo	7-DoF Sawyer arm
2018	Zhu et al. ¹⁸²	robotic manipulation	N	A	MuJoCo	Jaco robot arm
2018	Zeng et al. ¹⁰¹	robotic manipulation	H		V-REP	UR5 robot arm et al.
2018	Yu et al. ¹⁸³	robotic manipulation		C, D	MuJoCo	7-DoF PR2 arm et al.
2019	Yu et al. ¹⁸⁴	robotic manipulation	N, O, J	C	MuJoCo	none
2019	Zeng et al. ¹⁸⁵	robotic manipulation		B	None	Amazon Robotics Challenge
2019	Tsurumine et al. ¹⁸⁶	robotic manipulation	P		n-DoF simulated manipulator	15-DoF humanoid robot
2020	Singh et al. ¹⁸⁷	robotic manipulation		D	bullet physics engine	none

We classify the meta-learning methods into several classes. “A” represents recurrent network. “B” represents metric network. “C” represents MAML. “D” represents meta-imitation learning. “E” represents meta-RL. Similarly, we classify the RL methods into several classes. “F” represents Fitted Q-iteration. “G” represents soft Q-learning. “H” represents DQN. “I” represents DDPG. “J” represents soft AC. “K” represents A3C. “L” represents GPS. “M” represents asynchronous NAF (normalized advantage function).¹⁸⁸ “N” represents PPO (proximal policy optimization).¹⁸⁹ “O” represents TRPO (trust region policy optimization).¹⁹⁰ “P” represents DPP. DoF, degrees of freedom.

improve CycleGAN from the aspect of semantic information alignment at the feature level, Hoffman et al.⁷⁴ proposed CyCADA by combining domain adaptation and cycle-consistent adversarial, which uniformly considers feature-level and pixel-level adversarial domain adaptation and cycle-consistency constraints. CyCADA has achieved satisfactory results in some challenging tasks, such as from synthesis to practical conversion and seasonal conversion, which is very important for the generalization of autonomous systems. It was shown that CyCADA has a better transferability than the original CycleGAN model. Since these methods, such as CycleGAN and CyCADA, can only realize the translation between two domains, different models should be trained for each pair of domains in the case of handling multiple-domain translation tasks, which limits their wide application. By considering this point, Choi et al.¹⁹⁷ proposed StarGAN to perform image translations for multiple domains using a single generator and a discriminator. StarGAN takes both the image and its domain label as input, and learns to transfer the input image into the corresponding target domain. To further improve the existing adaptive image style transfer methods,

Gong et al.¹³³ proposed a domain flow generation (DLOW) model, which generates a series of intermediate domains to bridge two different domains. This method may be helpful for gradual changes, such as day or season, because it can generate a continuous sequence of intermediate samples ranging from the source to target samples. Recent image translation tasks focused on semantic consistency of images instead of image style and content. Royer et al.¹⁹⁸ proposed XGAN, which is an unsupervised semantic style transfer task for many-to-many mapping. Royer et al. used domain adaptation techniques to constrain the shared embedding and proposed a semantic consistency loss as a form of self-supervision to act on two domain translations. This method has a good generalization effect when there is a large domain shift between the two domains. In addition, to obtain fine-grained local information of images, Shen et al.¹³⁴ proposed the instance-aware image-to-image translation approach, which applies instance and global styles to the target image spatially, as shown in Figure 3. Similarly, the image style transfer was considered at the instance level by Ma et al.¹⁹⁹ and Mo et al.²⁰⁰

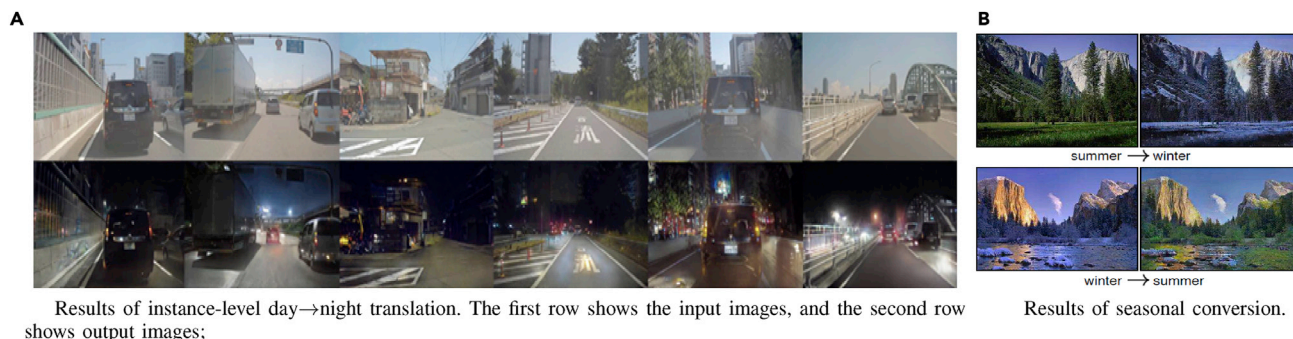


Figure 3. Generative Adversarial Networks for Image Style Transfer

(A) Results of instance-level day $\rightarrow$ night translation. Copyright (2019) IEEE. Reprinted, with permission, from Shen et al.¹³⁴
(B) Results of seasonal conversion. Copyright (2017) IEEE. Reprinted, with permission, from Zhu et al.²⁴
See also Table 3.

As a data augmentation strategy, image style transfer can help the scene to be transferred in various lighting conditions, various weather conditions, simulation of real-world environment, and so forth. Image style transfer helps autonomous systems perform their perception and decision-making tasks in better lighting conditions, and effectively reduces hardware losses in the real-world environment, which is critical for autonomous systems. Although many traditional DL-based models have achieved good style transfer results, with the advent of GANs various research works have been extended based on GANs. The recent developments of image style transfer have focused on instance-level style transfer. We believe that future works should focus on the image style transfer for more complex scenes, such as changing the style of the specified instance without changing the background style in a wild environment. In addition, future research should also consider improving the accuracy of style transfer and the speed of the overall process, striving for real-time performance with good accuracy. In addition to the style transfer task, we further consider increasing the resolution of images, i.e., the SR task.

Super-resolution

SR is a challenging visual perception task to generate high-resolution (HR) images from low-resolution (LR) image inputs.²⁰¹ SR is crucial to understanding the environment at high level for autonomous systems. For example, SR is helpful in constructing a dense map. In this subsection, we first discuss the recent developments in SR by focusing on accuracy. We then summarize the new developments in SR by considering transferability.

There are a number of methods dedicated to improving image quality, such as single-image interpolation²⁰² and image restoration.²⁰³ It is worth pointing out that these are different from SR. On the one hand, single-image interpolation usually cannot restore high-frequency details.²⁰² In addition, image restoration often uses methods such as image sharpening, whereby the input image and output image remain the same size, although the output quality can be improved.²⁰³ SR does not only improve the output quality but also increases the number of pixels per unit area, i.e., the size of the image increases.²⁰¹ In some cases, the image SR can be regarded as a method of image enhancement.²⁰⁴ Recently, a large number of SR methods have been proposed, such as interpolation-based methods²⁰⁵ and reconstruction-based methods.²⁰⁶ Farsiu et al.²⁰⁷ introduced the advances and challenges of traditional methods for SR.

Traditional DL-Based SR. There are some results studying traditional DL-based methods without adversarial learning for SR, which are mainly CNN-based. Dong et al.⁷ considered using CNNs to handle SR tasks in an end-to-end manner. They presented the SR convolutional neural network (SRCNN), which has little extra pre-/post-processing beyond optimization. In addition, they confirmed that DL provides a better quality and speed for SR than the sparse coding method²⁰⁸ and the K-SVD-based method,²⁰⁹ while SRCNN only uses information on the luminance channel. Dong et al.⁸ then extended SRCNN to process three color channels simultaneously to improve the accuracy of SR results. Considering the poor real-time performance of SRCNN, Dong et al.¹³⁵ utilized a compact hourglass-shape CNN structure to accelerate the current SRCNN. In fact, most learning-based SR methods use the per-pixel loss between the output image and the ground-truth image.^{7,8} Johnson et al.¹³¹ considered the use of perceptual loss to achieve a better SR, which is able to better reconstruct details than the per-pixel loss. Note that the aforementioned SR methods often rely on specific training data. When there are non-ideal imaging conditions due to noise or compression artifacts, these methods usually fail to provide good SR results. Therefore, Shocher et al.¹³⁷ considered “zero-shot” SR (ZSSR), which does not rely on prior training. To the best of our knowledge, ZSSR is the first unsupervised CNN-based SR method, which achieves reasonable SR results in some complex or unknown imaging conditions. Due to the lack of recurrence of blurry LR images, ZSSR is less effective for SR when facing very blurry LR images. By taking into account this issue, Zhang et al.²¹⁰ proposed a deep plug-and-play SR framework for LR images with arbitrary blur kernels. This modified framework is flexible and effective in dealing with very blurry LR images. Recent trends in SR also include SR for stereo images²¹¹ and 3D appearance.²¹²

GANs-Based SR. In addition to the traditional DL-based SR methods, GANs show their promising results in SR. The use of GANs for SR has the advantage of bringing the generated results closer to the natural image manifold, which may improve the accuracy of the result.²⁶ The representative work on GANs-based SR (SRGAN) was presented by Ledig et al.,²⁶ which combines a content loss with an adversarial loss by training a GAN. This method is capable of reconstructing photo-realistic natural images for an upscaling factor of 4 $\times$. Although the SRGAN

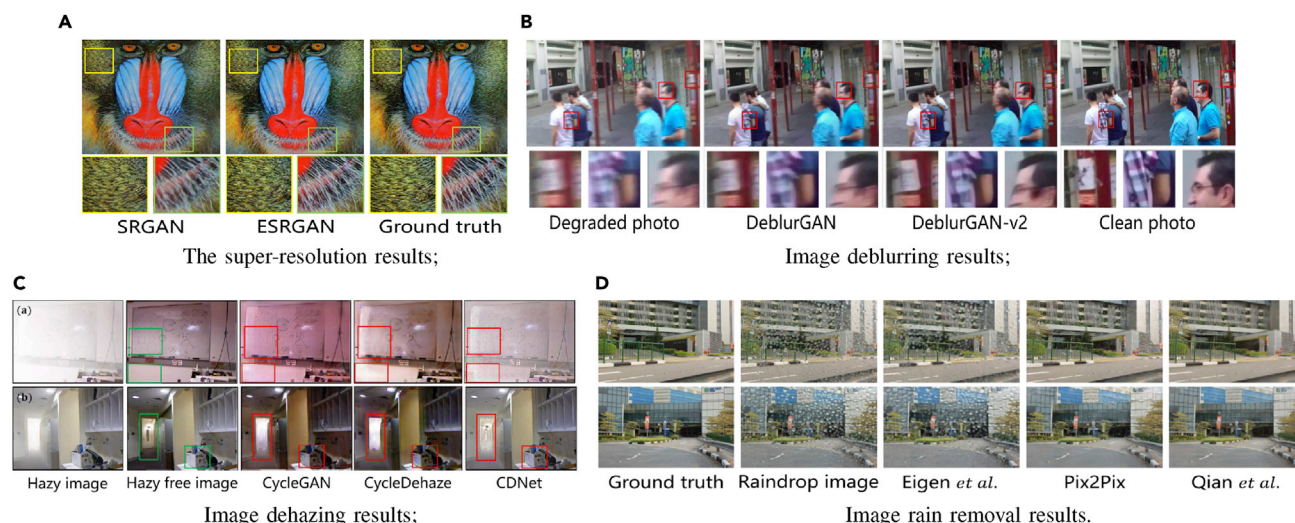


Figure 4. Generative Adversarial Networks for SR and Image Deblurring/Dehazing/Rain Removal

(A) The super-resolution results of $\times 4$ for SRGAN, ESRGAN, and the ground truth. Reprinted by permission of Wang et al.¹³⁸ Copyright.

(B) Image deblurring results. Copyright (2019) IEEE. Reprinted, with permission, from Kupyn et al.¹⁴²

(C) Image dehazing results. Copyright (2019) IEEE. Reprinted, with permission, from Dudhane and Murala.¹⁴⁷

(D) Image rain removal results. Copyright (2018) IEEE. Reprinted, with permission, from Qian et al.³⁵

See also Table 3.

achieves good SR results, the local matching of texture statistics is not considered, which may restrict the improvement of the SR results to some extent. By considering this point, Sajjadi et al.¹³⁶ focused on creating realistic textures to achieve SR. They proposed EnhanceNet, which combines adversarial training, perceptual loss, and a newly proposed texture transfer loss to achieve HR results with realistic textures. To further improve the accuracy of SRGAN, Wang et al.¹³⁸ extended SRGAN to ESRGAN by introducing residual-in-residual dense block and improving the discriminator and a perceptual loss. ESRGAN consistently has a better visual quality and natural texture than SRGAN,²⁶ as shown in Figure 4.

HR images are conducive to improving the accuracy of perception tasks in autonomous systems. In autonomous systems more complicated situations may be encountered, such as when HR datasets are unavailable or the input LR images are noisy and blurry, which means that SR cannot be achieved with paired data. Inspired by the cycle consistency of CycleGAN, Yuan et al.¹³⁹ tackled these issues with a cycle-in-cycle network (CinCGAN), which consists of two CycleGANs. The first CycleGAN maps LR images to the clean LR space, in which the proper denoising/deblurring processing is implemented on the original LR input. They then stacked another well-trained deep model to up-sample the intermediate results to the desired size. Finally, they used adversarial learning to fine-tune the network in an end-to-end manner. The second CycleGAN contains the first one to achieve the purpose of mapping from the original LR to the HR. CinCGAN achieves results comparable with those of the supervised method.¹³⁵ Most SR methods trained on synthetic datasets are not effective in the real world. SRGAN and EnhanceNet increase the perceptual quality by enhancing textures, which may produce fake details and unnatural artifacts. Soh et al.¹⁴⁰ focused on the naturalness of the results to reconstruct realistic HR images. On further considering the transferability of

the model, to solve the domain shift between synthetic data and real-world data, Gong et al.¹⁴¹ proposed to further minimize the domain gap by aligning the feature distribution while achieving SR. Specifically, they proposed a method to learn real-world SR images from a set of unpaired LR and HR images, which achieves satisfactory SR results on both paired and unpaired datasets. It is difficult to directly extend the image SR methods to video SR. Recent developments included using the same framework to implement image SR and video SR,²¹³ and real-time video SR using GANs.²¹⁴

Image SR is used to increase the resolution of images, which helps to improve the accuracy of perception tasks. Although various SR models focus on improving accuracy, recent works have focused on the transferability of the model, such as the transfer from synthetic datasets to real-world data. Future works may consider combining SR task with other tasks so that one model can achieve multiple tasks including SR. In addition to image SR tasks, we further consider image restoration, such as image deblurring/dehazing/rain removal.

Image Deblurring, Image Dehazing, and Image Rain Removal

Autonomous systems often encounter poor weather conditions, such as rain and fog. There also exist blurry images due to poor shooting conditions or fast-moving objects. It is well recognized that the accuracy of computer vision tasks heavily depends on the quality of input images. Hence, it is of great importance to study image deblurring/dehazing/rain removal for autonomous systems, which make the high-level understanding of tasks such as semantic segmentation and depth estimation possible in practical applications of autonomous systems. It should be noted that although some image deblurring/dehazing/rain removal tasks use image enhancement algorithms,^{215–217} they are more relevant to image restoration than image enhancement. According to Maini and Aggarwal,²¹⁸ the aim of image

enhancement is to improve the viewer's perception of the image in a way that improves the information content, and it is designed to give emphasis to features of the image. Image restoration aims to restore the noisy/corrupt image to its corresponding clean image, and the corruption may include motion blur and noise.²¹⁹ Therefore, image deblurring/dehazing/rain removal can be regarded as image restoration tasks. When adversarial learning, such as GANs, is used for image deblurring/dehazing/rain removal tasks, it can not only generate realistic images to improve the accuracy of image restoration but also improve the transferability of the models by considering the transfer from synthetic datasets to real-world images.

