



人工智能原理与技术



生成式对抗网络



【GAN这4.5年】Ian Goodfellow 谈论生成对抗网络在人脸生成进展，5篇经典论文，从黑白模糊到惟妙惟肖

原创：专知 专知 今天

【导读】生成式对抗网络(Generative Adversarial Networks, GAN)自在2014年被Ian Goodfellow提出后，取得巨大的进展，在理论算法和应用方面有着丰富成果。近日，GAN之父Ian Goodfellow在twitter谈论了关于GAN在人脸生成的4年半进展，包含5篇代表性文章，值得关注，人脸生成从早起的模糊阶段进化到现在逼真的程度。

GAN在人脸生成这4.5年进展



目录

CONTENTS

01

背景

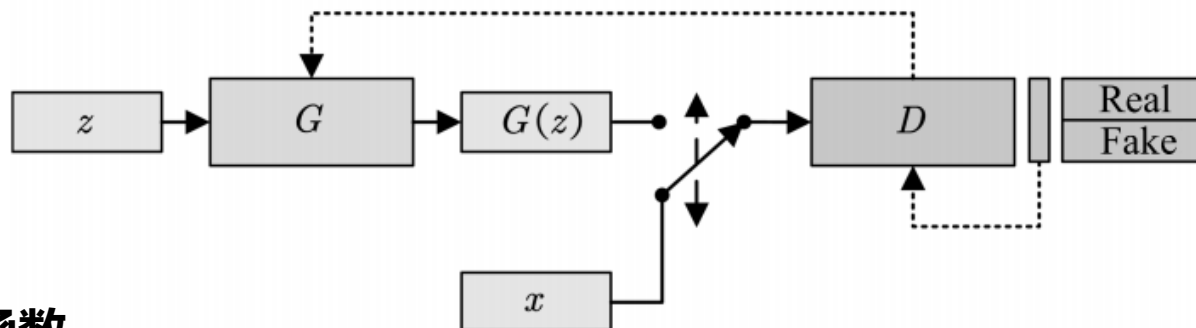
02

示例

背景

Ian Goodfellow

模型



目标函数

$$\min_G \max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

作者 Ian Goodfellow 于2014年NIPS首次提出了GAN的概念



背景

Yann LeCun

What are some recent and potentially upcoming breakthroughs in deep learning?



Yann LeCun, Director of AI Research at Facebook and Professor at NYU

Answered Jul 29, 2016 · Upvoted by Joaquin Quiñero Candela, studied Machine Learning and Gokul Krishnan, M.Sc Computer Science & Machine Learning, ETH Zurich (2018)

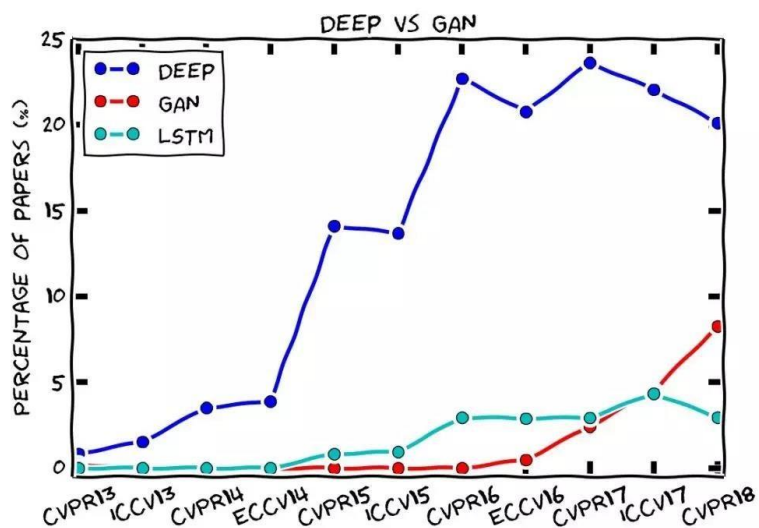
The most important one, in my opinion, is adversarial training (also called GAN for Generative Adversarial Networks). This is an idea that was originally proposed by Ian Goodfellow when he was a student with Yoshua Bengio at the University of Montreal (he since moved to Google Brain and recently to OpenAI). This, and the variations that are now being proposed is the most interesting idea in the last 10 years in ML, in my opinion.



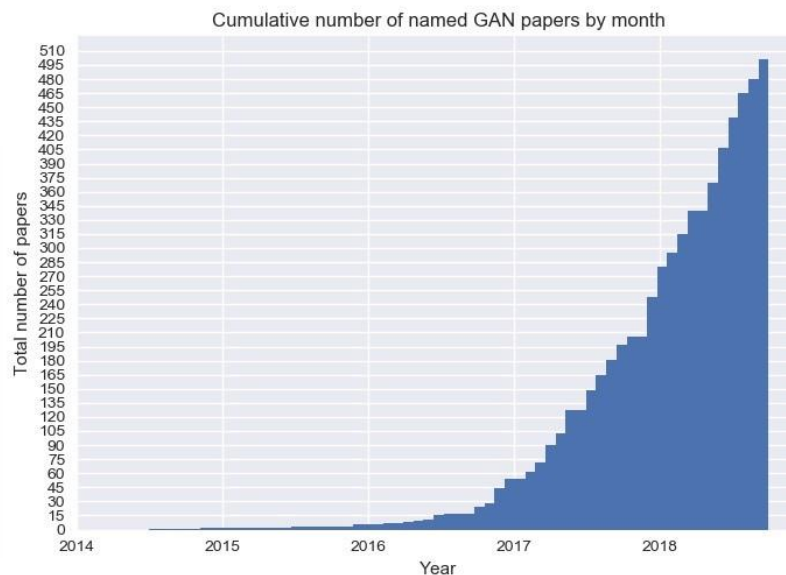
Yann LeCun 是美国工程院院士
Facebook人工智能研究院院长
纽约大学Sliver教授
创立卷积神经网络
人工智能三巨头之一

背景

论文数量



计算机视觉顶级会议 (CVPR, ICCV, ECCV) 上, GAN已经超过LSTM, 成为重要关键词。



ArXiv上, 每个月和GAN相关的论文数量

<https://github.com/hindupuravinash/the-gan-zoo>

目录

CONTENTS

01

背景

02

示例

GAN案例

图像着色



用途：旧图像修复

GAN案例

图像超像素



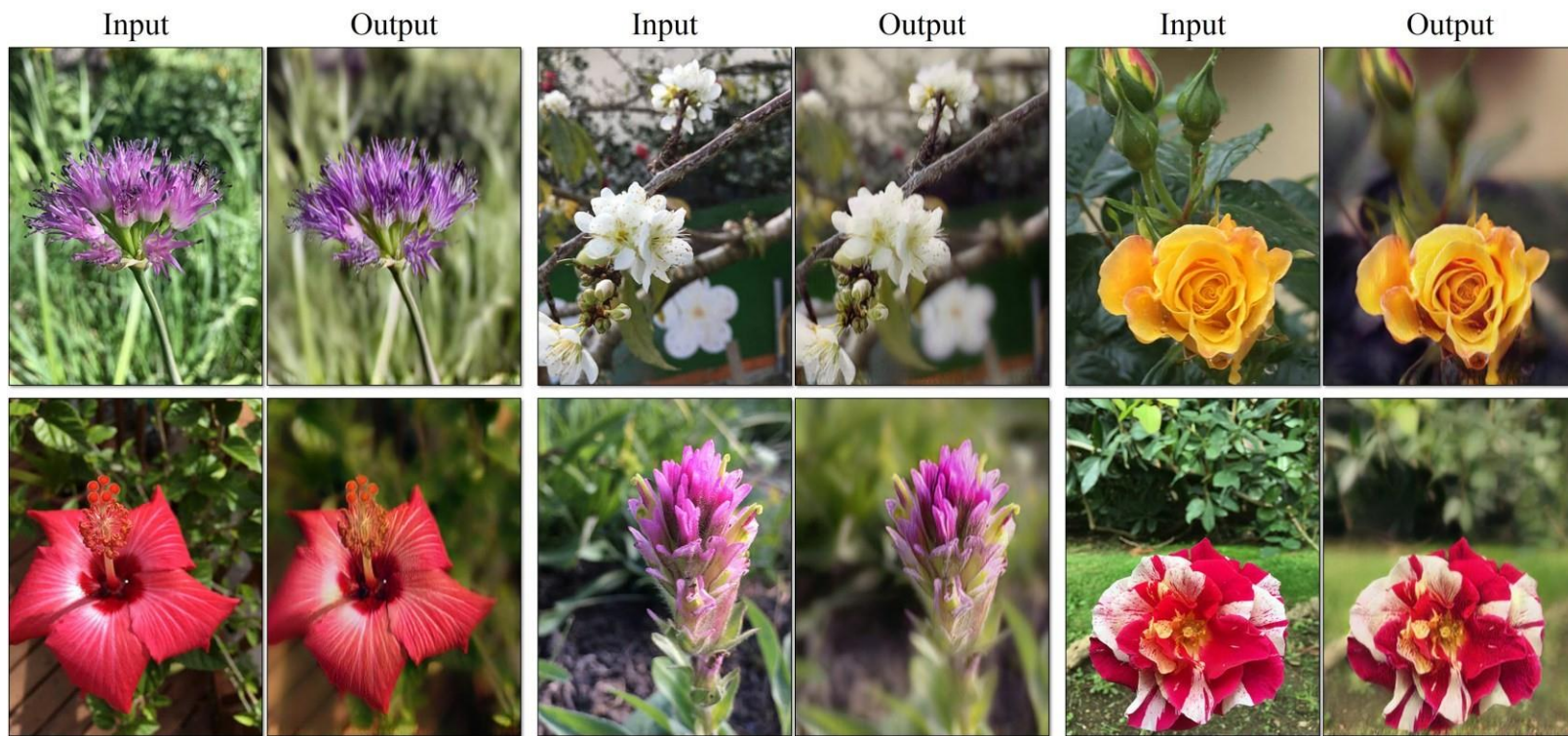
用途：手机相机

https://github.com/JustinhoCHN/SRGAN_Wasserstein

超分辨率算法已经广泛的应用于我们的智能手机中

GAN案例

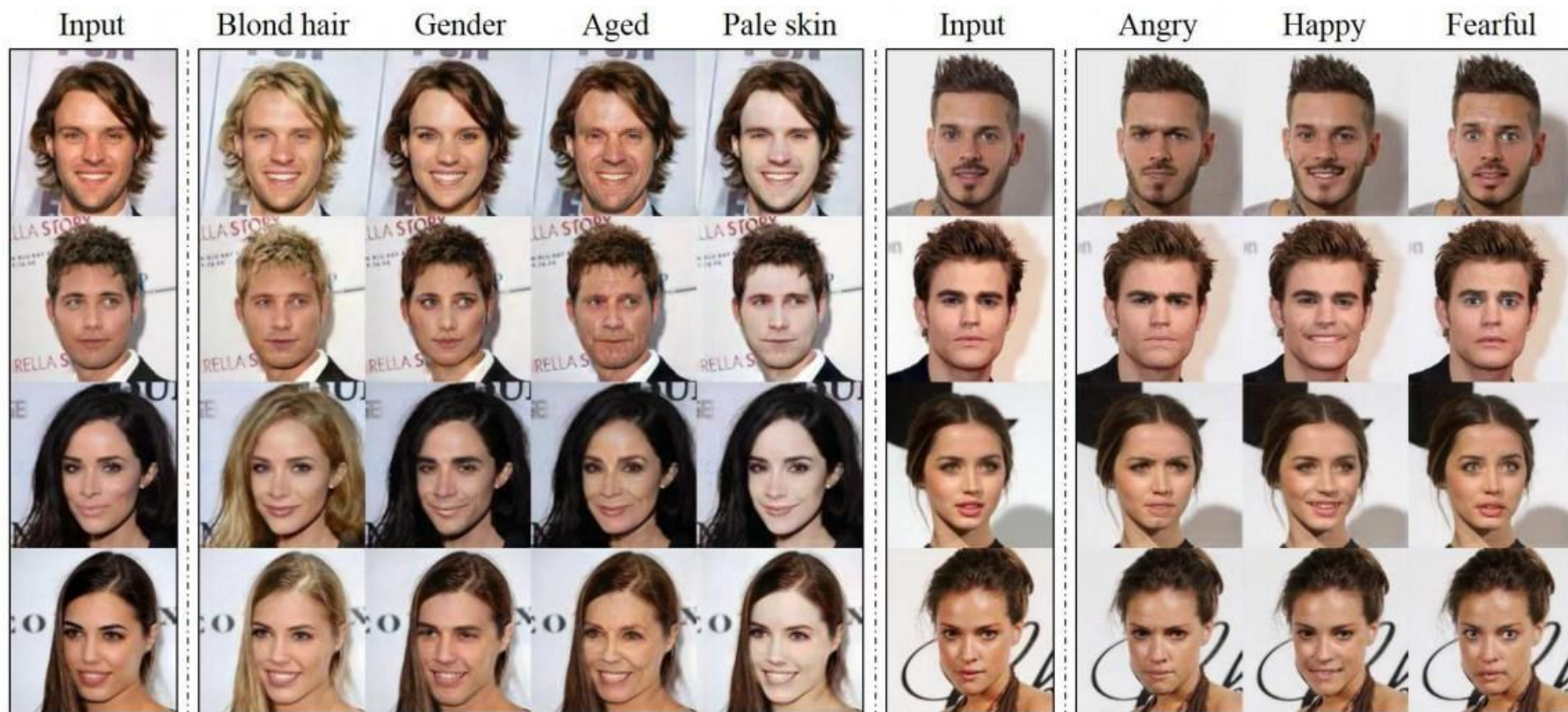
背景模糊



<https://github.com/junyanz/CycleGAN>

GAN案例

人脸生成



<https://github.com/yunjey/stargan>

用途：娱乐软件，抖音，美图等

GAN案例

人脸定制

INSTRUCTION: press +/- to adjust feature, toggle feature name to lock the feature



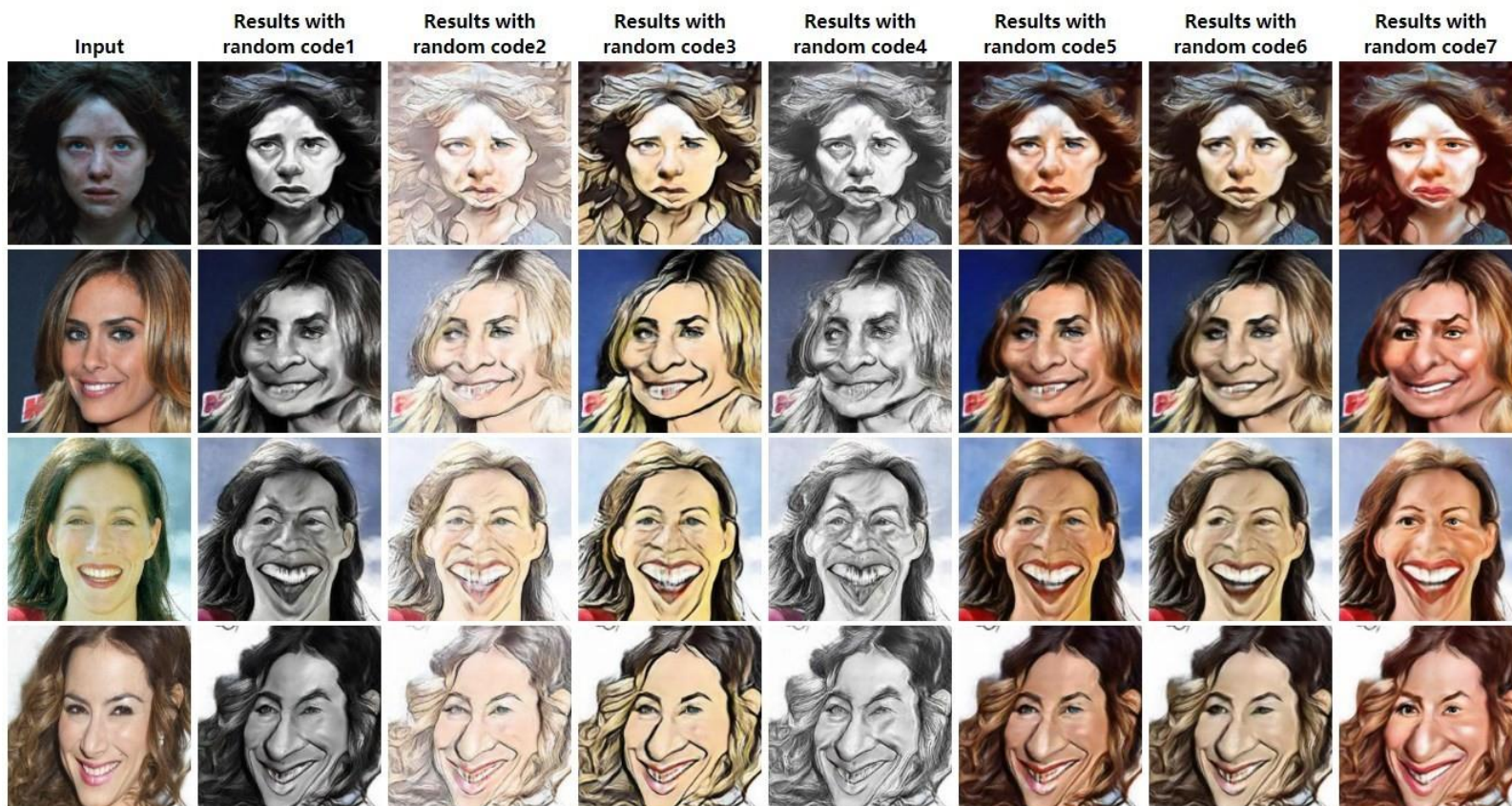
random face

Male	Age	Skin_Tone
- +	- +	- +
Bangs	Hairline	Bald
- +	- +	- +
Big_Nose	Pointy_Nose	Makeup
- +	- +	- +
Smiling	Mouth_Open	Wavy_Hair
- +	- +	- +
Beard	Goatee	Sideburns
- +	- +	- +
Blond_Hair	Black_Hair	Gray_Hair
- +	- +	- +
Eyeglasses	Earrings	Necktie
- +	- +	- +

https://github.com/SummitKwan/transparent_latent_gan

GAN案例

卡通头像生成



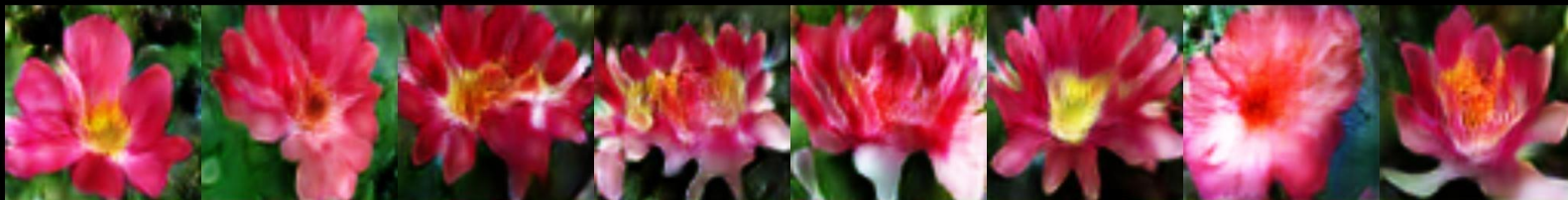
CariGANs: Unpaired Photo-to-Caricature Translation

GAN案例

文本生成图片

This flower has overlapping pink pointed petals surrounding a ring of short yellow filaments