Image Deblurring. Image blur, which heavily affects the understanding of the surroundings, is widely observed in autonomous systems. To tackle the problem of image deblurring, several traditional DL-based methods without adversarial learning have been successively proposed.^{9,220,221} Considering the convincing performance of GANs in preserving image textures and creating realistic images, as well as being inspired by image-to-image translation with GANs, Kupyn et al.²⁷ regarded image deblurring as a special image-to-image translation task. They proposed DeblurGAN, which is an end-to-end deblurring learning method based on cGANs. This method considers both accuracy and transferability, i.e., DeblurGAN improves deblurring results and it is 5-fold faster than the approach used by Nah et al.²²¹ for both synthetic and real-world blurry images. Kupyn et al.¹⁴² further improved DeblurGAN by adding a feature pyramid network to G and adopting a double-scale D , which is called DeblurGAN-v2. DeblurGAN-v2 achieves better accuracy than DeblurGAN while being 1- to 100-fold faster than competitors, which will make it applicable to real-time video deblurring, as shown in Figure 4. Recently, Aljadaany et al.¹⁴³ presented Dr-Net, which combines Douglas-Rachford iterations and Wasserstein-GAN²²² to solve image deblurring without knowing the specific blurring kernel. In addition, Lu et al.²²³ extracted the content and blur features separately from blurred images to encode the blur features accurately into the deblurring framework. They also utilized the cycle-consistency loss to preserve the content structure of the original images. Considering that stereo cameras are more commonly used in unmanned aerial vehicles, Zhou et al.²²⁴ focused their research on the deblurring of stereo images.

Image Dehazing. Haze is a typical weather phenomenon with poor visibility, which forms a major obstacle for computer vision applications. Image dehazing is designed to recover clear scene reflections, atmospheric light colors, and transmission maps from input images.¹⁴⁵ In recent years, a series of learning-based image dehazing methods have been proposed.^{10,225,226} Although these methods do not require prior information, their dependence on parameters and models may severely cause an impact on the quality of dehazing images. To reduce the effects of intermediate parameters on the model and to establish an image dehazing method with good transferability, a series of GANs-based methods have been proposed for image dehazing. Li et al.¹⁴⁴ tackled image dehazing based on cGAN. Different from the basic cGAN, the generator in this method includes an encoder and decoder architecture, which helps the generator to capture more useful features to generate realistic results. The addition of cGAN makes the method in Li et al.¹⁴⁴ achieve

ideal results on both synthetic datasets and real-world hazy images. Considering the transferability of different scenarios and datasets as well as independence of paired images, Engin et al.¹⁴⁵ proposed the Cycle-Dehaze network by utilizing CycleGAN. This approach adds the cyclic perception-consistency loss and the cycle-consistency loss, thereby achieving image dehazing across datasets with unpaired images. Similar bidirectional GANs for dehazing have also been studied by Kim et al.¹⁴⁶ It is difficult for the Cycle-Dehaze network to reconstruct real scene information without color distortion. Therefore, Dudhane and Murala¹⁴⁷ proposed the cycle-consistent generative adversarial network (CDNet), which utilized the optical model to find the haze distribution from the depth information. CDNet ensures that the fog-free scene is obtained without color distortion. The image dehazing results of Cycle-Dehaze and CDNet are shown in Figure 4. Most image dehazing methods only consider objects at the same scale-space, which will make dehazed images suffer from blurriness and halo artifacts. Sharma et al.¹⁴⁸ considered improving the accuracy and transferability of image dehazing, and presented an approach that can remove haze based on per-pixel difference between Laplacians of Gaussian of hazy images and original haze-free images at a scale-space. The model showed compelling results from simulated datasets to real-world maps, from indoors to outdoors. Recent developments in image dehazing also included targeting different channels, such as color channel,²²⁷ dark channel,²²⁸ and multi-scale networks.²²⁹

Image Rain Removal. Image rain removal is a challenging task because the size, number, and shape of raindrops are usually uncertain and difficult to learn. A number of methods have been proposed for image rain removal, but most of them require stereo image pairs,²³⁰ image sequences,²³¹ or motion-based images.²³² Eigen et al.¹¹ proposed a single-image rain removal method, which is limited to dealing with relatively sparse and small raindrops.

To improve the accuracy of the image rain removal results and considering the outstanding performance of GANs in the image in painting or completion problems, a series of GANs-based methods have been used for image rain removal. Qian et al.³⁵ tackled the heavy raindrop removal from a single image using an attentive GAN. This method uses an attention map in both the generator and the discriminator. The generator produces an attention map through an attention-recurrent network and generates a raindrop-free image together with the input image. The discriminator evaluates the validity of the generation both globally and locally. The rain removal results of Eigen et al.¹¹ and Qian et al.³⁵ are shown in Figure 4. Nevertheless, this method is not suitable for torrential rain removal and is limited to raindrop removal. Heavy rain, strongly visible streaks, or dense rain accumulations make the scene less visible. Li et al.¹⁴⁹ considered the heavy-rain situation and introduced an integrated two-stage CNN, which is able to remove rain streaks and rain accumulation simultaneously. In the first physics-based stage, a streak-aware decomposition module was proposed to decompose entangled rain streaks and rain accumulation to extract joint features. The second refinement stage utilized a cGAN that inputs the reconstructed map of the previous level and generates the final clean image. This method considered the transferability between the synthetic datasets and real-world

images, and has achieved convincing results in both synthetic and real heavy-rain scenarios. To improve the stability of GANs and reduce artifacts introduced by GANs in the output images, Zhang et al.³⁶ proposed an image deraining conditional generative adversarial network (ID-CGAN), which uses a multi-scale discriminator to leverage features from different scales to determine whether the derained image is from real data or generated data. ID-CGAN has obtained satisfactory image rain removal results on both the synthetic dataset and real-world images. Jin et al.¹⁵⁰ considered that existing methods may cause over-smoothing in derained images, and therefore solved the problem from the perspective of feature disentanglement. They introduced an asynchronous interactive GAN (AI-GAN), which not only has achieved good results for image rain removal but also has strong generalization capabilities, which can be used for image/video encoding, action recognition, and person re-ID.

Image deblurring/dehazing/rain removal tasks help to extract more useful information from bad-weather scenes, which can help autonomous systems to better perceive the scene. We focus on introducing GANs-based models, which improve the accuracy or transferability or both of these tasks. Future works may include the image deblurring/dehazing/rain removal tasks as the premise of perception and then integrate a deeper model to achieve scene perception. After the introduction of tasks including image style transfer, image SR, and image deblurring/dehazing/rain removal, we consider high-level perception tasks such as semantic segmentation.

Semantic Segmentation

In emerging autonomous systems, such as autonomous driving and indoor navigation, scene understanding is required by means of semantic segmentation. Semantic segmentation is a pixel-level prediction method that can classify each pixel into different categories corresponding to their labels, such as airplanes, cars, traffic signs, or even backgrounds.²³³ In addition, instance segmentation combines semantic segmentation and object detection to further distinguish object categories in the scene.¹⁵¹ Some traditional DL-based methods without adversarial learning have been proposed and have achieved good accuracy of semantic segmentation^{13,152} and instance segmentation.^{151,153} In practice, such annotations of pixel-level semantic information are usually expensive to obtain. Considering that the semantic labels of synthetic datasets are easy to obtain, it is helpful to consider semantic segmentation on labeled synthetic datasets and then transfer the results to real-world applications. Due to the domain shift between synthetic datasets and real-world images, it is worth exploring how to transfer the model trained on synthetic datasets to real-world images. By considering this point, adversarial learning is used to implement domain adaptation to improve the transferability of the model. Like other computer vision tasks in this review, the trend is now moving from improving accuracy to enhancing transferability. In this subsection, we focus on accuracy or transferability, or both, to review semantic segmentation and instance segmentation tasks.

Traditional DL-Based Semantic Segmentation. Traditional DL-based semantic segmentation algorithms are mainly based on end-to-end convolutional network frameworks. To the best of our knowledge, Long et al.¹² were the first to train an end-to-end fully convolutional network (FCN) for semantic segmenta-

tion. The main insight is to replace fully connected layers with fully convolutional layers to output spatial maps. In addition, they defined a skip architecture to enhance the segmentation results. More importantly, the framework is suitable for input images of arbitrary size and can produce the correspondingly sized output. This work is well recognized as a milestone for semantic segmentation using DL. However, because the encoder network of this method has a large number of trainable parameters, the overall size of the network is large, which results in the difficulty to train FCN. Badrinarayanan et al.¹³ proposed SegNet, which has significantly fewer trainable parameters and can be trained in an end-to-end manner using SGD. SegNet is important in that the decoder performs the non-linear upsampling using the pooling index computed in the max-pooling step of the corresponding encoder, which eliminates the need to learn upsampling. Based on the encoder-decoder network of SegNet, DeepLab uses multi-scale contextual information to enrich semantic information. DeepLab proposed a series of semantic segmentation methods, such as DeepLabv3+,²³⁴ which combines a spatial pyramid pooling module and an encoder-decoder structure for semantic segmentation. In addition, the depthwise separable convolution is applied to both atrous spatial pyramid pooling and the decoder module to make the encoder-decoder network faster and stronger.

The accuracy of unsupervised semantic segmentation tasks is usually worse than that of supervised methods, while supervised semantic segmentation often requires a lot of manual labeling, which is very costly. Note that a synthetic dataset with computer simulation such as Grand Theft Auto²³⁵ can automatically label a large number of semantic tags, which is very important to improving the accuracy of the semantic segmentation model. However, due to the domain shift between the synthetic dataset and the real-world scene, it is necessary to consider domain adaptation in the semantic segmentation task. To address the domain gap problem and improve the transferability of the model, Hoffman et al.⁷¹ proposed a domain adaptation framework with FCN for semantic segmentation, as shown in Figure 5. This method considers aligning both global and local features through some specific adaptation techniques. The method makes full use of the label information of the synthesized dataset and successfully transfers the results from a synthetic dataset to the real scene, in which a satisfactory semantic segmentation result is achieved in practical applications. The same combination of FCN with domain adaptation for semantic segmentation was also presented by Zhang et al.,¹⁵² whereby fully convolutional adaptation networks also successfully explored domain adaptation for semantic segmentation. The model combines appearance adaptation networks and representation adaptation networks to synthesize images for domain adaptation at both the visual appearance level and the representation level. Recent developments also involved 3D semantic segmentation^{236,237} and 3D instance segmentation.²³⁸

Traditional DL-Based Instance Segmentation. The more challenging task is instance segmentation, which combines both object detection and semantic segmentation.¹⁵¹ Li et al.²³⁹ first proposed an end-to-end fully convolutional method for instance-aware semantic segmentation. However, the method produced spurious edges on overlapping instances. He et al.¹⁵¹ proposed Mask R-CNN, which is a classic instance

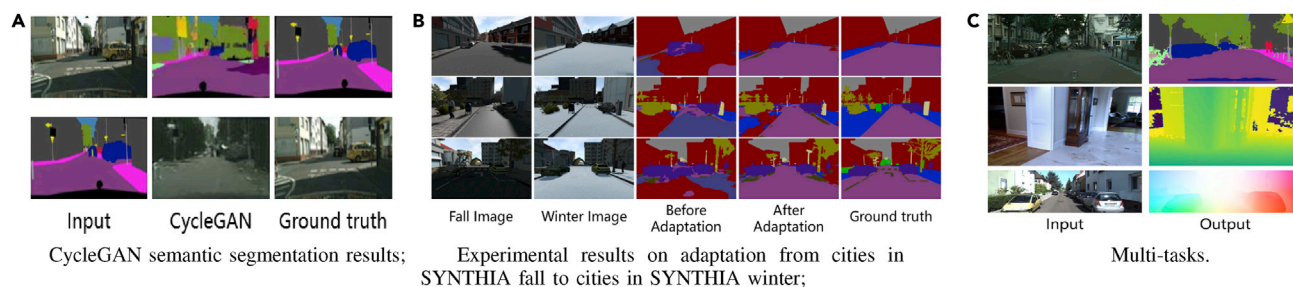


Figure 5. Generative Adversarial Networks for Semantic Segmentation and Multi-Task

(A) CycleGAN for semantic segmentation Copyright (2017) IEEE. Reprinted, with permission, from Zhu et al.²⁴

(B) Qualitative results on adaptation from cities in SYNTHIA fall to cities in SYNTHIA winter. From Hoffman et al.⁷¹

(C) Multi-task includes semantic segmentation (top row), depth prediction (middle row), to optical flow estimation (bottom row). Copyright (2019) IEEE. Reprinted, with permission, from Chen et al.¹⁵⁴

See also Table 3.

segmentation algorithm. Mask R-CNN is easy to train and to generalize to other tasks, i.e., it performs breakthrough results in instance segmentation, bounding-box object detection, and person keypoint detection. This method includes two stages. The first stage proposes a candidate object bounding box. In the second stage, the prediction class and the box offset are in parallel, and the network outputs a binary mask for each region of interest. Mask R-CNN implements instance segmentation in a supervised manner, which is very expensive to semantic labels. In view of this, Hu et al.¹⁵³ proposed a solution to large-scale instance segmentation by developing a partially supervised learning paradigm, in which only a small part of the training process has instance masks and the rest have box annotations. This method has demonstrated exciting new research directions in large-scale instance segmentation.

GANs-Based Semantic Segmentation. GANs are flexible enough to reduce the differences between the segmentation result and the ground truth, and further improve the accuracy of the semantic segmentation results without manual labeling in some cases.²⁴⁰ Regarding the use of GANs for semantic segmentation, the typical methods are Pix2Pix³⁴ and CycleGAN.²⁴ The semantic segmentation result for CycleGAN is shown in Figure 5. There are several variants based on Pix2Pix and CycleGAN.^{74,241,242} These methods not only achieve satisfactory results in image style transfer but also work well in semantic segmentation. Most of the adversarial domain-adaptive semantic segmentation methods for subsequent improvements of CycleGAN and Pix2Pix improve the training stability and transferability by improving loss functions or network layers. Hong et al.³⁹ proposed a method based on cGAN for semantic segmentation. The network integrated cGAN into the FCN framework to reduce the gap between source and target domains. In practical tasks, objects often appear in an occluded state, which brings great challenges to the perception tasks of autonomous systems. To solve this problem, Ehsani et al.³⁸ proposed SeGAN, which jointly generated the appearance and segmentation mask for invisible and visible regions of objects. Luo et al.¹⁵⁵ further considered a joint distribution at the category level that was different from global alignment strategies such as CycleGAN. They proposed a category-level adversarial network to enhance local semantic consistency in the case of global feature alignment. Note that traditional semantic segmentation methods

may suffer from the unsatisfactory quality of image-to-image conversion. Once the image-to-image conversion fails, nothing can be done to obtain satisfactory results in the subsequent stage of semantic segmentation. Li et al.¹⁵⁶ tackled this problem by introducing a bidirectional learning framework with self-supervised learning, in which both translation and segmentation adaptation models can promote each other in a closed loop. This segmentation adaptation model was trained on both synthetic and real-world datasets, which improved the segmentation performance of real-world data. In addition, Erkent and Laugier¹⁵⁷ considered a method of semantic segmentation adapted to different weather conditions, which can achieve satisfactory accuracy for semantic segmentation without the need of labeling the weather conditions of the source or target domain.