Stage-I



Stage-II



<https://github.com/hanzhanggit/StackGAN>

GAN案例

字体变换

千 萬 孤 獨
山 徑 舟 釣
鳥 人 蓑 寒
飛 踪 笠 江
絕 滅 翁 雪

<https://github.com/kaonashi-tyc/zi2zi>

GAN案例

风格变换

<https://github.com/junyanz/CycleGAN>



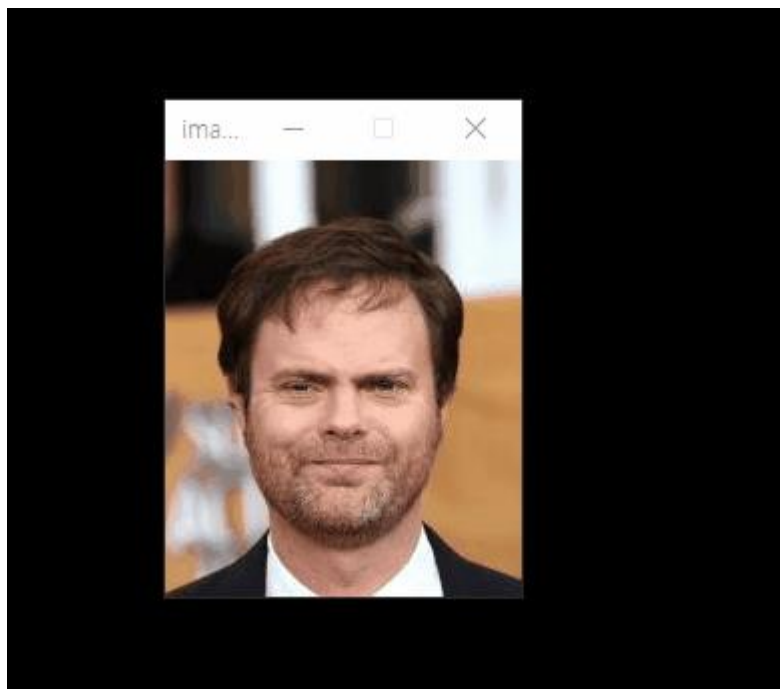
winter Yosemite → summer Yosemite



summer Yosemite → winter Yosemite

GAN案例

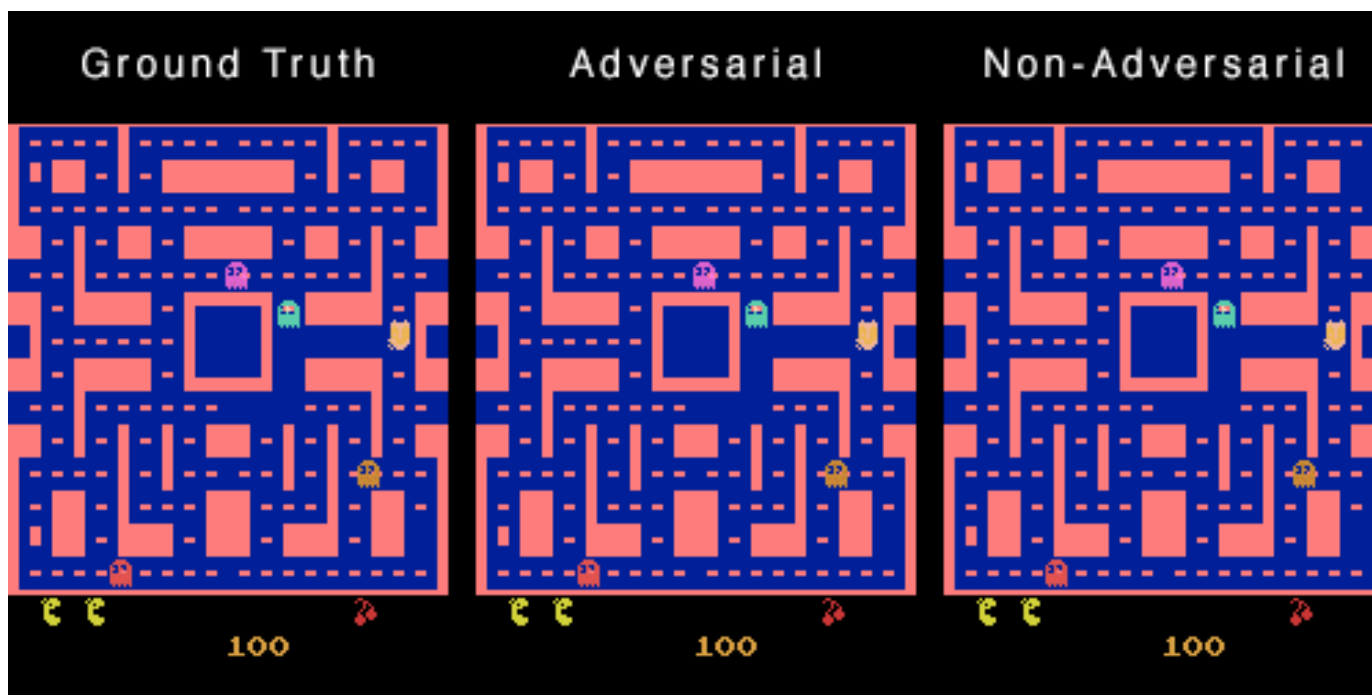
图像修复



https://github.com/shinseung428/GlobalLocalImageCompletion_TF

GAN案例

帧预测



https://github.com/dyelax/Adversarial_Video_Generation

目录

CONTENTS

01

GAN

Generative Adversarial Nets, 生成式对抗网络

02

cGAN

Conditional GAN, 条件生成式对抗网络

03

DCGAN

Deep Convolutional GAN, 深度卷积生成式对抗网络

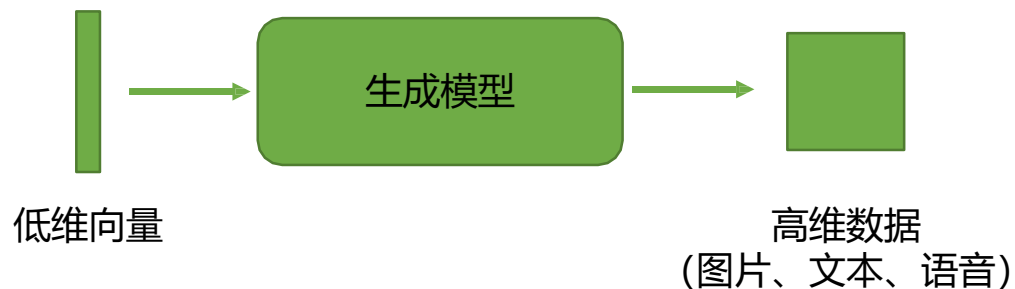
04

WGAN/WGAN-GP

Wasserstein GAN with Weight Clipping/ Gradient Penalty

GAN

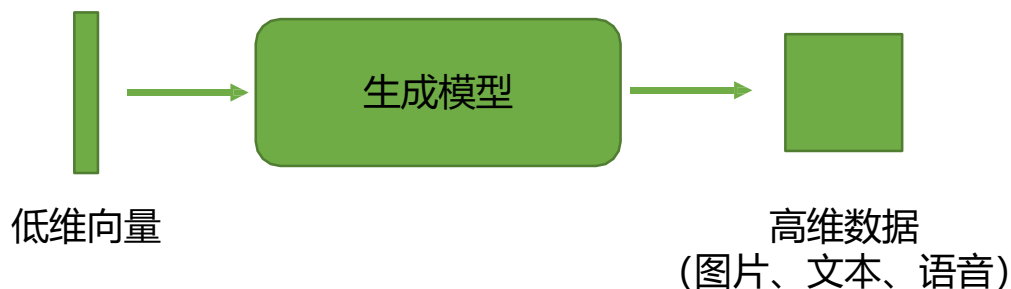
生成模型和GAN



不严谨的定义：一个能够生成我们想要的数据的模型(图模型、函数、神经网络)。

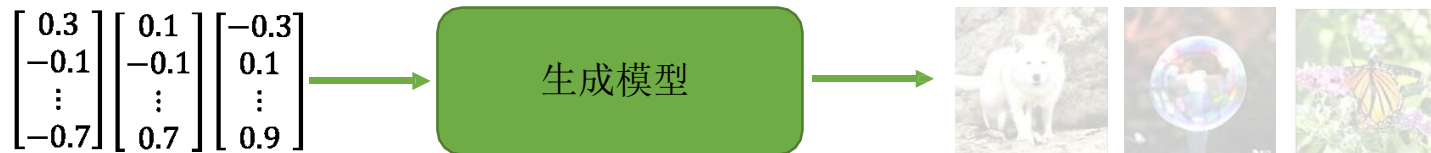
GAN

生成模型和GAN



不严谨的定义：一个能够生成我们想要的数据的模型(图模型、函数、神经网络)。

图像生成



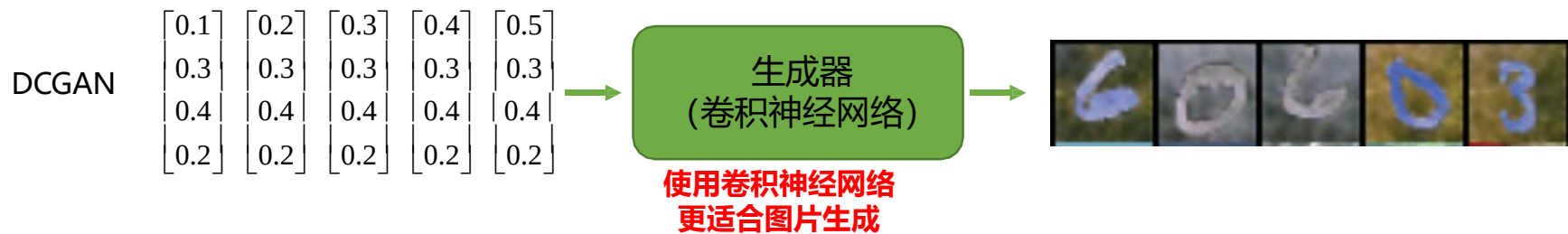
文本生成



生成式对抗网络 (GAN) 的目的是训练这样一个生成模型，生成我们想要的数据。
生成式对抗网络中**生成**一词的由来。

目录

简介



目录

CONTENTS

01

GAN

Generative Adversarial Nets, 生成式对抗网络

02

cGAN

Conditional GAN, 条件生成式对抗网络

03

DCGAN

Deep Convolutional GAN, 深度卷积生成式对抗网络

04

WGAN/WGAN-GP

Wasserstein GAN with Weight Clipping/ Gradient Penalty

GAN

零和博弈

为什么罪犯制造的假币越来越逼真？
为什么GAN可以生成数据？

罪犯



使用假币，购买商品
目标：制造逼真假币



商贩



收到真币，出售商品
收到假币，不出售商品
目标：准确区分真假币



正版人民币
是红色的！

正版人民币有
精美的花纹！

正版人民币有
毛爷爷头像！



正版人民币

GAN

零和博弈

罪犯



使用假币，购买商品
目标：制造逼真假币



大家后面学习时请思考两个问题：
为什么罪犯制造的假币越来越真？
为什么GAN可以生成数据？

收到真币，出售商品
收到假币，不出售商品
目标：准确区分真假币

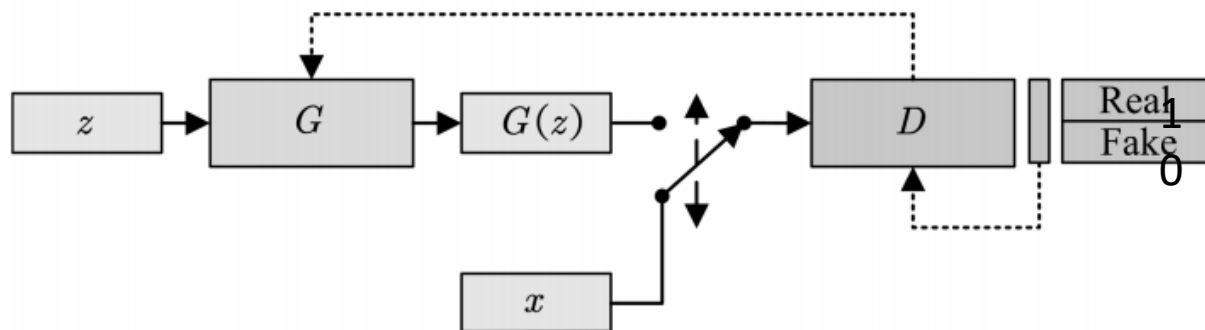


正版人民币

GAN

生成式对抗网络

GAN框架



GAN工作原理

随机噪声 z : 从一个先验分布（人为定义，一般是**均匀分布**或者**正态分布**）中随机采样的向量

真实样本 x : 从数据库中采样的样本；合成样本 $G(z)$: 生成模型 G 输出的样本。

判别器(Discriminator): **区分真实(real)样本和虚假(fake)样本**。对于真实样本，尽可能给出高的评分1；对于虚假数据，尽可能给出低个评分0。

生成器(Generator): **欺骗判别器**。生成虚假数据，使得判别器 D 能够尽可能给出高的评分1。

$\begin{bmatrix} 0.1 \\ 0.3 \\ 0.4 \\ 1 \end{bmatrix}$



GAN目标函数

$$\max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

$$\min_G V(G, D) = E_{z \sim p_z} [\log(1 - D(G(z)))]$$

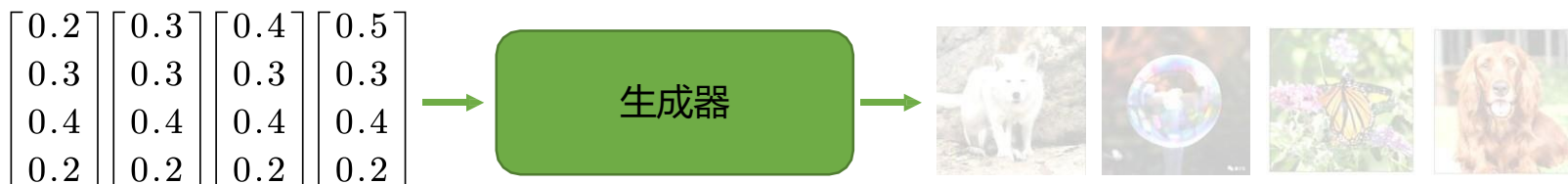
↑ 等价于

$$\min_G \max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

GAN

生成器

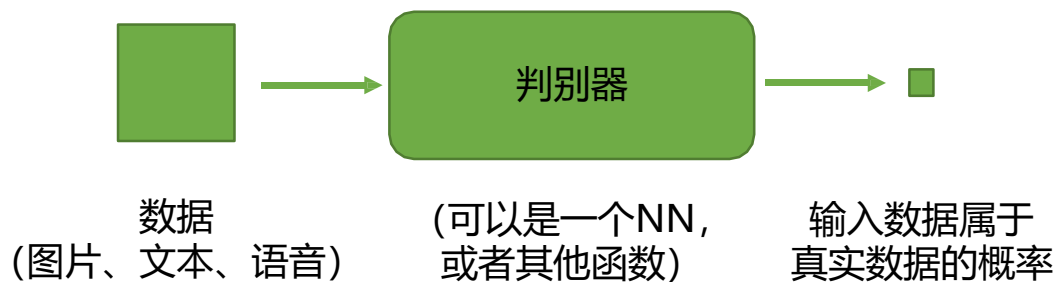
生成器：生成数据



GAN

判别器

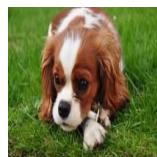
判别器：区分真样本和假样本



以图像为例



真实图片



真实图片



生成图片

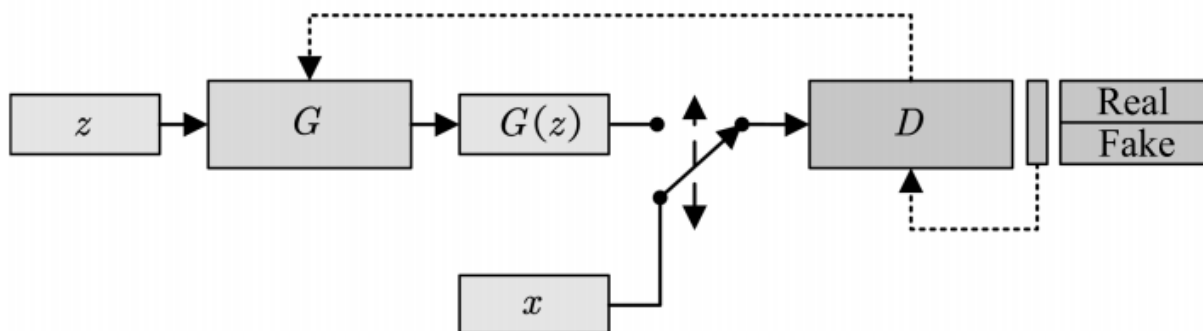


生成图片



GAN

GAN



【核心思想】 判别器：区分真假样本；生成器：欺骗判别器

目标函数

$$\max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

$$\min_G V(G, D) = E_{z \sim p_z} [\log(1 - D(G(z)))]$$

↕ 等价于

$$\min_G \max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

GAN

博弈

生成器
罪犯



使用假币，购买商品
目标：制造逼真假币

欺骗判别器

判别器
商贩



收到真币，出售商品
收到假币，不出售商品
目标：准确区分真假币

区分真假数据

对抗训练

更新生成器



合成数据



正版人民币
是红色的！



更新判别器



正版人民币有
精美的花纹！

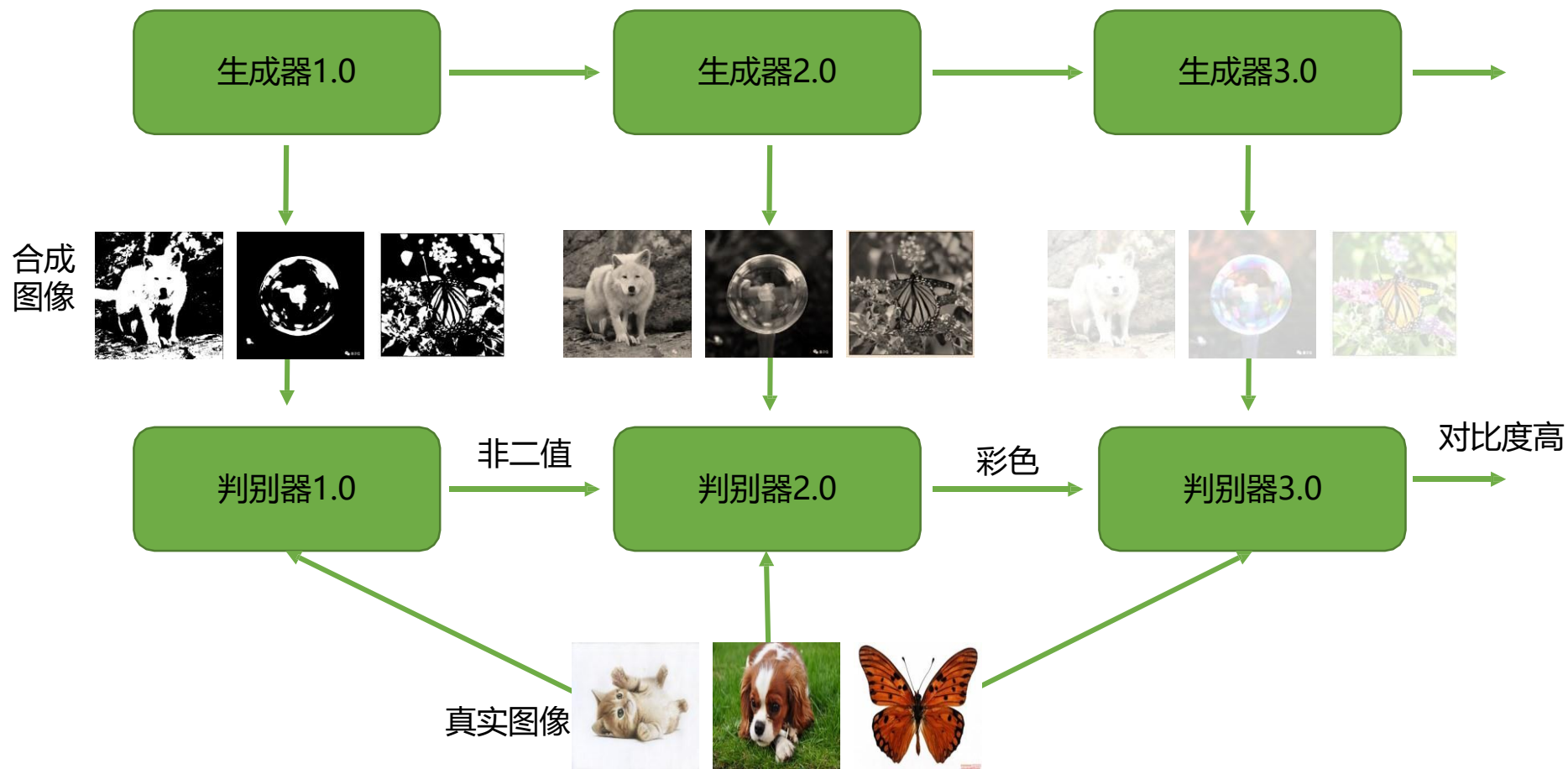
正版人民币有
毛爷爷头像！



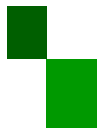
正版人民币 真实数据

GAN

对抗训练



生成器和判别器之间存在对抗的关系，这也是生成式对抗网络中**对抗**一词的由来



GAN

训练算法

【核心思想】 判别器：区分真假样本；生成器：欺骗判别器

$$\max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

$$\min_G V(G, D) = E_{z \sim p_z} [\log(1 - D(G(z)))]$$

GAN

【核心思想】 判别器：区分真假样本；生成器：欺骗判别器

$$\max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

$$\min_G V(G, D) = E_{z \sim p_z} [\log(1 - D(G(z)))]$$

训练算法

1. 随机初始化生成器和判别器

生成器

判别器

2. 交替训练判别器D 和 生成器 G，直到收敛

步骤1： **固定生成器G，训练判别器D区分真实图像与合成图像。**

赋予真实图像高分，赋予合成图像低分

GAN

【核心思想】判别器：区分真假样本；生成器：欺骗判别器

$$\max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

$$\min_G V(G, D) = E_{z \sim p_z} [\log(1 - D(G(z)))]$$

训练算法

1. 随机初始化生成器和判别器

生成器

判别器

2. 交替训练判别器D和生成器G，直到收敛

步骤1：固定生成器G，训练判别器D区分真实图像与合成图像。

赋予真实图像高分，赋予合成图像低分

从数据库中
随机采样的
真实图像



合成图像



判别器
(不固定，需要更新)

0.9

0.1

$$\begin{bmatrix} 0.2 \\ 0.3 \\ 0.4 \\ 0.2 \end{bmatrix} \begin{bmatrix} 0.3 \\ 0.3 \\ 0.4 \\ 0.2 \end{bmatrix} \begin{bmatrix} 0.4 \\ 0.3 \\ 0.4 \\ 0.2 \end{bmatrix} \begin{bmatrix} 0.5 \\ 0.3 \\ 0.4 \\ 0.2 \end{bmatrix}$$

从先验分布中采样随机噪声

生成器
(固定的、不需要更新)

请思考：如何训练呢？

GAN

【核心思想】判别器：区分真假样本；生成器：欺骗判别器

$$\max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

$$\min_G V(G, D) = E_{z \sim p_z} [\log(1 - D(G(z)))]$$

训练算法

1. 随机初始化生成器和判别器

生成器

判别器

2. 交替训练判别器D和生成器G，直到收敛

步骤1：固定生成器G，训练判别器D区分真实图像与合成图像。

赋予真实图像高分，赋予合成图像低分

从数据库中
随机采样的
真实图像



合成图像



判别器
(不固定，需要更新)

0.9

0.1

生成器不需要更新，只要不断生成数据就行

只需要用正常的方式训练一个两类别的分类器！

GAN

【核心思想】判别器：区分真假样本；生成器：欺骗判别器

$$\max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

$$\min_G V(G, D) = E_{z \sim p_z} [\log(1 - D(G(z)))]$$

训练算法

1. 随机初始化生成器和判别器

生成器

判别器

2. 交替训练判别器D 和 生成器 G，直到收敛

步骤2：固定判别器D，训练生成器G欺骗判别器D。
更新生成器的参数，使其合成的图片被生成器D赋予高分。

GAN

【核心思想】判别器：区分真假样本；生成器：欺骗判别器

$$\max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

$$\min_G V(G, D) = E_{z \sim p_z} [\log(1 - D(G(z)))]$$

训练算法

1. 随机初始化生成器和判别器

生成器

判别器

2. 交替训练判别器D 和 生成器 G，直到收敛

步骤2：固定判别器D，训练生成器G欺骗判别器D。
更新生成器的参数，使其合成的图片被生成器D赋予高分。



0.2	0.3	0.4	0.5
0.3	0.3	0.3	0.3
0.4	0.4	0.4	0.4
0.2	0.2	0.2	0.2

从先验分布中随机采样噪声

请思考：如何优化生成器呢？

GAN

【核心思想】判别器：区分真假样本；生成器：欺骗判别器

$$\max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

$$\min_G V(G, D) = E_{z \sim p_z} [\log(1 - D(G(z)))]$$

训练算法

1. 随机初始化生成器和判别器

生成器

判别器

2. 交替训练判别器D 和 生成器 G，直到收敛

步骤2：固定判别器D，训练生成器G欺骗判别器D。
更新生成器的参数，使其合成的图片被生成器D赋予高分。



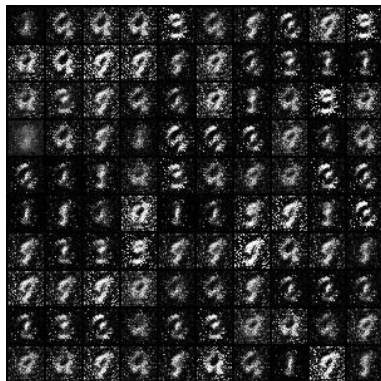
0.2	0.3	0.4	0.5
0.3	0.3	0.3	0.3
0.4	0.4	0.4	0.4
0.2	0.2	0.2	0.2