Semantic segmentation, as a high-level perception task of autonomous systems, predicts the semantic information of each pixel with a specific class label. Because early supervised algorithms are expensive in collecting labeled datasets from the real world, many algorithms consider the transferability between synthetic datasets and real-world data. Recent developments include semantic segmentation in more complex environments based on GANs, such as bad weather conditions. Meanwhile, we consider that instance segmentation based on GANs is also an open question. In addition to semantic segmentation, depth estimation is another high-level perception task of autonomous systems, whereby it is very challenging to estimate the depth value of each pixel in the image.

Depth Estimation

Depth estimation is an important task to help autonomous systems understand the 3D geometry of environments at high level. A series of classical and learning-based methods were proposed to estimate depth based on motion²⁴³ or stereo images,²⁴⁴ which is computationally expensive. As widely known, due to the lack of complete scene 3D information, estimating the depth from a single image is an ill-posed task.¹⁵⁸ For the monocular depth estimation task, a series of traditional DL-based algorithms without adversarial learning have been proposed to improve the accuracy of the model. However, considering that it is expensive to collect well-annotated datasets in depth estimation tasks, it is appealing to use adversarial learning methods, such as GANs, to achieve domain adaptation from synthetic datasets to real-world images. In addition, the adaptive method is used

to improve the transferability of the model, so that the model trained on the synthetic dataset can be well transferred to real-world images. Here, we introduce traditional DL-based depth estimation frameworks as well as methods to improve the transferability of depth estimation models by introducing adversarial learning.

Traditional DL-Based Depth Estimation. Traditional DL-based depth estimation methods mainly focus on improving the accuracy of the results by using deep convolution frameworks. Eigen et al.¹⁴ first proposed using a neural network to estimate depth from a single image in an end-to-end manner, which pioneeringly showed that it is promising for neural networks to estimate the depth from a single image. This framework consists of two components: the first one roughly estimated the global depth structure and the second one refined this global prediction using local information. Considering the continuous property of the monocular depth value, depth estimation is transformed into a learning problem of a continuous conditional random field (CRF). Liu et al.¹⁵⁸ presented a deep convolutional neural field model for single monocular depth estimation, which combined deep CNN and continuous CRF. This method achieved good results on both indoor and outdoor datasets. To reduce the dependence on the supervised signal and improve the transferability between different domains, unsupervised domain adaptation methods were presented for depth estimation by Nath Kundu et al.²⁴⁵ Some other developments in considering optical flow, camera pose, and intrinsic parameters from monocular video for depth estimation can be found in Chen et al.²⁴⁶ By considering the intrinsic parameters of the camera, similar to Gordon et al.,²⁴⁷ accurate depth information can be extracted from any video.

GANs-Based Depth Estimation. For the depth estimation task, it is too expensive to collect well-annotated image datasets. An appealing alternative is to use the unsupervised domain adaptation method via GANs to achieve domain adaptation from synthetic datasets to real-world images. Atapour-Abarghouei et al.²⁵ took advantage of the adversarial domain adaptation to train a depth estimation model in a synthetic city environment and transferred it to the real scene. The framework consists of two stages. At the first stage, a depth estimation model is trained with the dataset captured in the virtual environment. At the second stage, the proposed method transfers synthetic style images into real-world ones to reduce the domain discrepancy. Although this method considers the transfer of synthetic city environment to the real-world scene, it ignores the specific geometric structure of the image in the target domain, which is important for improving the accuracy of depth estimation. Motivated by this problem, Zhao et al.¹⁶⁰ proposed a geometry-aware symmetric domain adaptation network (GASDA), which produces high-quality results for both image style transfer and depth estimation. GASDA is based on CycleGAN,²⁴ which performs translations for both synthetic-realistic and realistic-synthetic simultaneously with a geometric consistency loss of real stereo images. Zhao et al.¹⁶¹ further considered high-level domain transformation, i.e., mixing a large number of synthetic images with a small amount of real-world images. They proposed the attend-remove-complete (ARC) method, which learns to attend, remove, and complete some challenging regions. The ARC method can ultimately make good use of synthetic data to generate accurate depth estimates.

Depth Estimation via Joint Tasks Learning. Each pixel in one image usually contains surface normal orientation vector information and semantic labels, and surface normal prediction, semantic segmentation, and depth estimation are related to the geometry of objects, which makes it possible to train different structured prediction tasks in a consistent manner. To the best of our knowledge, there are some works that apply a single model to multiple related tasks. Note that for different tasks, the model should be fine-tuned^{15,248} or use different loss functions,^{154,159} Applying a single model to multiple related tasks through fine-tuning or using different loss functions shows that the model has a good transferability. Eigen et al.¹⁵ developed a more general network for depth estimation and applied it to other computer vision tasks, such as surface normal estimation and per-pixel semantic labeling. It is worth noting that Eigen et al. used a single framework for depth estimation, surface normal estimation, and semantic segmentation with only fine-tuning, which improved the framework of their previous network¹⁴ by considering a third scale at a higher resolution. Considering that GANs perform well in structured prediction space, Hwang et al.¹⁵⁹ proposed adversarial structure matching (ASM), which trains a structured prediction network through an adversarial process. This method achieved ideal results on monocular depth estimation, semantic segmentation, and surface normal prediction. Although the ASM model has good transferability for multiple tasks through different loss functions, its limitation is that specified datasets should be used for specific tasks and cannot be generalized to other datasets. To solve this limitation, Chen et al.¹⁵⁴ embedded the pixel-level domain adaptation into the depth estimation task. Specifically, they proposed CrDoCo, a pixel-level adversarial domain-adaptive algorithm for dense prediction tasks. The core idea of this method is that although the image styles of two domains may be different during the domain-transfer process, the task predictions (e.g., depth estimation) should be exactly the same. Since CrDoCo is a pixel-level framework for dense prediction, it can be applied to semantic segmentation, depth prediction, and optical flow estimation, as shown in Figure 5. CrDoCo can be applied to multi-tasking only by adjusting its loss function, and it also shows a good transferability between different datasets for a specific task.

Depth estimation helps autonomous systems understand the 3D structure of the surrounding scene. The transferability of depth estimation includes not only the transfer of synthetic to real-world data but also the transfer of indoor to outdoor environments. Since depth, surface normals, and semantic labels are all related to object geometric information, recent works have also considered improving the accuracy of depth estimation by utilizing the interconnection between different tasks. We believe that future works should include the consideration of depth estimation under poor weather and light conditions. Following the autonomous systems perception tasks reviewed above, we now introduce pedestrian detection, re-ID, and tracking tasks involved in autonomous systems.

Pedestrian Detection, Re-identification, and Tracking

Pedestrian detection, re-ID, and tracking are very important for autonomous systems, especially for autonomous driving. The related works of pedestrian detection mainly focused on improving the accuracy of the results. Various developments have been made on improving the accuracy of complex visual

environments, such as nighttime and occlusion. Person re-ID, which is more complicated than pedestrian detection, requires matching pedestrians in disjointed camera views. Traditional learning-based methods of person re-ID mainly focused on improving the accuracy of results, while recent GANs-based algorithms focused on transferability between domains. To further complicate the person re-ID, some developments consider locating targets in a sequence of time, i.e., video tracking. The RL-based pedestrian tracking methods focus on not only accuracy but also transferability of the algorithms. In this subsection, we review pedestrian detection, re-ID, and tracking tasks, focusing on accuracy or transferability, or both.

Pedestrian Detection. In recent years, pedestrian detection has been widely taken into account in autonomous systems, especially for autonomous driving and robot movement.^{249,250} Pedestrian detection methods are generally divided into two categories: models based on hand-crafted features and deep models.¹⁶ Various models based on hand-crafted features have been proposed in the past few decades.^{251–253} Although these models have made good progress, models based on hand-crafted features fail to extract semantic information. Sermanet et al.¹⁶² used sparse convolutional feature hierarchies for pedestrian detection, which is named ConvNet. The network first performs layer-wise training on the whole multi-stage system, then uses the labeled data to fine-tune the complete architecture for the detection task. Although ConvNet learns features from training data, it treats pedestrian detection as a single binary classification task, which may confuse positive and negative samples. Therefore, Tian et al.¹⁶ proposed a task-assistant CNN (TA-CNN), which can learn features from multiple tasks and multiple datasets. TA-CNN combines semantic tasks, including pedestrian attributes and scene attributes, to optimize pedestrian detection results. To further improve the accuracy of pedestrian detection in natural scenes, Li et al.¹⁶³ considered that the problem of large variance in pedestrian scale with different spatial scales may cause dramatically different features. Therefore, they developed a scale-aware fast R-CNN (SAF R-CNN) framework, which combines a large-size subnetwork and a small-size subnetwork, as well as using the scale-aware weighting mechanism to deal with various sizes of pedestrian in scenes.¹⁶³ Although SAF R-CNN can detect pedestrian instances of different scales, it does not consider factors such as illumination conditions. To solve the problem of pedestrian detection under challenging illumination conditions at nighttime, Kim et al.⁴¹ used adversarial learning for cross-spectral pedestrian detection with unpaired setting. This method makes the color and thermal features of prominent areas where pedestrians exist to be similar by using adversarial learning, thereby improving the accuracy of pedestrian detection results at nighttime. Recent developments in pedestrian detection include tiny-scale pedestrian detection²⁵⁴ and occluded pedestrian detection.²⁵⁵

Person Re-identification. A similar while more difficult task than pedestrian detection, person re-ID requires matching pedestrians in disjointed camera views. At present, there are several learning-based methods focusing on person re-ID.^{17,256,257} However, these methods have poor transferability, i.e., the person re-ID models trained on one domain usually fail to generalize well to another domain. Considering that CycleGAN shows

impressive results in transferability using unpaired images, Deng et al.⁴² introduced the similarity preserving cycle-consistent generative adversarial network, an unsupervised domain adaptation approach to generate samples that not only have the target domain style but also preserve the underlying ID information. This method showed that applying domain adaptation to person re-ID can achieve competitive accuracy. Taking into account the data augmentation of different cameras, Zhong et al.¹⁶⁴ introduced the camera style (CamStyle) adaptation. CamStyle smoothes disparities in camera styles, transferring labeled training image styles to each camera to augment the training set. CamStyle helps to learn pedestrian descriptors through camera-invariant properties to improve re-ID experimental accuracy. The above approaches, such as SPGAN⁴² and CamStyle,¹⁶⁴ treated the domain gap as a black box and attempted to solve it by using a single style transformer. Liu et al.¹⁶⁵ proposed a novel adaptive transfer network (ATNet), which investigates the root causes of the domain gap. ATNet realizes the domain transfer of person re-ID by decomposing complicated cross-domain transfers and transferring features through sub-GANs separately. Recently, a theory-based analysis by Song et al.²⁵⁸ bridged the theoretical gap between unsupervised domain adaptation and re-ID tasks. Recent developments in person re-ID also involved considering occluded parts²⁵⁹ and different visual factors such as viewpoint, pose, illumination, and background.²⁶⁰

Pedestrian Tracking. Video tracking is an improvement in person re-ID and needs to locate the target in a sequence of time, which is difficult to handle because of tracking obstacles. When it comes to accuracy and transferability, RL-based methods in pedestrian tracking are concerned with whether the action in each frame is discrete or continuous and whether the labeled bounding boxes at each frame are limited or not. Tracking pedestrians by searching a series of discrete actions in each frame is a solution. Yun et al.¹⁶⁶ proposed the action-decision network (ADNet) to generate actions to find the location and the size of the target object in a new frame. The ADNet is updated by performing tracking simulation on the training sequence and utilizing action dynamics with the help of RL. After pre-training ADNet by supervised learning, online adaptation is applied to the network to accommodate the appearance changes or deformation of the target during tracking test sequences. Therefore, the pre-trained ADNet features can be transferred to a new frame during online adaptation. When the bounding boxes at each frame are limited, the algorithm can also be trained successfully and transferred to new frames with the help of larger training datasets. Supancic and Ramanan¹⁸ used RL to train trackers with more limited supervision on far more massive datasets. The results illustrated that the algorithm can track pedestrians on a never-before-seen video, and the video can be used for both evaluation of the current tracker and for training the tracker for future use. In brief, the learning structure is informative and the features contained in the video will be transferred to another training process. Furthermore, to exploit continuous actions for visual tracking, which improves training efficiency and accuracy, Chen et al.¹⁶⁷ introduced a real-time AC framework to exploit continuous action space for visual tracking. For online tracking, the “actor” model provides an offline dynamic search strategy to locate the target object in each frame efficiently by only one action output, and the “critic” model acts as a verification module to make the

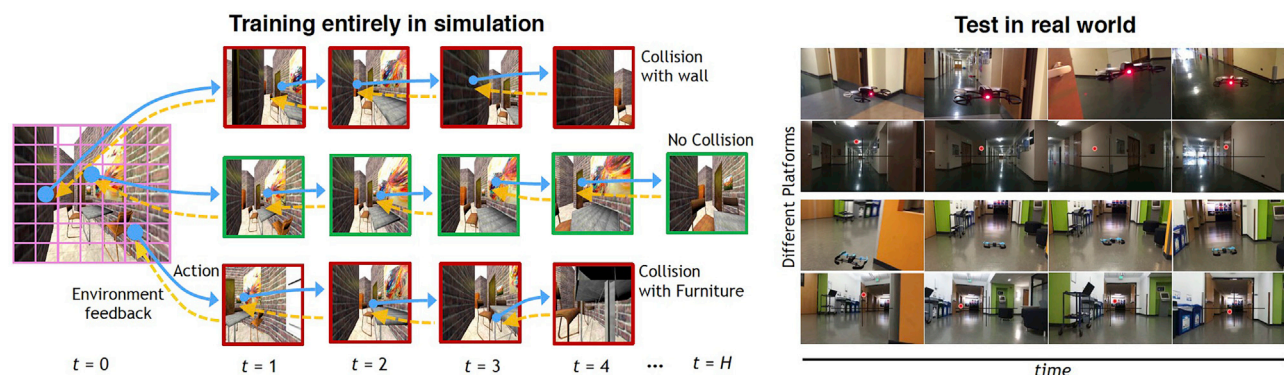


Figure 6. UAV Indoor Navigation via DRL Algorithm

DRL algorithm is entirely trained in a simulated 3D CAD model and generalized to real indoor flight environment. From Sadeghi and Levine.¹⁶⁹ See also Table 4.

tracker more robust. The real-time performance of the trackers is better than that of state-of-the-art methods such as MDNet²⁶¹ and ADNet.¹⁶⁶ Similar to Chen et al.,¹⁶⁷ Luo et al.⁵² used the same RL method to deal with continuous tracking problems. Moreover, they introduced an environment augmentation technique, i.e., virtual environments named ViZDoom,²⁶² to boost the tracker's generalization ability.

The current methods of pedestrian detection focus on improving the accuracy of the detection results, while the GANs-based person re-ID methods mentioned in this survey center on improving the transferability of the algorithm. The RL-based pedestrian tracking methods concentrate on equipping both the accuracy and transferability. Future works should include pedestrian detection, re-ID, and tracking in severe occlusion situations. In addition, future research may also involve pedestrian detection, re-ID, and tracking tasks at the semantic level. In addition to perception tasks, we now introduce some decision-making tasks, such as robot navigation. Robot navigation focuses on navigating a robot to avoid collision or to a target, considering accuracy or transferability, or both, of tasks.