从先验分布中随机采样噪声

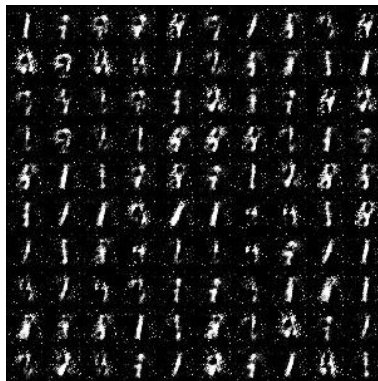
看成输入数据

GAN

实验结果



1 EPOCH



5 EPOCH



10 EPOCH



20 EPOCH



50 EPOCH



200 EPOCH

代码演示

GAN

Pytorch实现GAN

准备工作

定义网络结构：判别器和生成器的网络结构

定义目标函数：根据预测标签和真实标签，计算损失

定义优化器：优化器，自动计算梯度并优化网络权重

数据加载：加载数据集

```
class Discriminator(nn.Module):
    """判别器"""
    def __init__(self):
        super(Discriminator, self).__init__()

        layers = []
        # 32 * 32
        layers.append(nn.Linear(in_features=28*28, out_features=128))
        layers.append(nn.ReLU(True))
        # 16 * 16
        layers.append(nn.Linear(in_features=128, out_features=64))
        layers.append(nn.ReLU(True))
        # 8 * 8
        layers.append(nn.Linear(in_features=64, out_features=32))
        layers.append(nn.ReLU(True))
        # 4 * 4
        layers.append(nn.Linear(in_features=32, out_features=16))
        layers.append(nn.ReLU(True))
        # 1 * 1
        layers.append(nn.Linear(in_features=16, out_features=1))
        layers.append(nn.Sigmoid())

        self.model = nn.Sequential(layers)

    def forward(self, x):
        x = x.view(-1, 28*28)
        validity = self.model(x)
        return validity
```

判别器

```
class Generator(nn.Module):
    """生成器"""
    def __init__(self):
        super(Generator, self).__init__()

        layers = []
        # 32 * 32
        layers.append(nn.Linear(in_features=1, out_features=256))
        layers.append(nn.ReLU(True))
        # 16 * 16
        layers.append(nn.Linear(in_features=256, out_features=128))
        layers.append(nn.ReLU(True))
        # 8 * 8
        layers.append(nn.Linear(in_features=128, out_features=64))
        layers.append(nn.ReLU(True))
        # 4 * 4
        layers.append(nn.Linear(in_features=64, out_features=32))
        layers.append(nn.ReLU(True))
        # 1 * 1
        layers.append(nn.Linear(in_features=32, out_features=28*28))
        layers.append(nn.Tanh())

        self.model = nn.Sequential(layers)

    def forward(self, x):
        x = self.model(x)
        x = x.view(-1, 1, 28, 28)
        return x
```

生成器

```
# 加载数据集
data_loader = DataLoader(
    dataset=dataset,
    batch_size=batch_size,
    shuffle=True,
    num_workers=num_workers,
)

# 训练过程
for epoch in range(1, num_epochs + 1):
    for i, (real_images, real_labels) in enumerate(data_loader):
        # 生成假数据
        fake_images = generator(real_images)

        # 计算损失
        loss = discriminator_loss(real_images, real_labels, fake_images)

        # 优化器
        optimizer.zero_grad()
        loss.backward()
        optimizer.step()
```

其他准备工作

训练直到收敛

加载数据

提取一个batch的真实数据

生成一个batch的虚假数据

计算损失并优化模型

计算判别器损失，并优化判别器

计算生成器损失，并优化生成器

```
# 训练过程
for epoch in range(1, num_epochs + 1):
    for i, (real_images, real_labels) in enumerate(data_loader):
        # 生成假数据
        fake_images = generator(real_images)

        # 计算损失
        loss = discriminator_loss(real_images, real_labels, fake_images)

        # 优化器
        optimizer.zero_grad()
        loss.backward()
        optimizer.step()

    # 生成假数据
    fake_images = generator(real_images)

    # 计算损失
    loss = generator_loss(real_images, fake_images)

    # 优化器
    optimizer.zero_grad()
    loss.backward()
    optimizer.step()

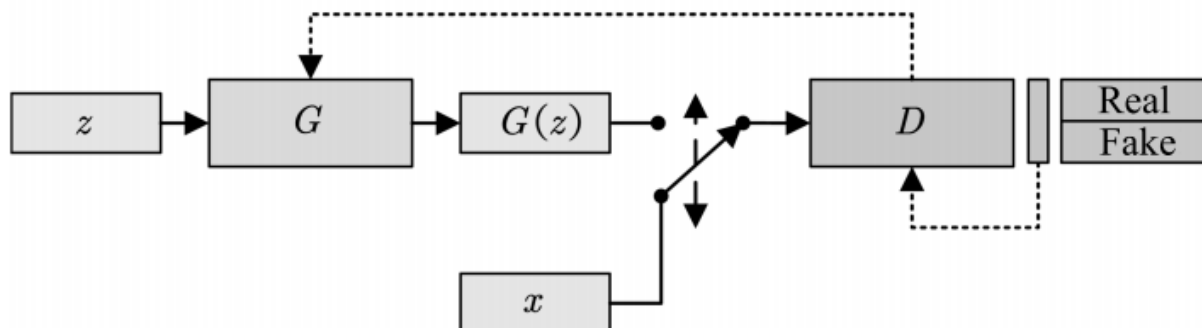
    # 打印损失
    print(f"Epoch {epoch} | Discriminator Loss: {discriminator_loss} | Generator Loss: {generator_loss} | Total Loss: {total_loss}")
```

训练过程

GAN

生成式对抗网络

GAN框架



GAN工作原理

判别器(Discriminator): 区分真实(real)数据和虚假(fake)数据。对于真实数据, 尽可能给出高的评分1; 对于虚假数据, 尽可能给出低个评分0.

生成器(Generator): 生成虚假数据, 使得判别器D能够尽可能给出高的评分1。

GAN目标函数

$$\min_G \max_D = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

GAN

KL散度和JS散度

KL散度：一种衡量两个概率分布的匹配程度的指标

$$KL(P_1||P_2) = \mathbb{E}_{x \sim P_1} [\log \frac{P_1}{P_2}]$$

对于**连续**样本 $KL(P_1||P_2) = \int_{\mathcal{X}} P_1(x) \log \frac{P_1(x)}{P_2(x)} dx$

对于**离散**样本 $KL(P_1||P_2) = \sum_x P_1(x) \log \frac{P_1(x)}{P_2(x)}$

当且 $P_1=P_2$ 的时候，KL散度为零

KL散度具有**非负性**（吉布斯不等式）：

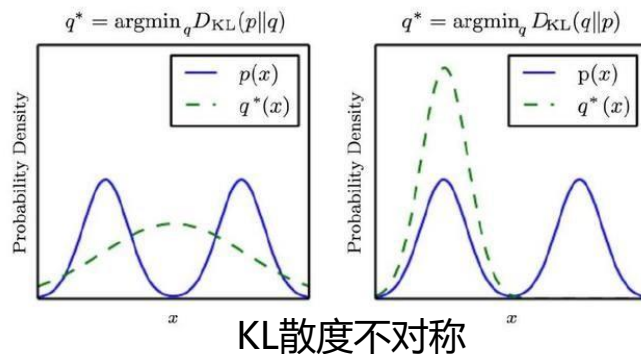
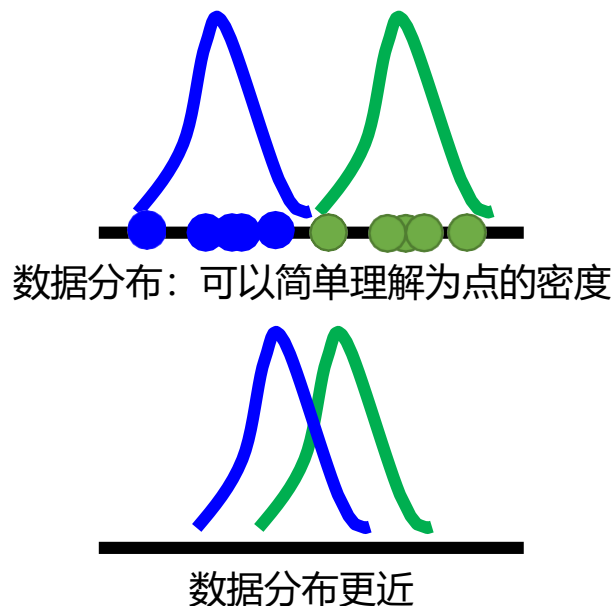
已知 $\log(x) \leq x - 1$

$$-KL(P_1||P_2) = \sum_x P_1(x) \log \frac{P_2(x)}{P_1(x)} \leq \sum_x P_1(x) \left(\frac{P_2(x)}{P_1(x)} - 1 \right) = \sum_x P_2(x) - \sum_x P_1(x) = 0$$

JS散度：

$$JS(P_1||P_2) = \frac{1}{2} KL(P_1||\frac{P_1+P_2}{2}) + \frac{1}{2} KL(P_2||\frac{P_1+P_2}{2})$$

JS散度具有 非负性、以及对称性



GAN

数据分布

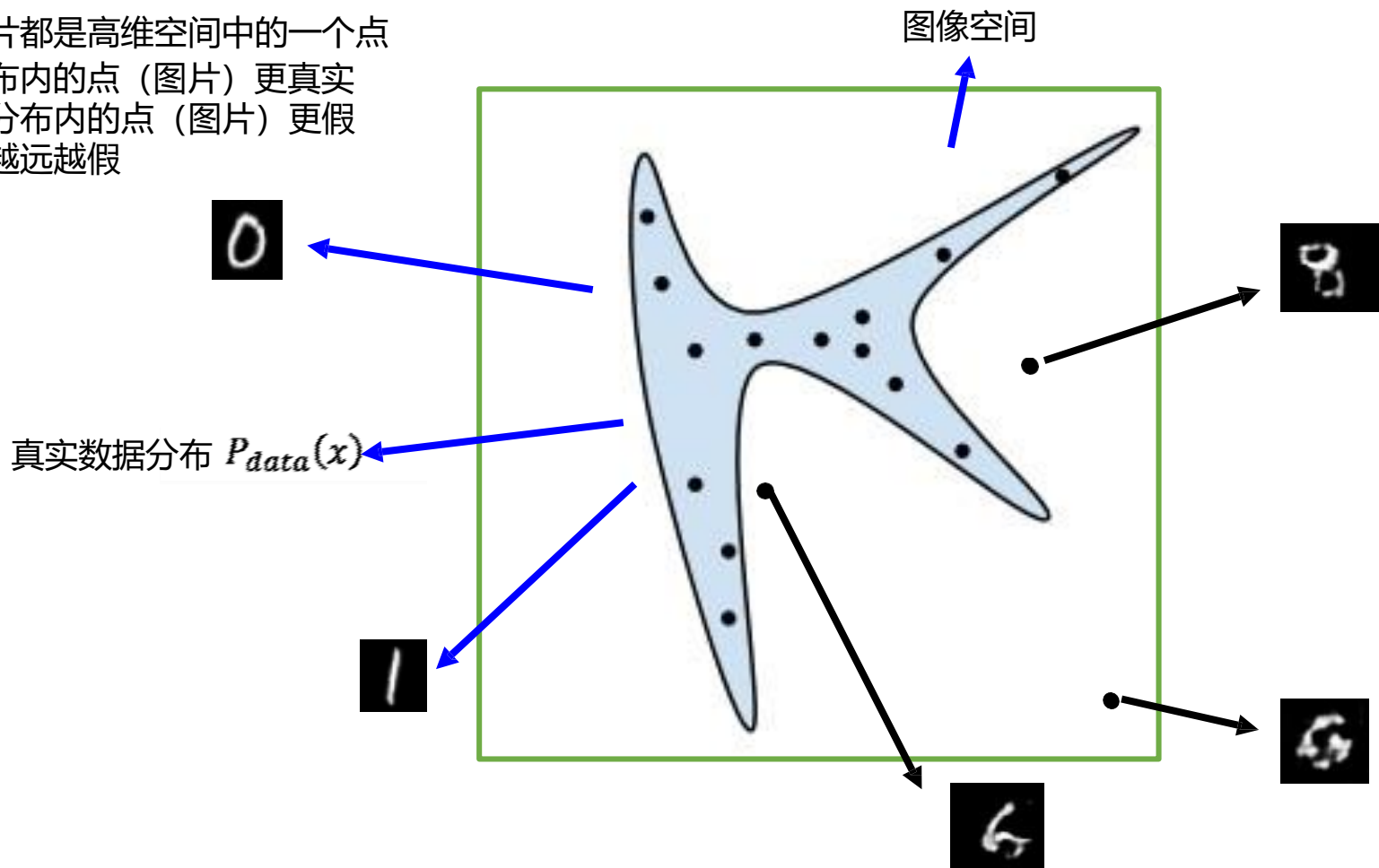
生成模型其实就是在学习（找到）数据的分布

每一个图片都是高维空间中的一个点

在真实分布内的点（图片）更真实

不在真实分布内的点（图片）更假

而且距离越远越假



GAN

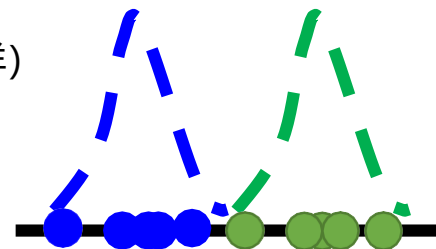
极大似然估计

我们有一个真实数据分布 $P_{\text{data}}(x)$ (不知道他的表达形式, 但是可以从中采样)

我们还有一个由参数 θ 定义的数据分布 $P_G(x, \theta)$

我们的目的是, 找到这样的参数, 使得 $P_G(x, \theta)$ 和 $P_{\text{data}}(x)$ 接近

例如 $P_G(x, \theta)$ 是高斯混合模型, 那么 θ 就是均值和方差



$$N(x|u, \sigma^2) = \frac{1}{\sqrt{2\pi\sigma^2}} \exp\left[-\frac{1}{2\sigma^2}(x - u)^2\right]$$

极大似然估计

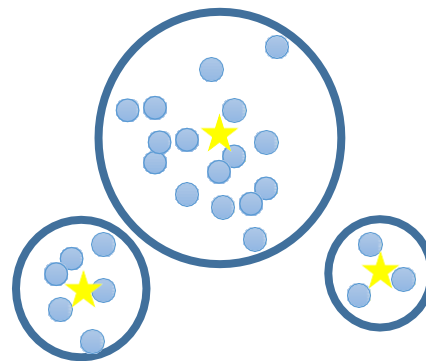
从数据集中采样 m 的真实样本 $\{x^1, x^2, \dots, x^m\}$

我们可以计算 $P_G(x^i, \theta)$

$$L = \prod_{i=1}^m P_G(x^i; \theta)$$

通过最大化 L , 找到最优的 θ^*

$P_G(x, \theta^*)$ 就是我们学习到的数据分布



GAN

极大似然估计=最小化KL

$$\begin{aligned}\theta^* &= \arg \max_{\theta} \prod_{i=1}^m P_G(x^i; \theta) &= \arg \max_{\theta} \log \prod_{i=1}^m P_G(x^i; \theta) \\ &= \arg \max_{\theta} \sum_{i=1}^m \log P_G(x^i; \theta) \\ &\approx \arg \max_{\theta} E_{x \sim P_{data}} [\log P_G(x; \theta)] &= \arg \max_{\theta} \int P_{data}(x) \log P_G(x; \theta) dx \\ &= \arg \max_{\theta} \int P_{data}(x) \log P_G(x; \theta) dx &- \int P_{data}(x) \log P_{data}(x) dx \\ &= \arg \min_{\theta} KL(P_{data} || P_G)\end{aligned}$$

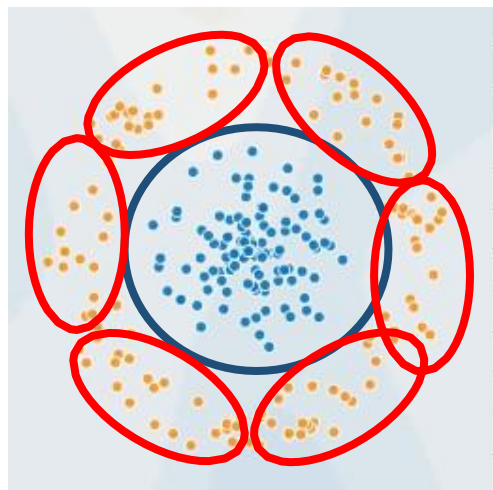
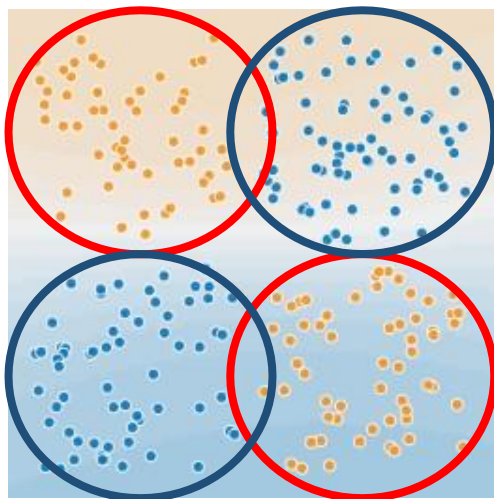
极大似然估计 等价于 **最小化生成数据分布和真实分布的KL散度**。
其中KL散度是两个分布的距离度量。

很多数据分布并不符合高斯分布的假设。

我们需要一个更通用的生成模型

GAN

极大似然估计=最小化KL



很多情况下，假设数据符合高斯分布是不合理的，数据分布是无法用公式显示的写出来的
因此用高斯模型去拟合数据分布，生成的结果总是不令人们以

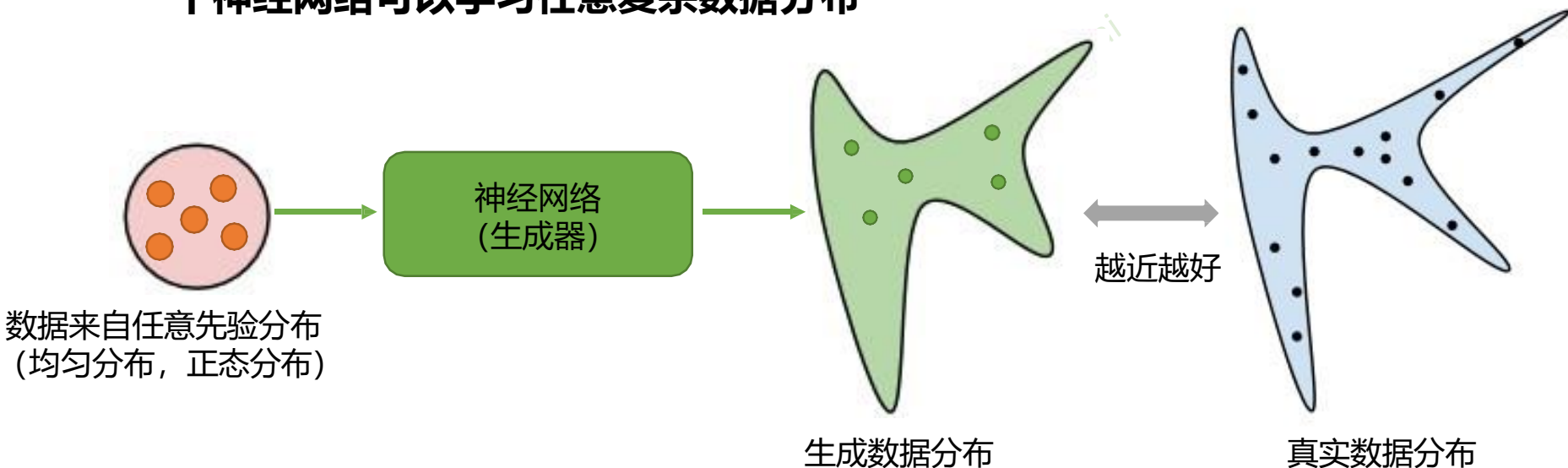
我们需要一个更通用的生成模型，可以拟合任意数据分布



GAN

生成器

一个神经网络可以学习任意复杂数据分布



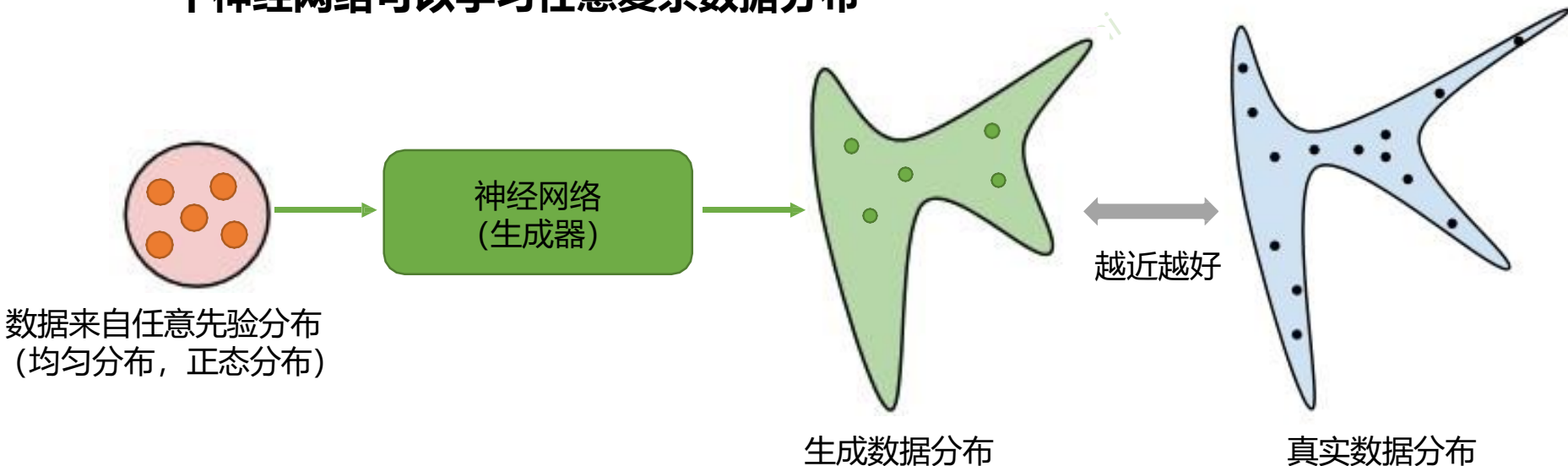
通过最小化生成数据分布 P_G 和真实数据分布 P_{data} 的KL散度, 就可以求出最优的生成器