Robot Navigation

Robot navigation, which has recently become crucial and a hot topic in autonomous systems, mainly focuses on navigating a robot to a target position or to avoid obstacles in a known/unknown environment. Here we consider whether the trained model can accurately learn the task features or successfully transfer the previous information to new tasks or domains. A variety of RL and meta-learning methods, such as DQN,¹⁷¹ LSTM structure,²⁶³ and MAML,¹²⁷ can accurately or transferably handle the changes arising from the environment or task when using the previously trained model. As shown in Table 4, we summarize the RL and meta-learning methods to handle accurate learning and domain-transfer tasks in robot and unmanned and autonomous vehicle (UAV) navigation issues. As for the experimental platform, AI2-THOR (the house of interactions)⁶¹ performs well due to its shared task features and datasets, ensuring the learned skills transfer to new tasks. Moreover, meta-learning methods usually have more satisfactory transferability than RL methods when lacking training and testing data by means of extracting or memorizing previous training data in simulation.

RL-Based Robot Navigation. To improve training efficiency and accuracy, dividing a single task to several subtasks and training them separately is a solution. Polvara et al.⁹⁹ proposed two distinct DQNs, called double DQNs, which were used to train two subtasks: landmark detection and vertical landing. Due to the separate training of each single task at the same time, training efficiency and accuracy were improved to an extent. Moreover, training the model with various auxiliary tasks, such as pixel control,²⁶⁴ reward prediction,²⁶⁵ and value function replay,¹⁰⁰ also helps the robot to adapt to the target faster and more accurately.

To equip the model with better transferability when encountering a new situation, task features^{53,171} and training policies⁴⁷ can be transferred to novel tasks in the same domain or across domains. Parisotto et al.⁵³ and Rusu et al.⁵⁴ transferred useful features among different ATARI games and then the corresponding features were utilized to train a new ATARI game in the same domain. In addition, when dealing with the tasks whose trials in the real world are usually time-consuming or expensive, the characteristic of tasks can be transferred across domains effectively. Zhang et al.¹⁷¹ proposed a shared DQN between tasks in order to learn informative features of tasks, which can be transferred from simulation to the real world. Similarly, as shown in Figure 6, Sadeghi and Levine¹⁶⁹ proposed a novel realistic translation network, which transforms virtual image inputs into real images with a similar scene structure. Moreover, policies can be transferred from simulation to simulation. Similar to Polvara et al.,⁹⁹ the primary training policy of Sadeghi and Levine¹⁶⁹ can be divided into several secondary policies, which acquire certain behaviors. These behaviors are then combined to train the primary policy, which helps to make the primary policy more transferable across domains. Chen et al.⁴⁷ used AC networks to train the secondary policies as well as the primary policy. In navigation, the primary behavior learned by a high-degree-freedom robot is to navigate straight to the target within a sample environment. Chen et al. then randomized the non-essential aspects of every secondary behavior, such as the appearance, the positions, and the number of obstacles in the scene, to improve generalization ability of the final policy.

Due to the sampling constraints of model-free RL methods and transferring limits of model-based RL methods as mentioned in Preliminaries, it is difficult to equip a model with

good transferability and sampling efficiency at the same time. An easy way to handle this contradiction is to combine model-free methods with model-based methods. Kahn et al.⁴⁸ used a generalized computation graph to find the navigation policies from scratch by inserting specific instantiations between model-free and model-based ones. Therefore, the algorithm not only learns high-dimensional tasks but also has promising sampling efficiency.

Meta-Learning-Based Robot Navigation. RL-based methods tend to need sufficient training data to acquire transferability. When a new task has insufficient data during training and testing, meta-learning methods can also promote the model to be transferable across domains. Firstly, recurrent models, such as LSTM structure, weaken the long-term dependency of sequential data, which acts as an optimizer to learn an optimization method for the gradient descent models. Mirowski et al.²⁶⁶ proposed a multi-city navigation network with LSTM structure. The main task of the LSTM structure was used to encode and encapsulate region-specific features and structures to add multiple paths in each city. After training in multiple cities, it was proved that the network is sufficiently versatile. Moreover, metric learning can be utilized to extract image information and generalize the specific information, which is helpful in navigation. Zhu et al.⁶¹ combined Siamese network with AC network to navigate the robot to the target only with 3D images. Siamese network captures and compares the special characteristics from the observation image and target image. The joint representation of images is then kept in scene-specific layers. AC network uses the features in scene-specific layers to generate policy and value outputs in navigation. To sum up, the deep Siamese AC network shares parameters across different tasks and domains so that the model can be generalized across targets and scenes. Even if the models trained by these two meta-learning methods acquire both accuracy and transferability, when the models encounter new cross-domain tasks they also need plenty of data to be re-trained. To fine-tune a new model with few data, MAML is beneficial. Finn et al.⁶⁰ verified that MAML performs well in 2D navigation and locomotion simulation compared with traditional policy gradient algorithms. It was shown that MAML could learn a model that adapts much more quickly in a single gradient update while continuing to improve with additional updates without overfitting. When the training process is unsupervised, MAML is not applicable and needs to be adjusted, such as by constructing a reward function during the meta-training process¹²⁹ and labeling data using clustering methods.²⁶⁷ Wortsman et al.¹⁷⁵ proposed a self-adaptive visual navigation (SAVN) method derived from MAML to learn adaptation to new environments without any supervision. Specifically, SAVN optimizes two objective functions: self-supervised interaction loss and navigation loss. During training, the interaction and navigation gradients are back-propagated through the network, and the parameters of the self-supervised loss are updated at the end of each episode using navigation gradients, which are trained by MAML. During testing, the parameters of the interaction loss remain fixed while the rest of the network is updated using interaction gradients. Therefore, the model equips the MAML methods with good transferability in a no-supervision environment.

RL or meta-learning methods help the robot navigate to targets or avoid obstacles. When using RL methods, separating

tasks or adding auxiliary tasks during the training process will improve the accuracy of the navigation results. Moreover, there are many ways to improve transferability in robot navigation, including task transfer, parameter transfer, and policy transfer. Compared with RL methods, meta-learning methods promote the transferability well, especially when the training and testing data are limited. In the future, with the popularity of MAML, we believe that MAML will become capable of handling more complex tasks in reality and achieving more satisfactory transferability by means of combining with state-of-art methods such as metric learning and LSTM structure. After outlining robot navigation tasks in autonomous systems, we now focus on another robotic issue, robotic manipulation.

Robotic Manipulation

We now focus on transferability in robotic manipulation issues according to domain-transfer tasks implemented by various robotic arms. Compared with robot navigation, robotic manipulation mainly considers precise control of robotic arms by means of multiple degrees of freedom. RL methods enable the robotic arm to transfer between different environments and tasks by means of special inputs²⁶⁸ and reformed training networks.²⁶⁹ Moreover, meta-learning and imitation learning can be utilized to handle difficult tasks with few or even one demonstration during the meta-testing process in the same domain or across domains^{115,183} in order to speed up the learning process and transfer previous task features. Table 4 summarizes RL and meta-learning methods to deal with domain-transfer robotic manipulation problems. As experimental platforms, MuJoCo (multi-joint dynamics with contact)²⁷⁰ and the PR2 arm are popular because robotic arms with multiple degrees of freedom and shared information have better accuracy and transferability. Moreover, compared with RL, meta-learning is capable of training with fewer data and adapting more quickly to new tasks to acquire model transferability.

RL-Based Robotic Manipulation. When considering improvement of the transferability of robotic arm systems, synthetic data as input¹⁷⁹ and separate networks in training⁹⁸ are possible RL-based solutions. Synthetic inputs help to transfer experience learned from different settings in simulation to the real world. Zhang et al.¹⁷⁹ were the first to produce a three-joint robot arm that learned control via DQN merely from raw-pixel images without any prior knowledge. The robot arm reaches the target in the real world successfully only when it takes synthetic images that are generated by the 2D simulator as inputs according to real-time joint angles. Therefore, the input of synthetic images inevitably offsets the gap between the simulation and real world, thereby improving the transferability. Moreover, when the data are limited and unable to be synthesized, DQN can be divided into perception and control modules, which are trained separately. The perception skills and the controller obtained from simulation will then be transferable.⁹⁸ Similarly, DQN can also train several networks and combine the experience learned together. Zeng et al.¹⁰¹ used DQN to jointly train two fully convolutional networks mapping from visual observations to actions. The experience transfers between robot pushing and grasping processes, and thus these synergies are learned. To compare some popular RL methods focusing on generalization ability in robotic manipulation, Quillen et al.²⁷¹ evaluated simulated benchmark tasks, whereby robot arms

were used to grasp random targets in comparison with some DRL algorithms, such as double Q-learning (DQL), DDPG, path consistency learning, and Monte Carlo (MC) policy evaluation. In the experiment, the trained robot arms coped with grasping unseen targets. The results revealed that DQL performs better than other algorithms in low-data regimes and has a relatively higher robustness to the choice of hyperparameters. When data are becoming plentiful, MC policy evaluation achieves a slightly better performance.

MAML-Based Robotic Manipulation. However, in robotic manipulation issues, traditional RL methods tend to need plenty of training data. Even if they can transfer to new tasks or domains, they also have poor generalization ability.^{180,272} MAML combined with imitation learning is able to utilize past experience across different tasks or domains, which can learn new skills from a very small number of demonstrations in various fields of application. Duan et al.⁶² let the robot arm demonstrate itself in simulation, i.e., the input and output samples were collected by the robot arm itself. The inputs of the model are the position information of each block rather than images or videos. They first sampled a demonstration from one of the training tasks, then sampled another pair of observation and action from a second demonstration of the same task. Considering both the first demonstration and second observation, the network was trained to output the corresponding action. In a manipulation network, the soft attention structure allows the model to generalize to conditions and tasks that are invisible in training data. Thus, Finn et al.¹¹⁵ used visual inputs from raw pixels as demonstration. The model requires data from significantly fewer prior demonstrations in training and merely one demonstration in testing to learn new skills effectively. Moreover, it not only performs well in simulation but also works in a real robotic arm system. MAML is modified to two-head architecture, which means that the algorithm is flexible for both learning to adapt policy parameters and learning the expert demonstration. Therefore, the number of demonstrations needed for an individual task is reduced by sharing the data across tasks. Taking robot arm pushing as an example, during the training process the robot arm can see various pushing demonstrations that contain different objects, and each object may have different quality and friction. In the testing process, the robot arm needs to push the object that it has never seen during training. It needs to learn which object to push and how to push it according to merely one demonstration. As shown in Figure 7, compared with Finn et al.,¹¹⁵ Yu et al.¹⁸³ increased the difficulty of imitation learning, i.e., only using a single video demonstration from a human as input, whereby the robot arm needs to accomplish the same work as in Finn et al.¹¹⁵ by domain adaptation. The authors proposed a domain-adaptive meta-learning method that transfers the data from human demonstrations to robot arm demonstrations. MAML was utilized to deal with the setting of learning from video demonstrations of humans. Due to the clone of behavior across the domain, the loss function also needs to be reconstructed, and TCN is used to construct the loss network in MAML structure in the robotic arm domain. Specifically, the robot arm will learn a set of initial parameters in the video domain, then after one or a few steps of gradient descent on merely one human demonstration the robot arm is able to perform the new task effectively. Recently, based on the study by Finn et al.¹¹⁵ Singh et al.¹⁸⁷

improved the one-shot imitation model by using additional autonomously collected data instead of manually collecting data. It is novel that they put forward an embedding network to distinguish whether two demonstration embeddings are close to each other. By the use of metric learning, they compute the Euclidean distance to find the distance between two videos. If they are close, it is regarded that the demonstrations falls into the same task. Therefore, the demonstrations from the same task are viewed as autonomously collected data that can be used to be trained in different tasks.

In robotic manipulation tasks, synthetic data as input and separate networks in training are RL-based ideas to equip robotic arms with transferability. Moreover, MAML with imitation learning methods do well in task and domain transfer with relatively few data. In the future, training with unlabeled data will be a trend, at which point autonomous systems need to label the training data by means of unsupervised methods. On the other hand, the testing demonstrations in meta-imitation learning will be much fewer, even with no demonstrations, so that an accurate and transferable model is needed.

Discussion and Future Works

This review shows the powerful effects of traditional DL, adversarial learning, RL, and meta-learning on complex visual and control tasks in autonomous systems. In particular, some traditional DL-based methods may not guarantee the accuracy when transferred to another domain; however, adversarial learning, RL, and meta-learning are able to treat transferability well. Although adversarial learning, such as GANs, produce better, clearer, and more transferable results than other traditional DL-based methods, meta-learning methods or combining them with RL and imitation learning methods tend to be equipped with an efficiency or transferability, or both of these.

Discussion

In this review, we introduce several typical perception and decision-making tasks of autonomous systems from the perspectives of accuracy and transferability. Since autonomous systems may have a better perception accuracy under good lighting environments than in harsh environments, we first introduce image style transfer, which can change the training data from night to day, rain to sunny, and so forth. Moreover, image style transfer can realize the transfer of synthetic datasets to real-world images, which greatly reduces the hardware loss caused by directly using autonomous systems for real scenes. We then review image enhancement and image restoration. Autonomous systems usually involve tasks such as image SR and image deblurring/dehazing/rain removal. We review recent developments from the perspectives of accuracy and transferability. When the image quality reaches a good perceptible state, we consider two typical high-level perception tasks of autonomous systems, i.e., semantic segmentation and depth estimation. Since obtaining the ground-truth labels is difficult with these two tasks, various methods have been proposed for the transferability between the synthetic datasets and real-world data. In addition, we review the tasks of pedestrian detection, person re-ID, and pedestrian tracking that are often involved in autonomous systems. Among them, pedestrian detection mainly aims to improve the accuracy of the detection results; person re-ID is a similar task

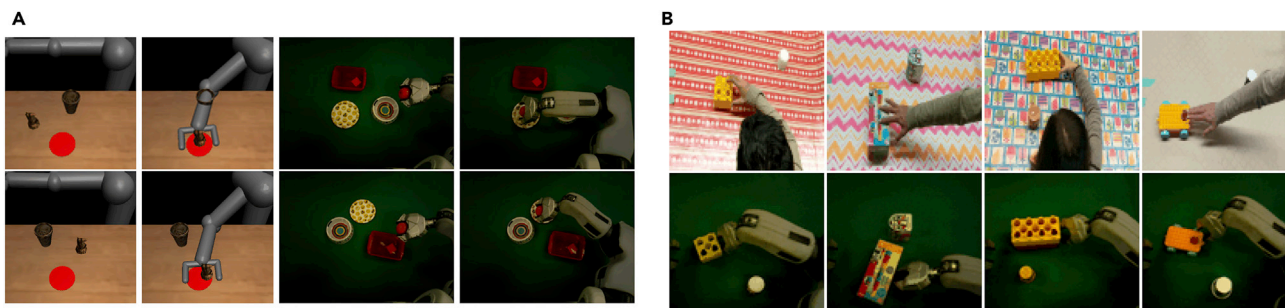


Figure 7. Demonstrations and Robotic Actions in Simulation and Real World

(A) Robot demonstrations used for meta-imitation learning. From Finn et al.¹¹⁵
(B) Human and robot demonstrations used for meta-imitation learning with large domain shift. From Yu et al.¹⁸³
See also Table 4.

but is more difficult, requiring the matching of pedestrians in disjointed camera views; pedestrian tracking pays attention to transferring network and video features and promoting framework accuracy. Furthermore, we consider two perception and decision-making tasks of robotic systems, i.e., robot navigation and robotic manipulation. Robot navigation tasks focus on accurately learning task features and transferring information across tasks or domains by RL or meta-learning. Robotic manipulation deals with domain-transfer tasks with more precise robotic arm control. These two tasks have simulation platforms or practice platforms, or both, to verify the accuracy or transferability.

Future Works

There are still important challenges and future works worthy of our attention. In this subsection, we summarize some trends and challenges for autonomous systems.

GANs with Good Stability, Quick Convergence, and Controllable Mode. GANs employ the gradient descent method to iterate the generator and discriminator to solve the minimax game problem. In the game, the mutual game between the generator and discriminator may cause model training to be unstable and difficult to converge, and even cause the mode to collapse. Although there are some preliminary studies aimed at improving these deficiencies of GANs,^{273,274} there is still much room for improvement in terms of the modal diversity and real-time performance. In addition, controlling the mode of data enhancement remains an open question. How to make the generated data mode controllable by controlling additional conditions and keep the model stable, and to achieve purposeful data enhancement, in particular for the computer vision tasks in autonomous systems, is an interesting research direction in the future.

GANs for Complex Multi-Tasking. Although GANs have achieved great results in some typical computer vision tasks of autonomous systems, it still remains difficult to consider the development of more complex multi-tasking in the future. Since some visual tasks are often related to each other, this phenomenon makes it possible to seamlessly reuse supervision between related tasks or solve different tasks in one system without adding complexity.²⁷⁵ For example, it is promising to consider training a general-purpose network that can be used for multi-task image restoration in a bad weather condition with only fine-tuning, such as image rain removal, snow removal, dehazing, seasonal change, and light adjustment. In addition, in severe

rain and foggy weather, how to perform image SR while removing rain/dehazing at the same time is challenging. In short, the use of GANs for more complex multi-tasking remains an open question and is worth exploring.

GANs for More Challenging Domain Adaptation. In autonomous systems, transferability is important for computer vision tasks. Although some results introduce GANs into domain adaptation to improve domain transfer,^{25,276} there is still much room for development. When considering more diverse domains, more differentiated cross-domains, and cross-style domains, such as road scenarios in different countries, the existing methods often cannot guarantee good transferability among these domains. However, GANs are promising for development of more diverse domain adaptations by showing unprecedented effectiveness in domain transfer. It is of interest to study the further use of GANs for more differentiated cross-domain transferability.

Multi-Modal, Multi-Task, and Multi-Agent RL. Most of the RL methods in applications focus primarily on visual input only. However, when considering information from multiple models, such as voice, text, and video, agents that can better understand the scenes and the performance in experiments will be more accurate and satisfactory.^{49,277} Moreover, in multi-task RL models, the agent is simultaneously trained on both auxiliary tasks and target tasks,^{264,265} so that the agent has the ability to transfer experience between tasks. Furthermore, thanks to the distributed nature of the multi-agent, multi-agent RL can achieve learning efficiency from sharing experience, such as communication, teaching, and imitation.²⁷⁸

Meta-Learning for Unsupervised Tasks. Traditional meta-learning consists of supervised learning during training and testing whereby both training data and testing data are labeled. However, if we use the unlabeled training data—in other words, there is no reward generated in training—how can we also achieve better results on specific tasks during testing? Leveraging unsupervised embeddings to automatically construct tasks or losses for unsupervised meta-learning is a solution,^{129,175,267} after which the training tasks for meta-learning are constructed. Therefore, meta-learning issues can be transformed into a wider unsupervised application. It will be interesting to use unsupervised meta-learning methods in more realistic task distributions so that the agent can explore and adapt to

new tasks more intelligently and the model can solve real-world tasks more effectively.

The Application Performance of RL and Meta-Learning. To deal with the differences between simulation environments and real scenes, the tasks or the networks can be transferred successfully using RL or meta-learning. Chances are that most of the existing algorithms with good performance in simulation cannot perform as well in the real world,²⁷⁹ which limits the applications of the models in simulation. Therefore, content-rich and flexible simulation frameworks, namely physics engines such as AI2-THOR,⁶¹ MuJoCo,²⁷⁰ or GymFC,¹⁷⁷ synthetic datasets such as SUNCG,¹⁷¹ and robot operating platforms such as V-REP (virtual robot experiment platform),²⁸⁰ will help to keep the learned information in more detail so that when transferred into the real world the performance is potentially good.^{115,183} In the future, more informative simulation environments and more flexible real-world platforms will shorten the gap between simulation and real world, thereby making the model more accurate and transferable. For example, a humanoid robotic hand with multiple degrees of freedom is able to deal with tasks more accurately,²⁸¹ 3D simulation involving shared tasks, simulators, and datasets ensures that the learned skills can be transferred successfully to reality.²⁸² In sum, due to the high similarity between simulation and real-world platforms, various high-complexity applications trained in simulation can be accurately transferred into practice.

Conclusion

In this review, we aim to contribute to the evolution of autonomous systems by exploring the impacts of accuracy or transferability, or both of them, on complex computer vision tasks and decision-making problems. To this end, we mainly focus on basic challenging perception and decision-making tasks in autonomous systems, such as image style transfer, image SR, image deblurring/dehazing/rain removal, semantic segmentation, depth estimation, pedestrian detection, person re-ID, pedestrian tracking, robot navigation, and robotic manipulation. We introduce some basic concepts and methods of transfer learning and its special case domain adaptation, then briefly discuss three typical adversarial learning networks, namely GANs, cGANs, and CycleGAN. We also present some basic concepts of RL, explain the idea of meta-learning, and discuss the relationship between adversarial learning, RL and meta-learning. Additionally, we analyze some typical DL methods and focus on the powerful performance of GANs in computer vision tasks, and discuss RL and meta-learning methods in robot control tasks in both simulation and the real world. Moreover, we provide summary tables of learning-based methods for different tasks in autonomous systems, which include the training manners, loss function of models, and experimental platforms in visual and robot control tasks. Finally, we discuss main challenges and future works from the aspects of perception and decision-making of autonomous systems by considering accuracy and transferability.

ACKNOWLEDGMENTS

The authors would like to thank the Editor-in-Chief, Scientific Editor, and anonymous referees for their helpful comments and suggestions, which have greatly improved this paper. This work was supported by the National Key Research and Development Program of China under grant 2018YFC0809302; the National Natural Science Foundation of China under

grant nos. 61988101, 61751305, and 61673176; and by the Program of Introducing Talents of Discipline to Universities (the 111 Project) under grant B17017.

AUTHOR CONTRIBUTIONS

Ideas design, Y.T. and C. Zhang; Reference Collection, Y.T., C. Zhang, and J.W.; Writing – Original Draft, C. Zhang and J.W.; Writing – Review & Editing, Y.T., G.G.Y., C. Zhao, Q.S., J.K, and F.Q.

REFERENCES

1. Silver, D., Schrittwieser, J., Simonyan, K., Antonoglou, I., Huang, A., Guez, A., Hubert, T., Baker, L., Lai, M., Bolton, A., et al. (2017). Mastering the game of go without human knowledge. *Nature* 550, 354–359.
2. Rahwan, I., Cebrian, M., Obradovich, N., Bongard, J., Bonnefon, J.-F., Breazeal, C., Crandall, J.W., Christakis, N.A., Couzin, I.D., Jackson, M.O., et al. (2019). Machine behaviour. *Nature* 568, 477–486.
3. LeCun, Y., Bengio, Y., and Hinton, G. (2015). Deep learning. *Nature* 521, 436–444.
4. Ferdowsi, A., Challita, U., and Saad, W. (2019). Deep learning for reliable mobile edge analytics in intelligent transportation systems: an overview. *IEEE Vehicular Technology Mag.* 14, 62–70.
5. McFarlane, D., Giannikas, V., and Lu, W. (2016). Intelligent logistics: involving the customer. *Comput. Industry* 87, 105–115.
6. J. Forlizzi and C. DiSalvo, (2006). Service robots in the domestic environment: a study of the roomba vacuum in the home. In *Proceedings of the 1st ACM SIGCHI/SIGART Conference on Human-robot Interaction*, pp. 258–265.
7. C. Dong, C.C. Loy, K. He, and X. Tang, “Learning a deep convolutional network for image super-resolution. In *European Conference on Computer Vision*, 2014, pp. 184–199.
8. Dong, C., Loy, C.C., He, K., and Tang, X. (2015). Image super-resolution using deep convolutional networks. *IEEE Trans. Pattern. Anal. Mach. Intell.* 38, 295–307.
9. J. Sun, W. Cao, Z. Xu, and J. Ponce, Learning a convolutional neural network for non-uniform motion blur removal. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2015, pp. 769–777.
10. Cai, B., Xu, X., Jia, K., Qing, C., and Tao, D. (2016). DehazeNet: an end-to-end system for single image haze removal. *IEEE Trans. Image Process.* 25, 5187–5198.
11. D. Eigen, D. Krishnan, and R. Fergus, “Restoring an image taken through a window covered with dirt or rain. In *Proceedings of the IEEE International Conference on Computer Vision*, 2013, pp. 633–640.
12. J. Long, E. Shelhamer, and T. Darrell, Fully convolutional networks for semantic segmentation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2015, pp. 3431–3440.
13. Badrinarayanan, V., Kendall, A., and Cipolla, R. (2017). Segnet: A deep convolutional encoder-decoder architecture for image segmentation. *IEEE Trans. Pattern. Anal. Mach. Intell.* 39, 2481–2495.
14. D. Eigen, C. Puhrsch, and R. Fergus, “Depth map prediction from a single image using a multi-scale deep network. In *Advances in Neural Information Processing Systems 27 (NIPS 2014)*, 2014, pp. 2366–2374.
15. D. Eigen and R. Fergus, “Predicting depth, surface normals and semantic labels with a common multi-scale convolutional architecture. In *Proceedings of the IEEE International Conference on Computer Vision*, 2015, pp. 2650–2658.
16. Y. Tian, P. Luo, X. Wang, and X. Tang, Pedestrian detection aided by deep learning semantic tasks. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2015, pp. 5079–5087.
17. M. Ye, A.J. Ma, L. Zheng, J. Li, and P.C. Yuen, Dynamic label graph matching for unsupervised video re-identification. In *Proceedings of*

- the IEEE International Conference on Computer Vision, 2017, pp. 5142–5150.
18. J. Supancic III and D. Ramanan, “Tracking as online decision-making: learning a policy from streaming videos with reinforcement learning. In Proceedings of the IEEE International Conference on Computer Vision, 2017, pp. 322–331.
19. Kober, J., Bagnell, J.A., and Peters, J. (2013). Reinforcement learning in robotics: a survey. *Int. J. Robotics Res.* 32, 1238–1274.
20. Polydoros, A.S., and Nalpantidis, L. (2017). Survey of model-based reinforcement learning: applications on robotics. *J. Intell. Robotic Syst.* 86, 153–173.
21. Gupta, A., Devin, C., Liu, Y., Abbeel, P., and Levine, S. (2017). Learning invariant feature spaces to transfer skills with reinforcement learning. *arXiv*, 1703.02949.
22. A. Faust, K. Oslund, O. Ramirez, A. Francis, L. Tapia, M. Fiser, and J. Davidson, “Prm-rl: Long-range robotic navigation tasks by combining reinforcement learning and sampling-based planning. In 2018 IEEE International Conference on Robotics and Automation (ICRA), 2018, pp. 5113–5120.
23. C. Tan, F. Sun, T. Kong, W. Zhang, C. Yang, and C. Liu, A survey on deep transfer learning. In International Conference on Artificial Neural Networks, 2018, pp. 270–279.
24. J.-Y. Zhu, T. Park, P. Isola, and A.A. Efros, Unpaired image-to-image translation using cycle-consistent adversarial networks. In Proceedings of the IEEE International Conference on Computer Vision, 2017, pp. 2223–2232.
25. A. Atapour-Abarghouei and T.P. Breckon, Real-time monocular depth estimation using synthetic data with domain adaptation via image style transfer. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 2800–2810.
26. C. Ledig, L. Theis, F. Huszár, J. Caballero, A. Cunningham, A. Acosta, A. Aitken, A. Tejani, J. Totz, Z. Wang *et al.*, “Photo-realistic single image super-resolution using a generative adversarial network. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2017, pp. 4681–4690.
27. O. Kupyn, V. Budzan, M. Mykhailych, D. Mishkin, and J. Matas, DeblurgAN: blind motion deblurring using conditional adversarial networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 8183–8192.
28. Gui, J., Sun, Z., Wen, Y., Tao, D., and Ye, J. (2020). A review on generative adversarial networks: algorithms, theory, and applications. *arXiv*, 2001.06937.
29. Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., Courville, A., and Bengio, Y. (2014). Generative adversarial nets. *Advances in Neural Information Processing Systems*, 2672–2680.
30. Goodfellow, I. (2016). NIPS 2016 tutorial: generative adversarial networks. *arXiv*, 1701.00160.
31. Csurka, G. (2017). Domain adaptation for visual applications: a comprehensive survey. *arXiv*, 1702.05374.
32. Reed, S., Akata, Z., Yan, X., Logeswaran, L., Schiele, B., and Lee, H. (2016). Generative adversarial text to image synthesis. *arXiv*, 1605.05396.
33. T. Xu, P. Zhang, Q. Huang, H. Zhang, Z. Gan, X. Huang, and X. He, “AttnGAN: fine-grained text to image generation with attentional generative adversarial networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 1316–1324.
34. 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, 2017, pp. 1125–1134.
35. R. Qian, R.T. Tan, W. Yang, J. Su, and J. Liu, Attentive generative adversarial network for raindrop removal from a single image. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 2482–2491.
36. Zhang, H., Sindagi, V., and Patel, V.M. (2019). Image de-raining using a conditional generative adversarial network. *IEEE Trans. Circuits Syst. Video Technol.*
37. J. Li, X. Liang, Y. Wei, T. Xu, J. Feng, and S. Yan, “Perceptual generative adversarial networks for small object detection. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2017, pp. 1222–1230.
38. K. Ehsani, R. Mottaghi, and A. Farhadi, SeGAN: segmenting and generating the invisible. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 6144–6153.
39. W. Hong, Z. Wang, M. Yang, and J. Yuan, Conditional generative adversarial network for structured domain adaptation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 1335–1344.
40. S. Sankaranarayanan, Y. Balaji, A. Jain, S. Nam Lim, and R. Chellappa, “Learning from synthetic data: addressing domain shift for semantic segmentation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 3752–3761.
41. M. Kim, S. Joung, K. Park, S. Kim, and K. Sohn, “Unpaired cross-spectral pedestrian detection via adversarial feature learning. In 2019 IEEE International Conference on Image Processing (ICIP), 2019, pp. 1650–1654.
42. W. Deng, L. Zheng, Q. Ye, G. Kang, Y. Yang, and J. Jiao, “Image-image domain adaptation with preserved self-similarity and domain-dissimilarity for person re-identification. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 994–1003.
43. C. Vondrick, H. Pirsiavash, and A. Torralba, “Generating videos with scene dynamics. In Advances in Neural Information Processing Systems, 2016, pp. 613–621.
44. Jaradat, M.A.K., Al-Rousan, M., and Quadan, L. (2011). Reinforcement based mobile robot navigation in dynamic environment. *Robotics and Computer-Integrated Manufacturing* 27, 135–149.
45. N. Kohl and P. Stone, “Policy gradient reinforcement learning for fast quadrupedal locomotion. In IEEE International Conference on Robotics and Automation, 2004. Proceedings. ICRA’04. 2004, vol. 3, 2004, pp. 2619–2624.
46. Xie, L., Wang, S., Markham, A., and Trigoni, N. (2017). Towards monocular vision based obstacle avoidance through deep reinforcement learning. *arXiv*, 1706.09829.
47. X. Chen, A. Ghadirzadeh, J. Folkesson, M. Björkman, and P. Jensfelt, Deep reinforcement learning to acquire navigation skills for wheel-legged robots in complex environments. In 2018 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS), 2018, pp. 3110–3116.
48. G. Kahn, A. Villafior, B. Ding, P. Abbeel, and S. Levine, Self-supervised deep reinforcement learning with generalized computation graphs for robot navigation. In 2018 IEEE International Conference on Robotics and Automation (ICRA), 2018, pp. 1–8.
49. P. Anderson, Q. Wu, D. Teney, J. Bruce, M. Johnson, N. Sünderhauf, I. Reid, S. Gould, and A. van den Hengel, Vision-and-language navigation: interpreting visually-grounded navigation instructions in real environments. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 3674–3683.
50. Levine, S., Finn, C., Darrell, T., and Abbeel, P. (2016). End-to-end training of deep visuomotor policies. *J. Machine Learn. Res.* 17, 1334–1373.
51. V. Mnih, A.P. Badia, M. Mirza, A. Graves, T. Lillicrap, T. Harley, D. Silver, and K. Kavukcuoglu, Asynchronous methods for deep reinforcement learning. In International Conference on Machine Learning, 2016, pp. 1928–1937.
52. Luo, W., Sun, P., Zhong, F., Liu, W., Zhang, T., and Wang, Y. (2020). End-to-end active object tracking and its real-world deployment via reinforcement learning. *IEEE Trans. Pattern. Anal. Mach. Intell.* 42, 1317–1332.
53. Parisotto, E., Ba, J.L., and Salakhutdinov, R. (2015). Actor-mimic: deep multitask and transfer reinforcement learning. *arXiv*, 1511.06342.
54. Rusu, A.A., Colmenarejo, S.G., Gulcehre, C., Desjardins, G., Kirkpatrick, J., Pascanu, R., Mnih, V., Kavukcuoglu, K., and Hadsell, R. (2015). Policy distillation. *arXiv*, 1511.06295.