$$G^* = \arg \min_G Div(P_{data} || P_G)$$

GAN

生成器

一个神经网络可以学习任意复杂数据分布



通过最小化生成数据分布 P_G 和真实数据分布 P_{data} 的KL散度, 就可以求出最优的生成器

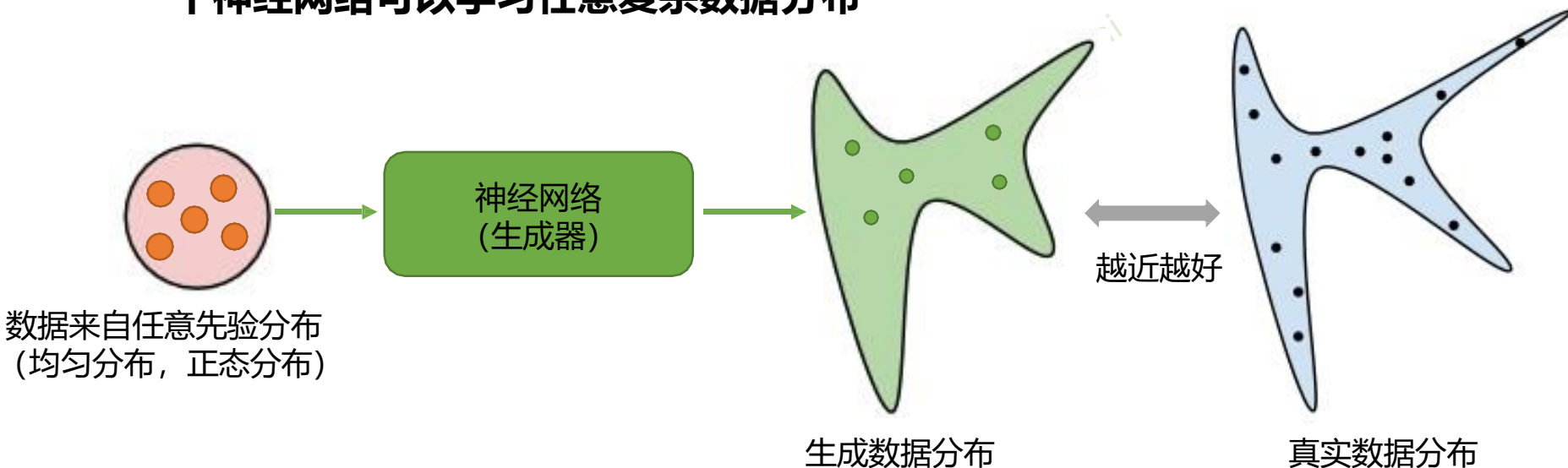
$$G^* = \arg \min_G Div(P_{data} || P_G)$$

**但是, 生成数据和真实数据分布的形式我们不知道,
无法计算散度, 无法优化生成器!**

GAN

生成器

一个神经网络可以学习任意复杂数据分布



通过最小化生成数据分布 P_G 和真实数据分布 P_{data} 的KL散度, 就可以求出最优的生成器

$$G^* = \arg \min_G Div(P_{data} || P_G)$$

GAN通过对抗训练, 间接计算出JS散度, 使得模型可以优化!

GAN

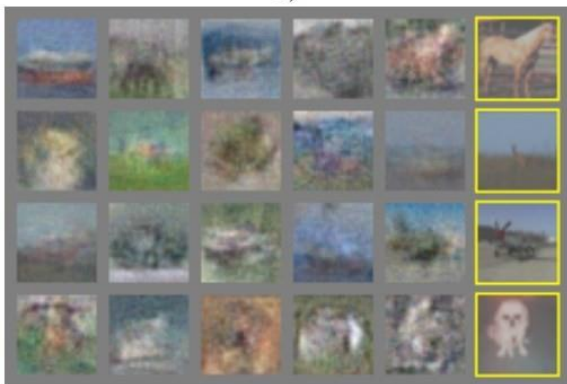
实验结果



a)



b)



c)



d)

生成图片比较逼真

生成图片并非简单记忆数据库中的图片

目录

CONTENTS

01

GAN

Generative Adversarial Nets, 生成式对抗网络

02

cGAN

Conditional GAN, 条件生成式对抗网络

03

DCGAN

Deep Convolutional GAN, 深度卷积生成式对抗网络

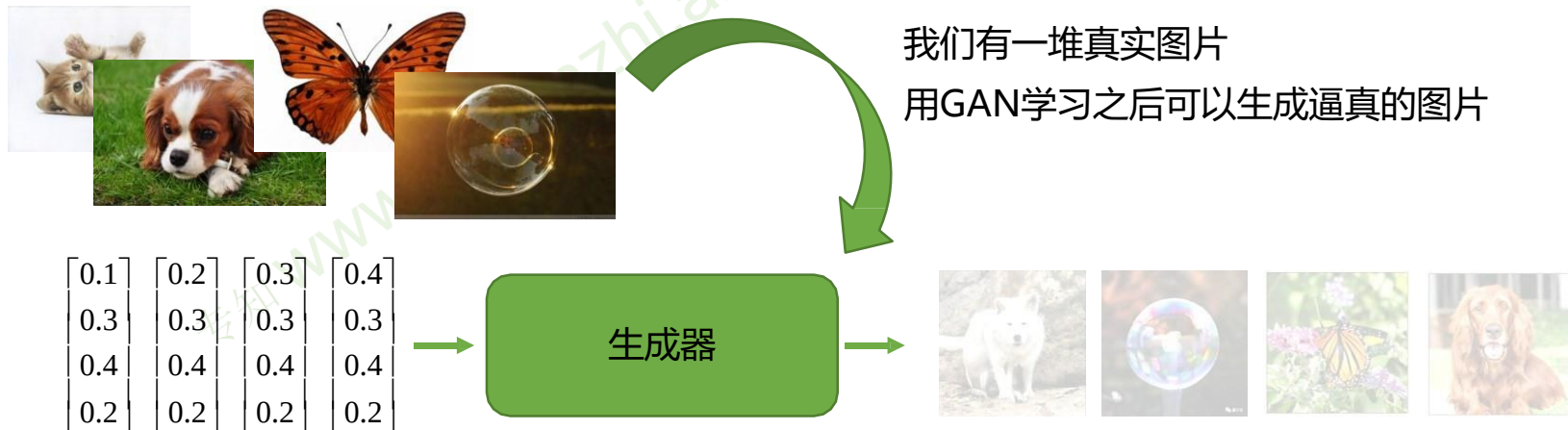
04

WGAN/WGAN-GP

Wasserstein GAN with Weight Clipping/ Gradient Penalty

cGAN

条件生成



实际应用中的问题：无法控制生成的类别！

cGAN

条件生成



我们有一堆真实图片
用GAN学习之后可以生成逼真的图片

0.1	0.2	0.3	0.4
0.3	0.3	0.3	0.3
0.4	0.4	0.4	0.4
0.2	0.2	0.2	0.2

生成器



为什么要控制生成数据的类别呢？

很多有意义的任务都需要学习条件生成模型！



条件



X



条件



X

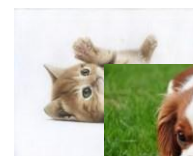
**条件生成
非常重要！！**

cGAN

条件生成



首先要有类别标签



猫



狗



蝴蝶



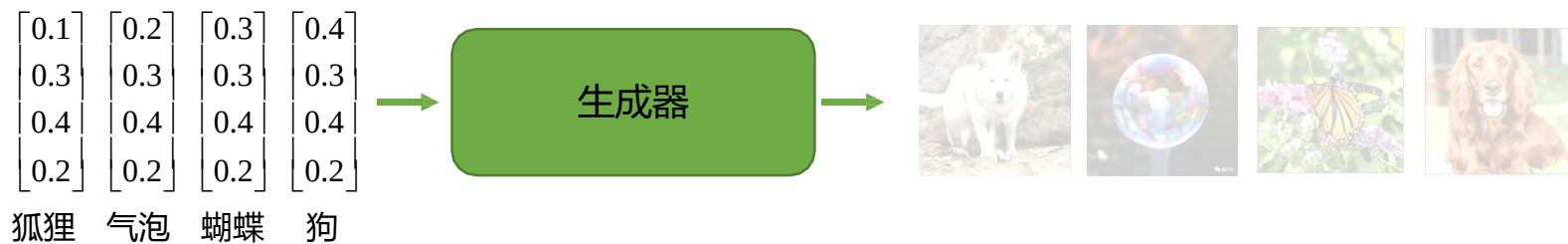
气泡

cGAN

条件生成



除了输入随机向量，还要数据类别

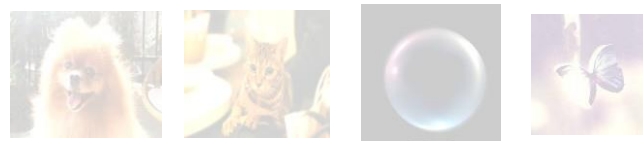
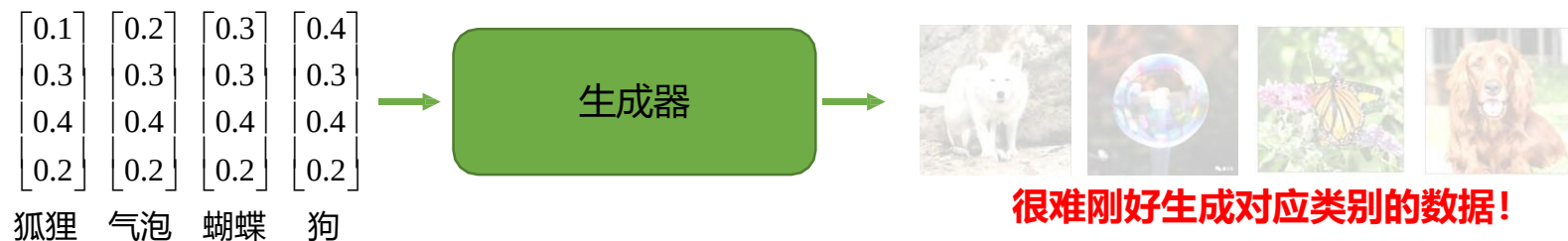