55. Olivecrona, M., Blaschke, T., Engkvist, O., and Chen, H. (2017). Molecular de-novo design through deep reinforcement learning. *J. Cheminformatics* 9, 48.
56. Popova, M., Isayev, O., and Tropsha, A. (2018). Deep reinforcement learning for de novo drug design. *Sci. Adv.* 4, eaap7885.
57. Botvinick, M., Ritter, S., Wang, J.X., Kurth-Nelson, Z., Blundell, C., and Hassabis, D. (2019). Reinforcement learning, fast and slow. *Trends Cogn. Sci.* 23, 408–422.
58. Vilalta, R., and Drissi, Y. (2002). A perspective view and survey of meta-learning. *Artif. Intelligence Rev.* 18, 77–95.
59. Koch, G., Zemel, R., and Salakhutdinov, R. (2015). Siamese neural networks for one-shot image recognition. *ICML Deep Learning Workshop* 2.
60. C. Finn, P. Abbeel, and S. Levine, “Model-agnostic meta-learning for fast adaptation of deep networks. In *Proceedings of the 34th International Conference on Machine Learning Volume 70*, 2017, pp. 1126–1135.
61. Y. Zhu, R. Mottaghi, E. Kolve, J.J. Lim, A. Gupta, L. Fei-Fei, and A. Farhadi, Target-driven visual navigation in indoor scenes using deep reinforcement learning. In *2017 IEEE International Conference on Robotics and Automation (ICRA)*, 2017, pp. 3357–3364.
62. Duan, Y., Andrychowicz, M., Stadie, B., Ho, O.J., Schneider, J., Sutskever, I., Abbeel, P., and Zaremba, W. (2017). One-shot imitation learning. *Advances in Neural Information Processing Systems*, 1087–1098.
63. Tang, Y., Zhao, C., Wang, J., Zhang, C., Sun, Q., Zheng, W., Du, W., Qian, F., and Kurths, J. (2020). An overview of perception and decision-making in autonomous systems in the era of learning. *arXiv*, 2001.02319.
64. Arulkumaran, K., Deisenroth, M.P., Brundage, M., and Bharath, A.A. (2017). A brief survey of deep reinforcement learning. *arXiv*, 1708.05866.
65. Weiss, K., Khoshgoftaar, T.M., and Wang, D. (2016). A survey of transfer learning. *J. Big Data* 3, 9.
66. Pan, S.J., and Yang, Q. (2009). A survey on transfer learning. *IEEE Trans. Knowledge Data Eng.* 22, 1345–1359.
67. Pan, S.J., Tsang, I.W., Kwok, J.T., and Yang, Q. (2010). Domain adaptation via transfer component analysis. *IEEE Trans. Neural Networks* 22, 199–210.
68. T. Evgeniou and M. Pontil, Regularized multi-task learning. In *Proceedings of the Tenth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2004, pp. 109–117.
69. R. Raina, A. Battle, H. Lee, B. Packer, and A.Y. Ng, Self-taught learning: transfer learning from unlabeled data. In *Proceedings of the 24th International Conference on Machine learning*, 2007, pp. 759–766.
70. Ganin, Y., and Lempitsky, V. (2014). Unsupervised domain adaptation by backpropagation. *arXiv*, 1409.7495.
71. Hoffman, J., Wang, D., Yu, F., and Darrell, T. (2016). FCNs in the wild: pixel-level adversarial and constraint-based adaptation. *arXiv*, 1612.02649.
72. M. Wulfmeier, A. Bewley, and I. Posner, Addressing appearance change in outdoor robotics with adversarial domain adaptation. In *2017 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*, 2017, pp. 1551–1558.
73. S. Bak, P. Carr, and J.-F. Lalonde, “Domain adaptation through synthesis for unsupervised person re-identification. In *Proceedings of the European Conference on Computer Vision (ECCV)*, 2018, pp. 189–205.
74. Hoffman, J., Tzeng, E., Park, T., Zhu, J.-Y., Isola, P., Saenko, K., Efros, A.A., and Darrell, T. (2017). CyCADA: cycle-consistent adversarial domain adaptation. *arXiv*, 1711.03213.
75. Y. Zhu, Y. Chen, Z. Lu, S.J. Pan, G.-R. Xue, Y. Yu, and Q. Yang, Heterogeneous transfer learning for image classification. In *Twenty-Fifth AAAI Conference on Artificial Intelligence*, 2011.
76. Yang, L., Jing, L., Yu, J., and Ng, M.K. (2015). Learning transferred weights from co-occurrence data for heterogeneous transfer learning. *IEEE Trans. Neural Networks Learn. Syst.* 27, 2187–2200.
77. Patel, V.M., Gopalan, R., Li, R., and Chellappa, R. (2015). Visual domain adaptation: a survey of recent advances. *IEEE Signal Process. Mag.* 32, 53–69.
78. Chen, M., Xu, Z., Weinberger, K., and Sha, F. (2012). Marginalized denoising autoencoders for domain adaptation. *arXiv*, 1206.4683.
79. Long, M., Cao, Y., Wang, J., and Jordan, M.I. (2015). Learning transferable features with deep adaptation networks. *arXiv*, 1502.02791.
80. E. Tzeng, J. Hoffman, K. Saenko, and T. Darrell, “Adversarial discriminative domain adaptation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 7167–7176.
81. Ganin, Y., Ustinova, E., Ajakan, H., Germain, P., Larochelle, H., Laviolette, F., Marchand, M., and Lempitsky, V. (2016). Domain-adversarial training of neural networks. *J. Machine Learn. Res.* 17, 2096–3030.
82. Q. Chen, Y. Liu, Z. Wang, I. Wassell, and K. Chetty, “Re-weighted adversarial adaptation network for unsupervised domain adaptation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2018, pp. 7976–7985.
83. N. Dalvi, P. Domingos, S. Sanghai, and D. Verma, “Adversarial classification. In *Proceedings of the Tenth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2004, pp. 99–108.
84. M. Großhans, C. Sawade, M. Brückner, and T. Scheffer, Bayesian games for adversarial regression problems. In *International Conference on Machine Learning*, 2013, pp. 55–63.
85. Brückner, M., Kanzow, C., and Scheffer, T. (2012). Static prediction games for adversarial learning problems. *J. Machine Learn. Res.* 13, 2617–2654.
86. S. Mei and X. Zhu, Using machine teaching to identify optimal training-set attacks on machine learners. In *Twenty-Ninth AAAI Conference on Artificial Intelligence*, 2015.
87. Dasgupta, P., Collins, J.B., and McCarrick, M. (2020). Playing to learn better: repeated games for adversarial learning with multiple classifiers. *arXiv*, 2002.03924.
88. Creswell, A., White, T., Dumoulin, V., Arulkumaran, K., Sengupta, B., and Bharath, A.A. (2018). Generative adversarial networks: an overview. *IEEE Signal Process. Mag.* 35, 53–65.
89. Mirza, M., and Osindero, S. (2014). Conditional generative adversarial nets. *arXiv*, 1411.1784.
90. Kaelbling, L.P., Littman, M.L., and Moore, A.W. (1996). Reinforcement learning: a survey. *J. Artif. Intelligence Res.* 4, 237–285.
91. Sutton, R.S., and Barto, A.G. (2018). *Reinforcement Learning: An Introduction* (MIT Press).
92. Geffner, H. (2018). Model-free, model-based, and general intelligence. *arXiv*, 1806.02308.
93. Feinberg, V., Wan, A., Stoica, I., Jordan, M.I., Gonzalez, J.E., and Levine, S. (2018). Model-based value estimation for efficient model-free reinforcement learning. *arXiv*, 1803.00101.
94. Gläscher, J., Daw, N., Dayan, P., and O’Doherty, J.P. (2010). States versus rewards: dissociable neural prediction error signals underlying model-based and model-free reinforcement learning. *Neuron* 66, 585–595.
95. Mnih, V., Kavukcuoglu, K., Silver, D., Graves, A., Antonoglou, I., Wierstra, D., and Riedmiller, M. (2013). Playing Atari with deep reinforcement learning. *arXiv*, 1312.5602.
96. Lillicrap, T.P., Hunt, J.J., Pritzel, A., Heess, N., Erez, T., Tassa, Y., Silver, D., and Wierstra, D. (2015). Continuous control with deep reinforcement learning. *arXiv*, 1509.02971.
97. Azar, M.G., Gómez, V., and Kappen, H.J. (2012). Dynamic policy programming. *J. Machine Learn. Res.* 13, 3207–3245.
98. Zhang, F., Leitner, J., Milford, M., and Corke, P. (2016). Modular deep Q networks for sim-to-real transfer of visuo-motor policies. *arXiv*, 1610.06781.

99. Polvara, R., Patacchiola, M., Sharma, S., Wan, J., Manning, A., Sutton, R., and Cangelosi, A. (2017). Autonomous quadrotor landing using deep reinforcement learning. *arXiv*, 1709.03339.
100. Mirowski, P., Pascanu, R., Viola, F., Soyer, H., Ballard, A.J., Banino, A., Denil, M., Goroshin, R., Sifre, L., Kavukcuoglu, K., et al. (2016). Learning to navigate in complex environments. *arXiv*, 1611.03673.
101. A. Zeng, S. Song, S. Welker, J. Lee, A. Rodriguez, and T. Funkhouser, Learning synergies between pushing and grasping with self-supervised deep reinforcement learning. In 2018 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS), 2018, pp. 4238–4245.
102. Vanschoren, J. (2018). Meta-learning: a survey. *arXiv*, 1810.03548.
103. Li, Z., Zhou, F., Chen, F., and Li, H. (2017). Meta-SGD: learning to learn quickly for few-shot learning. *arXiv*, 1707.09835.
104. Hochreiter, S., and Schmidhuber, J. (1997). Long short-term memory. *Neural Comput.* 9, 1735–1780.
105. C. Lea, M.D. Flynn, R. Vidal, A. Reiter, and G.D. Hager, “Temporal convolutional networks for action segmentation and detection. In proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2017, pp. 156–165.
106. Kulis, B. (2013). Metric learning: a survey. *Foundations Trends Mach. Learn.* 5, 287–364.
107. Snell, J., Swersky, K., and Zemel, R. (2017). Prototypical networks for few-shot learning. *Advances in Neural Information Processing Systems*, 4077–4087.
108. Wang, Y., Kwok, J., Ni, L., and Yao, Q. (2019). Generalizing from a few examples: a survey on few-shot learning. *arXiv*, 1904.05046.
109. S. Chopra, R. Hadsell, and Y. LeCun, Learning a similarity metric discriminatively, with application to face verification. In 2005 IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR’05), vol. 1, 2005, pp. 539–546.
110. Vinyals, O., Blundell, C., Lillicrap, T., Kavukcuoglu, K., and Wierstra, D. (2016). Matching networks for one shot learning. *Advances in Neural Information Processing Systems*, 3630–3638.
111. F. Sung, Y. Yang, L. Zhang, T. Xiang, P.H. Torr, and T.M. Hospedales, “Learning to compare: relation network for few-shot learning. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 1199–1208.
112. Lee, Y., and Choi, S. (2018). Gradient-based meta-learning with learned layerwise metric and subspace. *arXiv*, 1801.05558.
113. Rakelly, K., Zhou, A., Quillen, D., Finn, C., and Levine, S. (2019). Efficient off-policy meta-reinforcement learning via probabilistic context variables. *arXiv*, 1903.08254.
114. Bottou, L. (2010). Large-scale machine learning with stochastic gradient descent. In Proceedings of COMPSTAT’2010, Y. Lechevallier and G. Saporta, eds. (Springer), pp. 177–186.
115. Finn, C., Yu, T., Zhang, T., Abbeel, P., and Levine, S. (2017). One-shot visual imitation learning via meta-learning. *arXiv*, 1709.04905.
116. Sutton, R.S., and Barto, A.G. (1998). *Reinforcement Learning: An Introduction* (MIT Press).
117. Pfau, D., and Vinyals, O. (2016). Connecting generative adversarial networks and actor-critic methods. *arXiv*, 1610.01945.
118. V.R. Konda and J.N. Tsitsiklis, Actor-critic algorithms. In *Advances in Neural Information Processing Systems*, 2000, pp. 1008–1014.
119. M. Sarmad, H.J. Lee, and Y.M. Kim, “RL-GAN-Net: a reinforcement learning agent controlled gan network for real-time point cloud shape completion. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 5898–5907.
120. Ganin, Y., Kulkarni, T., Babuschkin, I., Eslami, S., and Vinyals, O. (2018). Synthesizing programs for images using reinforced adversarial learning. *arXiv*, 1804.01118.
121. Ng, A.Y., and Russell, S.J. (2000). Algorithms for inverse reinforcement learning. *ICML ’00: Proceedings of the 17th International Conference on Machine Learning*, 663–670.
122. Ho, J., and Ermon, S. (2016). Generative adversarial imitation learning. *Advances in Neural Information Processing Systems*, 4565–4573.
123. Li, Y., Song, J., and Ermon, S. (2017). InfoGAIL: interpretable imitation learning from visual demonstrations. *Advances in Neural Information Processing Systems*, 3812–3822.
124. S. Hochreiter, A.S. Younger, and P.R. Conwell, “Learning to learn using gradient descent. In International Conference on Artificial Neural Networks, 2001, pp. 87–94.
125. Wang, J.X., Kurth-Nelson, Z., Tirumala, D., Soyer, H., Leibo, J.Z., Munos, R., Blundell, C., Kumaran, D., and Botvinick, M. (2016). Learning to reinforcement learn. *arXiv*, 1611.05763.
126. Nichol, A., and Schulman, J. (2018). Reptile: a scalable metalearning algorithm. *arXiv*, 1803.02999 <https://openai.com/blog/reptile/>.
127. Gupta, A., Mendonca, R., Liu, Y., Abbeel, P., and Levine, S. (2018). Meta-reinforcement learning of structured exploration strategies. *Advances in Neural Information Processing Systems*, 5302–5311.
128. Houthoofd, R., Chen, Y., Isola, P., Stadie, B., Wolski, F., Ho, O.J., and Abbeel, P. (2018). Evolved policy gradients. *Advances in Neural Information Processing Systems*, 5400–5409.
129. Gupta, A., Eysenbach, B., Finn, C., and Levine, S. (2018). Unsupervised meta-learning for reinforcement learning. *arXiv*, 1806.04640.
130. L.A. Gatys, A.S. Ecker, and M. Bethge, “Image style transfer using convolutional neural networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2016, pp. 2414–2423.
131. J. Johnson, A. Alahi, and L. Fei-Fei, Perceptual losses for real-time style transfer and super-resolution. In European Conference on Computer Vision, 2016, pp. 694–711.
132. Li, Y., Wang, N., Liu, J., and Hou, X. (2017). Demystifying neural style transfer. *arXiv*, 1701.01036.
133. R. Gong, W. Li, Y. Chen, and L.V. Gool, DLOW: domain flow for adaptation and generalization. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 2477–2486.
134. Z. Shen, M. Huang, J. Shi, X. Xue, and T.S. Huang, “Towards instance-level image-to-image translation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 3683–3692.
135. C. Dong, C.C. Loy, and X. Tang, “Accelerating the super-resolution convolutional neural network. In European Conference on Computer Vision, 2016, pp. 391–407.
136. M.S. Sajjadi, B. Scholkopf, and M. Hirsch EnhanceNet: single image super-resolution through automated texture synthesis. In Proceedings of the IEEE International Conference on Computer Vision, 2017, pp. 4491–4500.
137. A. Shocher, N. Cohen, and M. Irani, “Zero-shot” super-resolution using deep internal learning. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 3118–3126.
138. X. Wang, K. Yu, S. Wu, J. Gu, Y. Liu, C. Dong, Y. Qiao, and C. Change Loy, “Esrgan: Enhanced super-resolution generative adversarial networks. In Proceedings of the European Conference on Computer Vision (ECCV), 2018, pp. 0–0.
139. 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, 2018, pp. 701–710.
140. J.W. Soh, G.Y. Park, J. Jo, and N.I. Cho, Natural and realistic single image super-resolution with explicit natural manifold discrimination. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 8122–8131.
141. Gong, D., Sun, W., Shi, Q., Hengel, A.v. d., and Zhang, Y. (2020). Learning to zoom-in via learning to zoom-out: real-world super-resolution by generating and adapting degradation. *arXiv*, 2001.02381.

142. O. Kupyn, T. Martyniuk, J. Wu, and Z. Wang, DeblurGAN-v2: deblurring (orders-of-magnitude) faster and better. In Proceedings of the IEEE International Conference on Computer Vision, 2019, pp. 8878–8887.
143. R. Aljadaany, D.K. Pal, and M. Savvides, “Douglas-Rachford networks: learning both the image prior and data fidelity terms for blind image deconvolution. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 10 235–10 244.
144. R. Li, J. Pan, Z. Li, and J. Tang, “Single image dehazing via conditional generative adversarial network. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 8202–8211.
145. D. Engin, A. Genç, and H. Kemal Ekenel, “Cycle-dehaze: enhanced CycleGAN for single image dehazing. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops, 2018, pp. 825–833.
146. G. Kim, J. Park, S. Ha, and J. Kwon, “Bidirectional deep residual learning for haze removal. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops, 2019, pp. 46–54.
147. A. Dudhane and S. Murala, CDNet: single image de-hazing using unpaired adversarial training. In 2019 IEEE Winter Conference on Applications of Computer Vision (WACV), 2019, pp. 1147–1155.
148. P. Sharma, P. Jain, and A. Sur, Scale-aware conditional generative adversarial network for image dehazing. In The IEEE Winter Conference on Applications of Computer Vision, 2020, pp. 2355–2365.
149. R. Li, L.-F. Cheong, and R.T. Tan, “Heavy rain image restoration: integrating physics model and conditional adversarial learning. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 1633–1642.
150. Jin, X., Chen, Z., and Li, W. (2020). AI-GAN: asynchronous interactive generative adversarial network for single image rain removal. *Pattern Recogn.* 100, 107143.
151. K. He, G. Gkioxari, P. Dollár, and R. Girshick, “Mask r-CNN. In Proceedings of the IEEE International Conference on Computer Vision, 2017, pp. 2961–2969.
152. Y. Zhang, Z. Qiu, T. Yao, D. Liu, and T. Mei, Fully convolutional adaptation networks for semantic segmentation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 6810–6818.
153. R. Hu, P. Dollár, K. He, T. Darrell, and R. Girshick, Learning to segment every thing. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 4233–4241.
154. Y.-C. Chen, Y.-Y. Lin, M.-H. Yang, and J.-B. Huang, “CrDoCo: pixel-level domain transfer with cross-domain consistency. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 1791–1800.
155. Y. Luo, L. Zheng, T. Guan, J. Yu, and Y. Yang, “Taking a closer look at domain shift: category-level adversaries for semantics consistent domain adaptation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 2507–2516.
156. Y. Li, L. Yuan, and N. Vasconcelos, Bidirectional learning for domain adaptation of semantic segmentation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 6936–6945.
157. Erkent, Ö., and Laugier, C. (2020). Semantic segmentation with unsupervised domain adaptation under varying weather conditions for autonomous vehicles. *IEEE Robotics Automation Lett.* 1–8.
158. Liu, F., Shen, C., Lin, G., and Reid, I. (2015). Learning depth from single monocular images using deep convolutional neural fields. *IEEE Trans. Pattern. Anal. Mach. Intell.* 38, 2024–2039.
159. J.-J. Hwang, T.-W. Ke, J. Shi, and S.X. Yu, “Adversarial structure matching for structured prediction tasks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 4056–4065.
160. S. Zhao, H. Fu, M. Gong, and D. Tao, “Geometry-aware symmetric domain adaptation for monocular depth estimation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 9788–9798.
161. Zhao, Y., Kong, S., Shin, D., and Fowlkes, C. (2020). Domain decluttering: simplifying images to mitigate synthetic-real domain shift and improve depth estimation. *arXiv*, 2002.12114.
162. P. Sermanet, K. Kavukcuoglu, S. Chintala, and Y. LeCun, Pedestrian detection with unsupervised multi-stage feature learning. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2013, pp. 3626–3633.
163. Li, J., Liang, X., Shen, S., Xu, T., Feng, J., and Yan, S. (2017). Scale-aware fast r-CNN for pedestrian detection. *IEEE Trans. Multimedia* 20, 985–996.
164. Z. Zhong, L. Zheng, Z. Zheng, S. Li, and Y. Yang, “Camera style adaptation for person re-identification. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 5157–5166.
165. J. Liu, Z.-J. Zha, D. Chen, R. Hong, and M. Wang, “Adaptive transfer network for cross-domain person re-identification. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 7202–7211.
166. S. Yun, J. Choi, Y. Yoo, K. Yun, and J. Young Choi, Action-decision networks for visual tracking with deep reinforcement learning. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2017, pp. 2711–2720.
167. B. Chen, D. Wang, P. Li, S. Wang, and H. Lu, Real-time ‘actor-critic’ tracking. In Proceedings of the European Conference on Computer Vision (ECCV), 2018, pp. 318–334.
168. Schwenker, F., and Trentin, E. (2014). Pattern classification and clustering: a review of partially supervised learning approaches. *Pattern Recognition Lett.* 37, 4–14.
169. Sadeghi, F., and Levine, S. (2016). CAD2RL: real single-image flight without a single real image. *arXiv*, 1611.04201.
170. L. Tai, G. Paolo, and M. Liu, Virtual-to-real deep reinforcement learning: continuous control of mobile robots for mapless navigation. In 2017 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS), 2017, pp. 31–36.
171. J. Zhang, J.T. Springenberg, J. Boedecker, and W. Burgard, “Deep reinforcement learning with successor features for navigation across similar environments. In 2017 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS), 2017, pp. 2371–2378.
172. Banino, A., Barry, C., Uribe, B., Blundell, C., Lillicrap, T., Mirowski, P., Pritzel, A., Chadwick, M.J., Degris, T., Modayil, J., et al. (2018). Vector-based navigation using grid-like representations in artificial agents. *Nature* 557, 429–433.
173. F. Zhu, L. Zhu, and Y. Yang, “Sim-real joint reinforcement transfer for 3D indoor navigation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 11 388–11 397.
174. Niroui, F., Zhang, K., Kashino, Z., and Nejat, G. (2019). Deep reinforcement learning robot for search and rescue applications: exploration in unknown cluttered environments. *IEEE Robotics Automation Lett.* 4, 610–617.
175. M. Wortsman, K. Ehsani, M. Rastegari, A. Farhadi, and R. Mottaghi, “Learning to learn how to learn: self-adaptive visual navigation using meta-learning. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 6750–6759.
176. Jabri, A., Hsu, K., Gupta, A., Eysenbach, B., Levine, S., and Finn, C. (2019). Unsupervised curricula for visual meta-reinforcement learning. *Advances in Neural Information Processing Systems (NeurIPS 2019)* <http://papers.nips.cc/paper/9238-unsupervised-curricula-for-visual-meta-reinforcement-learning>.
177. Koch, W., Mancuso, R., West, R., and Bestavros, A. (2019). Reinforcement learning for UAV attitude control. *ACM Trans. Cyber-Physical Syst.* 3, 1–21.
178. Gaudet, B., Linares, R., and Furfaro, R. (2020). Adaptive guidance and integrated navigation with reinforcement meta-learning. *Acta Astronautica*. <https://doi.org/10.13140/RG.2.2.24778.52164>.
179. Zhang, F., Leitner, J., Milford, M., Upcroft, B., and Corke, P. (2015). Towards vision-based deep reinforcement learning for robotic motion control. *arXiv*, 1511.03791.

180. S. Gu, E. Holly, T. Lillicrap, and S. Levine, "Deep reinforcement learning for robotic manipulation with asynchronous off-policy updates. In 2017 IEEE International Conference on Robotics and Automation (ICRA), 2017, pp. 3389–3396.
181. T. Haarnoja, V. Pong, A. Zhou, M. Dalal, P. Abbeel, and S. Levine, Composable deep reinforcement learning for robotic manipulation. In 2018 IEEE International Conference on Robotics and Automation (ICRA), 2018, pp. 6244–6251.
182. Zhu, Y., Wang, Z., Merel, J., Rusu, A., Erez, T., Cabi, S., Tunyasuvunakool, S., Kramár, J., Hadsell, R., de Freitas, N., et al. (2018). Reinforcement and imitation learning for diverse visuomotor skills. *arXiv*, 1802.09564.
183. Yu, T., Abbeel, P., Levine, S., and Finn, C. (2018). One-shot hierarchical imitation learning of compound visuomotor tasks. *arXiv*, 1810.11043.
184. Yu, T., Quillen, D., He, Z., Julian, R., Hausman, K., Finn, C., and Levine, S. (2019). Meta-World: a benchmark and evaluation for multi-task and meta reinforcement learning. *arXiv*, 1910.10897.
185. A. Zeng, S. Song, K.-T. Yu, E. Donlon, F.R. Hogan, M. Bauza, D. Ma, O. Taylor, M. Liu, E. Romo *et al.*, Robotic pick-and-place of novel objects in clutter with multi-affordance grasping and cross-domain image matching. In 2018 IEEE International Conference on Robotics and Automation (ICRA), 2018, pp. 1–8.
186. Tsurumine, Y., Cui, Y., Uchibe, E., and Matsubara, T. (2019). Deep reinforcement learning with smooth policy update: application to robotic cloth manipulation. *Robotics Autonomous Syst.* 112, 72–83.
187. Singh, A., Jang, E., Irpan, A., Kappler, D., Dalal, M., Levine, S., Khansari, M., and Finn, C. (2020). Scalable multi-task imitation learning with autonomous improvement. *arXiv*, 2003.02636.
188. S. Gu, T. Lillicrap, I. Sutskever, and S. Levine, "Continuous deep Q-learning with model-based acceleration. In International Conference on Machine Learning, 2016, pp. 2829–2838.
189. Schulman, J., Wolski, F., Dhariwal, P., Radford, A., and Klimov, O. (2017). Proximal policy optimization algorithms. *arXiv*, 1707.06347.
190. J. Schulman, S. Levine, P. Abbeel, M. Jordan, and P. Moritz, "Trust region policy optimization. In International Conference on Machine Learning, 2015, pp. 1889–1897.
191. Shorten, C., and Khoshgoftaar, T.M. (2019). A survey on image data augmentation for deep learning. *J. Big Data* 6, 60.
192. A.A. Efros and W.T. Freeman, Image quilting for texture synthesis and transfer. In Proceedings of the 28th Annual Conference on Computer Graphics and Interactive Techniques, 2001, pp. 341–346.
193. A. Hertzmann, C.E. Jacobs, N. Oliver, B. Curless, and D.H. Salesin, Image analogies. In Proceedings of the 28th Annual Conference on Computer Graphics and Interactive Techniques, 2001, pp. 327–340.
194. D. Chen, L. Yuan, J. Liao, N. Yu, and G. Hua, Stereoscopic neural style transfer. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 6654–6663.
195. T. Kim, M. Cha, H. Kim, J.K. Lee, and J. Kim, "Learning to discover cross-domain relations with generative adversarial networks. In Proceedings of the 34th International Conference on Machine Learning-Volume 70, 2017, pp. 1857–1865.
196. Z. Yi, H. Zhang, P. Tan, and M. Gong, "Dualgan: Unsupervised dual learning for image-to-image translation. In Proceedings of the IEEE International Conference on Computer Vision, 2017, pp. 2849–2857.
197. Y. Choi, M. Choi, M. Kim, J.-W. Ha, S. Kim, and J. Choo, "StarGAN: unified generative adversarial networks for multi-domain image-to-image translation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 8789–8797.
198. Royer, A., Bousmalis, K., Gouw, S., Bertsch, F., Mosseri, I., Cole, F., and Murphy, K. (2020). XGAN: unsupervised image-to-image translation for many-to-many mappings. In Domain Adaptation for Visual Understanding, R. Singh, M. Vatsa, V.M. Patel, and N. Ratha, eds. (Springer), pp. 33–49.
199. S. Ma, J. Fu, C. Wen Chen, and T. Mei, "DA-GAN: instance-level image translation by deep attention generative adversarial networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 5657–5666.
200. Mo, S., Cho, M., and Shin, J. (2018). InstaGAN: instance-aware image-to-image translation. *arXiv*, 1812.10889.
201. Park, S.C., Park, M.K., and Kang, M.G. (2003). Super-resolution image reconstruction: a technical overview. *IEEE Signal Process. Mag.* 20, 21–36.
202. Hou, H., and Andrews, H. (1978). Cubic splines for image interpolation and digital filtering. *IEEE Trans. Acoust. Speech, Signal Process.* 26, 508–517.
203. K. Zhang, W. Zuo, S. Gu, and L. Zhang, "Learning deep CNN denoiser prior for image restoration. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2017, pp. 3929–3938.
204. X. Ding, Y. Wang, Z. Liang, J. Zhang, and X. Fu, "Towards underwater image enhancement using super-resolution convolutional neural networks. In International Conference on Internet Multimedia Computing and Service, 2017, pp. 479–486.
205. Keys, R. (1981). Cubic convolution interpolation for digital image processing. *IEEE Trans. Acoust. Speech, Signal Process.* 29, 1153–1160.
206. J. Sun, Z. Xu, and H.-Y. Shum, "Image super-resolution using gradient profile prior. In 2008 IEEE Conference on Computer Vision and Pattern Recognition, 2008, pp. 1–8.
207. Farsiu, S., Robinson, D., Elad, M., and Milanfar, P. (2004). Advances and challenges in super-resolution. *Int. J. Imaging Syst. Technology* 14, 47–57.
208. Yang, J., Wright, J., Huang, T.S., and Ma, Y. (2010). Image super-resolution via sparse representation. *IEEE Trans. Image Process.* 19, 2861–2873.
209. R. Zeyde, M. Elad, and M. Protter, On single image scale-up using sparse-representations. In International Conference on Curves and Surfaces, 2010, pp. 711–730.
210. K. Zhang, W. Zuo, and L. Zhang, "Deep plug-and-play super-resolution for arbitrary blur kernels. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 1671–1681.
211. L. Wang, Y. Wang, Z. Liang, Z. Lin, J. Yang, W. An, and Y. Guo, Learning parallax attention for stereo image super-resolution. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 12 250–12 259.
212. Y. Li, V. Tsiminaki, R. Timofte, M. Pollefeys, and L.V. Gool, 3D appearance super-resolution with deep learning. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 9671–9680.
213. Brifman, A., Romano, Y., and Elad, M. (2019). Unified single-image and video super-resolution via denoising algorithms. *IEEE Trans. Image Process.* 28, 6063–6076.
214. Lucas, A., Lopez-Tapia, S., Molina, R., and Katsaggelos, A.K. (2019). Generative adversarial networks and perceptual losses for video super-resolution. *IEEE Trans. Image Process.* 28, 3312–3327.
215. Joshi, N., Matusik, W., Adelson, E.H., and Kriegman, D.J. (2010). Personal photo enhancement using example images. *ACM Trans. Graph.* 29, 12–21.
216. R. Chen, "Image dehazing based on image enhancement algorithm. In 5th International Conference on Information Engineering for Mechanics and Materials, 2015.
217. Fu, X., Huang, J., Ding, X., Liao, Y., and Paisley, J. (2017). Clearing the skies: a deep network architecture for single-image rain removal. *IEEE Trans. Image Process.* 26, 2944–2956.
218. Maini, R., and Aggarwal, H. (2010). A comprehensive review of image enhancement techniques. *arXiv*, 1003.4053.

219. J. Kuruvilla, D. Sukumaran, A. Sankar, and S.P. Joy, A review on image processing and image segmentation. In *International Conference on Data Mining and Advanced Computing (SAPIENCE)*, 2016, pp. 198–203.
220. Schuler, C.J., Hirsch, M., Harmeling, S., and Schölkopf, B. (2015). Learning to deblur. *IEEE Trans. Pattern. Anal. Mach. Intell.* 38, 1439–1451.
221. S. Nah, T. Hyun Kim, and K. Mu Lee, “Deep multi-scale convolutional neural network for dynamic scene deblurring. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 3883–3891.
222. Gulrajani, I., Ahmed, F., Arjovsky, M., Dumoulin, V., and Courville, A.C. (2017). Improved training of Wasserstein GANs. *Advances in Neural Information Processing Systems*, 5767–5777.
223. B. Lu, J.-C. Chen, and R. Chellappa, Unsupervised domain-specific deblurring via disentangled representations. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2019, pp. 10 225–10 234.
224. S. Zhou, J. Zhang, W. Zuo, H. Xie, J. Pan, and J.S. Ren, DAVANet: stereo deblurring with view aggregation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2019, pp. 10996–11005.
225. W. Ren, S. Liu, H. Zhang, J. Pan, X. Cao, and M.-H. Yang, Single image dehazing via multi-scale convolutional neural networks. In *European Conference on Computer Vision*, 2016, pp. 154–169.
226. Zhang, H., Sindagi, V., and Patel, V.M. (2017). Joint transmission map estimation and dehazing using deep networks. *arXiv*, 1708.00581.
227. Ancuti, C.O., Ancuti, C., De Vleeschouwer, C., and Sbrer, M. (2019). Color channel transfer for image dehazing. *IEEE Signal Process. Lett.* 26, 1413–1417.
228. Golts, A., Freedman, D., and Elad, M. (2019). Unsupervised single image dehazing using dark channel prior loss. *IEEE Trans. Image Process.* 29, 2692–2701.
229. X. Liu, Y. Ma, Z. Shi, and J. Chen, “GridDehazeNet: attention-based multi-scale network for image dehazing. In *Proceedings of the IEEE International Conference on Computer Vision*, 2019, pp. 7314–7323.
230. A. Yamashita, Y. Tanaka, and T. Kaneko, “Removal of adherent water-drops from images acquired with stereo camera. In *2005 IEEE/RSJ International Conference on Intelligent Robots and Systems*, 2005, pp. 400–405.
231. A. Yamashita, I. Fukuchi, and T. Kaneko, Noises removal from image sequences acquired with moving camera by estimating camera motion from spatio-temporal information. In *2009 IEEE/RSJ International Conference on Intelligent Robots and Systems*, 2009, pp. 3794–3801.
232. You, S., Tan, R.T., Kawakami, R., Mukaigawa, Y., and Ikeuchi, K. (2015). Adherent raindrop modeling, detection and removal in video. *IEEE Trans. Pattern. Anal. Mach. Intell.* 38, 1721–1733.
233. Garcia-Garcia, A., Orts-Escolano, S., Oprea, S., Villena-Martinez, V., and Garcia-Rodriguez, J. (2017). A review on deep learning techniques applied to semantic segmentation. *arXiv*, 1704.06857.
234. L.-C. Chen, Y. Zhu, G. Papandreou, F. Schroff, and H. Adam, Encoder-decoder with atrous separable convolution for semantic image segmentation. In *Proceedings of the European Conference on Computer Vision (ECCV)*, 2018, pp. 801–818.
235. S.R. Richter, V. Vineet, S. Roth, and V. Koltun, “Playing for data: ground truth from computer games. In *European Conference on Computer Vision*, 2016, pp. 102–118.
236. Q.-H. Pham, B.-S. Hua, T. Nguyen, and S.-K. Yeung, “Real-time progressive 3D semantic segmentation for indoor scenes. In *2019 IEEE Winter Conference on Applications of Computer Vision (WACV)*, 2019, pp. 1089–1098.
237. Z. Liang, M. Yang, L. Deng, C. Wang, and B. Wang, “Hierarchical depth-wise graph convolutional neural network for 3D semantic segmentation of point clouds. In *2019 International Conference on Robotics and Automation (ICRA)*, 2019, pp. 8152–8158.
238. J. Lahoud, B. Ghanem, M. Pollefeys, and M.R. Oswald, 3D instance segmentation via multi-task metric learning. In *Proceedings of the IEEE International Conference on Computer Vision*, 2019, pp. 9256–9266.
239. Y. Li, H. Qi, J. Dai, X. Ji, and Y. Wei, “Fully convolutional instance-aware semantic segmentation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 2359–2367.
240. Luc, P., Couprie, C., Chintala, S., and Verbeek, J. (2016). Semantic segmentation using adversarial networks. *arXiv*, 1611.08408.
241. Liu, M.-Y., Breuel, T., and Kautz, J. (2017). Unsupervised image-to-image translation networks. *Advances in Neural Information Processing Systems*, 700–708.
242. Z. Murez, S. Kolouri, D. Kriegman, R. Ramamoorthi, and K. Kim, Image to image translation for domain adaptation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2018, pp. 4500–4509.
243. Ullman, S. (1979). The interpretation of structure from motion. *Proc. R. Soc. Lond. B Biol. Sci.* 47, 405–426.
244. Scharstein, D., and Szeliski, R. (2002). A taxonomy and evaluation of dense two-frame stereo correspondence algorithms. *Int. J. Comput. Vis.* 47, 7–42.
245. J. Nath Kundu, P. Krishna Uppala, A. Pahuja, and R. Venkatesh Babu, AdaDepth: unsupervised content congruent adaptation for depth estimation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2018, pp. 2656–2665.
246. Y. Chen, C. Schmid, and C. Sminchisescu, Self-supervised learning with geometric constraints in monocular video: connecting flow, depth, and camera. In *Proceedings of the IEEE International Conference on Computer Vision*, 2019, pp. 7063–7072.
247. A. Gordon, H. Li, R. Jonschkowski, and A. Angelova, “Depth from videos in the wild: unsupervised monocular depth learning from unknown cameras. In *Proceedings of the IEEE International Conference on Computer Vision*, 2019, pp. 8977–8986.
248. A. Mousavian, H. Pirsiavash, and J. Košecká, “Joint semantic segmentation and depth estimation with deep convolutional networks. In *2016 Fourth International Conference on 3D Vision (3DV)*, 2016, pp. 611–619.
249. B. Leibe, E. Seemann, and B. Schiele, “Pedestrian detection in crowded scenes. In *IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, vol. 1, 2005, pp. 878–885.
250. Y. Tian, P. Luo, X. Wang, and X. Tang, “Deep learning strong parts for pedestrian detection. In *Proceedings of the IEEE International Conference on Computer Vision*, 2015, pp. 1904–1912.
251. N. Dalal and B. Triggs, Histograms of oriented gradients for human detection. In *IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR’05)*, vol. 1, 2005, pp. 886–893.
252. Nam, W., Dollár, P., and Han, J.H. (2014). Local decorrelation for improved pedestrian detection. *Advances in Neural Information Processing Systems*, 424–432.
253. J. Cao, Y. Pang, and X. Li, Pedestrian detection inspired by appearance constancy and shape symmetry. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2016, pp. 1316–1324.
254. R. Yin, Multi-resolution generative adversarial networks for tiny-scale pedestrian detection. In *2019 IEEE International Conference on Image Processing (ICIP)*, 2019, pp. 1665–1669.
255. Xie, J., Pang, Y., Cholakkal, H., Anwer, R.M., Khan, F.S., and Shao, L. (2020). PSC-Net: learning part spatial co-occurrence for occluded pedestrian detection. *arXiv*, 2001.09252.
256. J. Wang, X. Zhu, S. Gong, and W. Li, Transferable joint attribute-identity deep learning for unsupervised person re-identification. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2018, pp. 2275–2284.
257. Fan, H., Zheng, L., Yan, C., and Yang, Y. (2018). Unsupervised person re-identification: clustering and fine-tuning. *ACM Trans. Multimedia Comput. Commun. Appl. (TOMM)* 14, 1–18.

258. Song, L., Wang, C., Zhang, L., Du, B., Zhang, Q., Huang, C., and Wang, X. (2020). Unsupervised domain adaptive re-identification: theory and practice. *Pattern Recogn.* 102, 107173.
259. R. Hou, B. Ma, H. Chang, X. Gu, S. Shan, and X. Chen, "VRSTC: occlusion-free video person re-identification. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 7183–7192.
260. X. Sun and L. Zheng, "Dissecting person re-identification from the viewpoint of viewpoint. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 608–617.
261. H. Nam and B. Han, Learning multi-domain convolutional neural networks for visual tracking. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2016, pp. 4293–4302.
262. M. Kempka, M. Wydmuch, G. Runc, J. Toczek, and W. Jaśkowski, ViZ-Doom: A Doom-based AI research platform for visual reinforcement learning. In 2016 IEEE Conference on Computational Intelligence and Games (CIG). IEEE, 2016, pp. 1–8.
263. A. Alahi, K. Goel, V. Ramanathan, A. Robicquet, L. Fei-Fei, and S. Savarese, Social LSTM: human trajectory prediction in crowded spaces. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2016, pp. 961–971.
264. Bellemare, M., Dabney, W., Dadashi, R., Taiga, A.A., Castro, P.S., Le Roux, N., Schuurmans, D., Lattimore, T., and Lyle, C. (2019). A geometric perspective on optimal representations for reinforcement learning. *Advances in Neural Information Processing Systems*, 4360–4371.
265. Jaderberg, M., Mnih, V., Czarnecki, W.M., Schaul, T., Leibo, J.Z., Silver, D., and Kavukcuoglu, K. (2016). Reinforcement learning with unsupervised auxiliary tasks. *arXiv*, 1611.05397.
266. Mirowski, P., Grimes, M., Malinowski, M., Hermann, K.M., Anderson, K., Teplyashin, D., Simonyan, K., Kavukcuoglu, K., Zisserman, A., and Hadsell, R. (2018). Learning to navigate in cities without a map. *Advances in Neural Information Processing Systems*, 2419–2430.
267. Hsu, K., Levine, S., and Finn, C. (2018). Unsupervised learning via meta-learning. *arXiv*, 1810.02334.
268. Kompella, V.R., Stollenga, M., Luciw, M., and Schmidhuber, J. (2017). Continual curiosity-driven skill acquisition from high-dimensional video inputs for humanoid robots. *Artif. Intelligence* 247, 313–335.
269. Andrychowicz, M., Wolski, F., Ray, A., Schneider, J., Fong, R., Welinder, P., McGrew, B., Tobin, J., Abbeel, O.P., and Zaremba, W. (2017). Hind-sight experience replay. *Advances in Neural Information Processing Systems*, 5048–5058.
270. E. Todorov T. Erez Y. Tassa MuJoCo: a physics engine for model-based control. In 2012 IEEE/RSJ International Conference on Intelligent Robots and Systems, 2012, pp. 5026–5033.
271. D. Quillen E. Jang O. Nachum C. Finn J. Ibarz S. Levine Deep reinforcement learning for vision-based robotic grasping: a simulated comparative evaluation of off-policy methods. In 2018 IEEE International Conference on Robotics and Automation (ICRA), 2018, pp. 6284–6291.
272. Kalashnikov, D., Irpan, A., Pastor, P., Ibarz, J., Herzog, A., Jang, E., Quillen, D., Holly, E., Kalakrishnan, M., Vanhoucke, V., et al. (2018). QT-Opt: scalable deep reinforcement learning for vision-based robotic manipulation. *arXiv*, 1806.10293.
273. Che, T., Li, Y., Jacob, A.P., Bengio, Y., and Li, W. (2016). Mode regularized generative adversarial networks. *arXiv*, 1612.02136.
274. A. Ghosh, V. Kulharia, V.P. Namboodiri, P.H. Torr, and P.K. Dokania, "Multi-agent diverse generative adversarial networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 8513–8521.
275. A.R. Zamir, A. Sax, W. Shen, L.J. Guibas, J. Malik, and S. Savarese, "Taskonomy: Disentangling task transfer learning. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2018, pp. 3712–3722.
276. K. Bousmalis, N. Silberman, D. Dohan, D. Erhan, and D. Krishnan, "Unsupervised pixel-level domain adaptation with generative adversarial networks. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2017, pp. 3722–3731.
277. X. Wang, Q. Huang, A. Celikyilmaz, J. Gao, D. Shen, Y.-F. Wang, W.Y. Wang, and L. Zhang, "Reinforced cross-modal matching and self-supervised imitation learning for vision-language navigation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2019, pp. 6629–6638.
278. Buşoniu, L., Babuška, R., and De Schutter, B. (2010). Multi-agent reinforcement learning: an overview. In *Innovations in Multi-Agent Systems and Applications – 1*, D. Srinivasan and L.C. Jain, eds. (Springer), pp. 183–221.
279. Li, Y. (2017). Deep reinforcement learning: an overview. *arXiv*, 1701.07274.
280. E. Rohmer, S.P. Singh, and M. Freese, "V-REP: a versatile and scalable robot simulation framework. In 2013 IEEE/RSJ International Conference on Intelligent Robots and Systems, 2013, pp. 1321–1326.
281. Andrychowicz, O.M., Baker, B., Chociej, M., Jozefowicz, R., McGrew, B., Pachocki, J., Petron, A., Plappert, M., Powell, G., Ray, A., et al. (2020). Learning dexterous in-hand manipulation. *Int. J. Robotics Res.* 39, 3–20.
282. M. Savva, A. Kadian, O. Maksymets, Y. Zhao, E. Wijmans, B. Jain, J. Straub, J. Liu, V. Koltun, J. Malik et al., Habitat: a platform for embodied AI research. In Proceedings of the IEEE International Conference on Computer Vision, 2019, pp. 9339–9347.