cGAN

条件生成



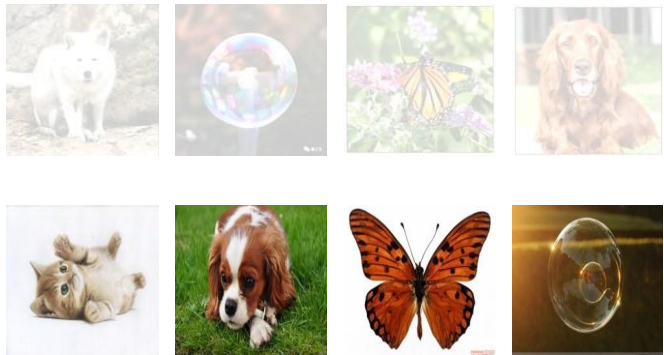
除了输入随机向量，还要数据类别



cGAN

条件生成

判别器：区分真假样本



判别器

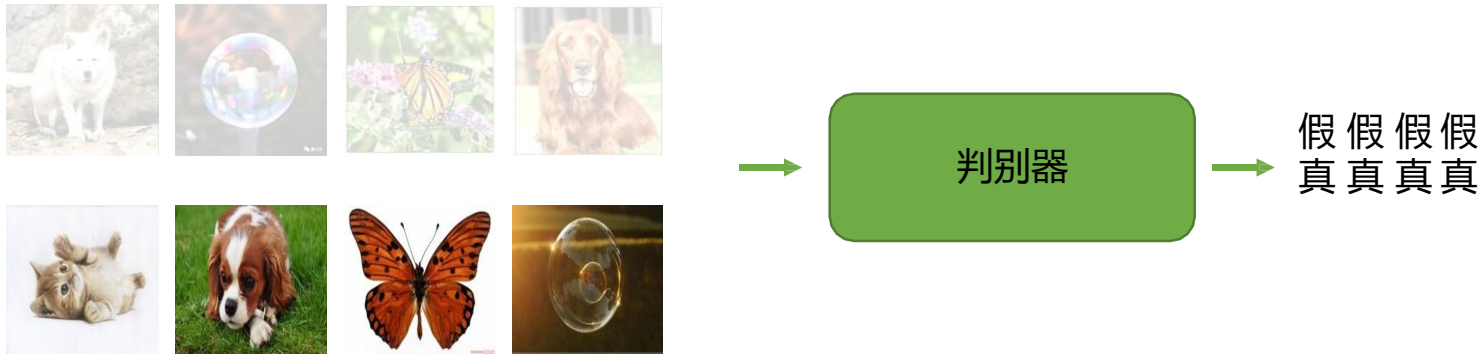


假 假 假 假
真 真 真 真

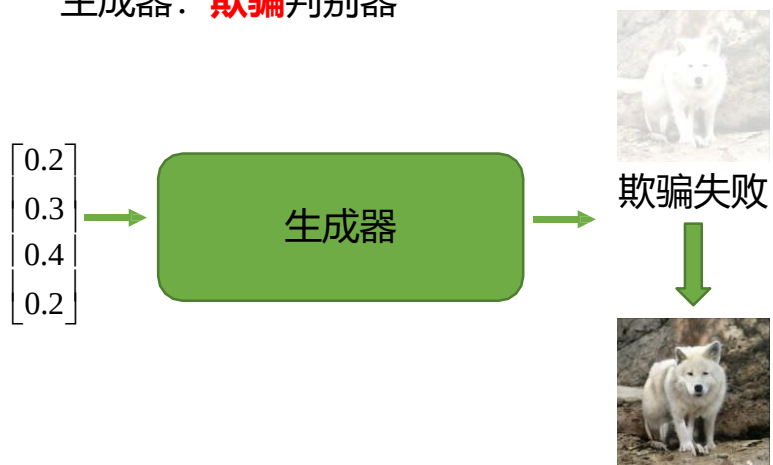
cGAN

条件生成

判别器：区分真假样本



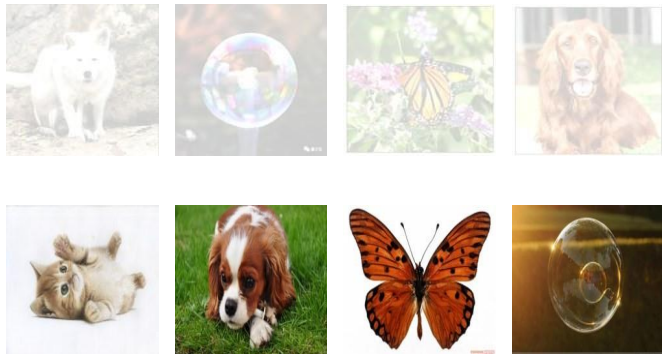
生成器：欺骗判别器



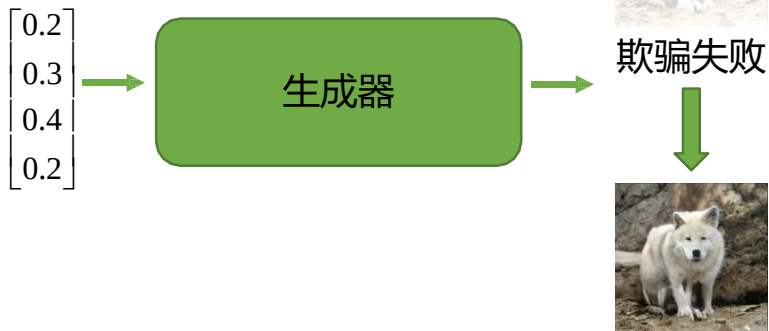
cGAN

条件生成

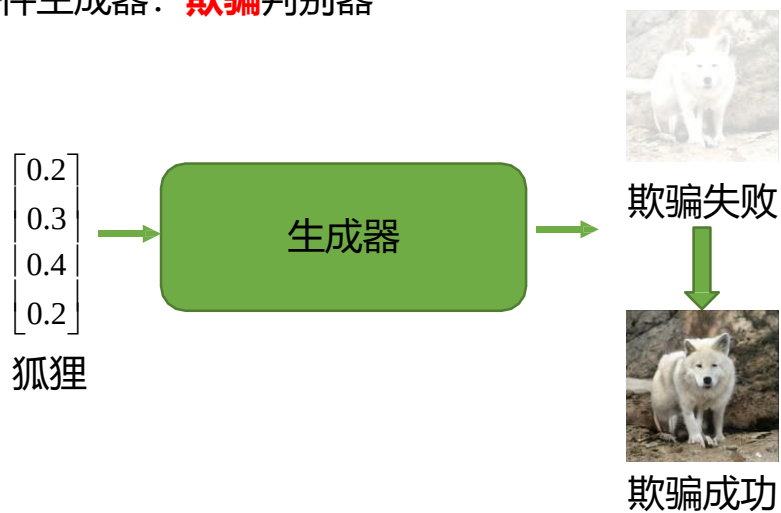
判别器：区分真假样本



生成器：欺骗判别器



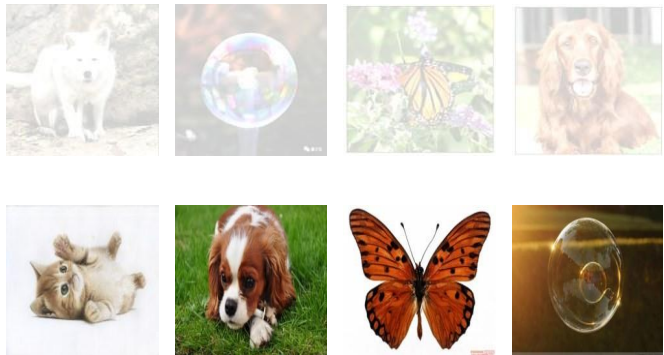
条件生成器：欺骗判别器



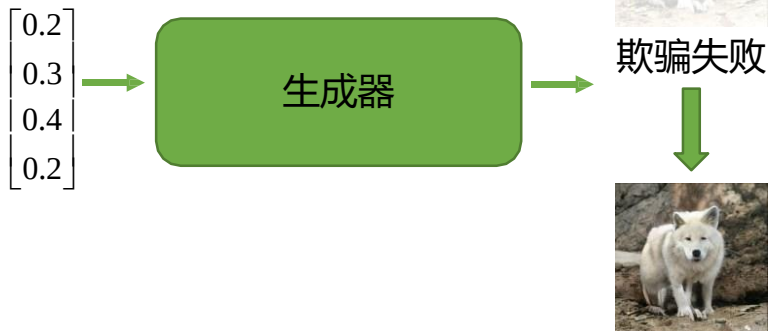
cGAN

条件生成

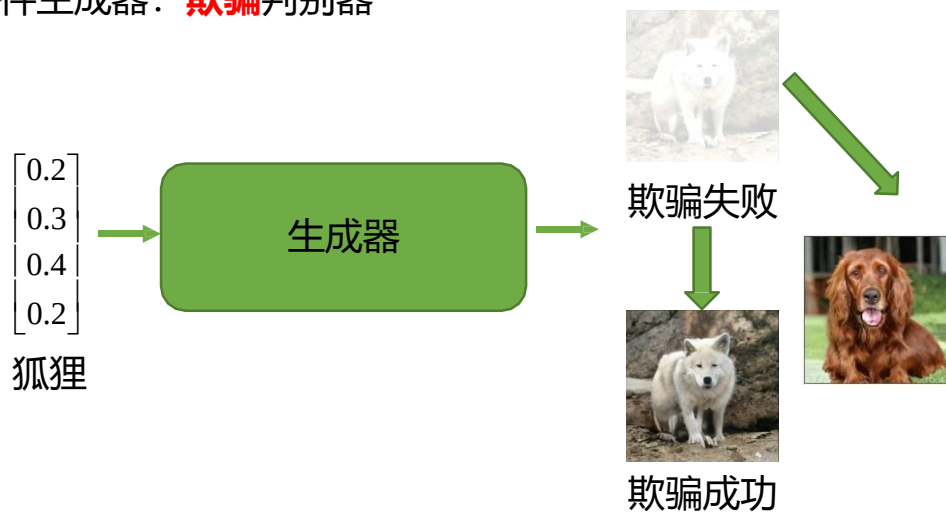
判别器：区分真假样本



生成器：欺骗判别器



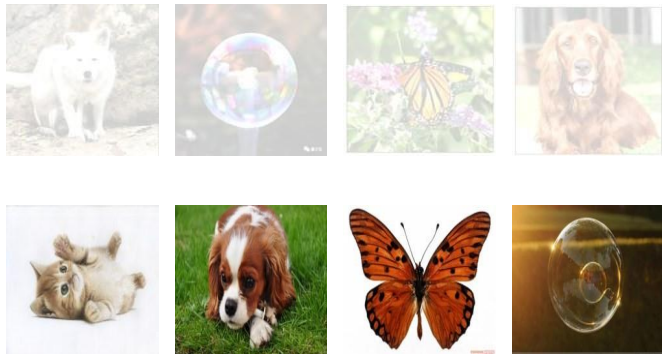
条件生成器：欺骗判别器



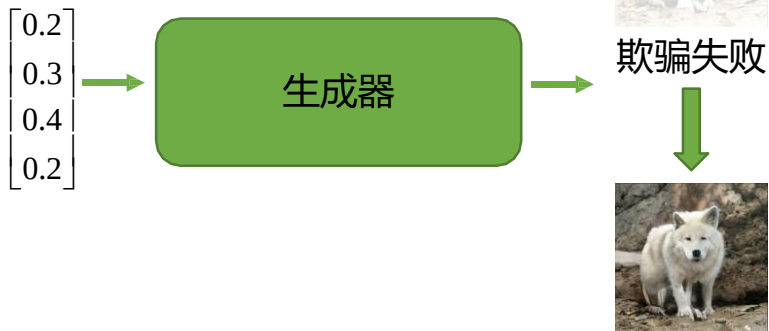
cGAN

条件生成

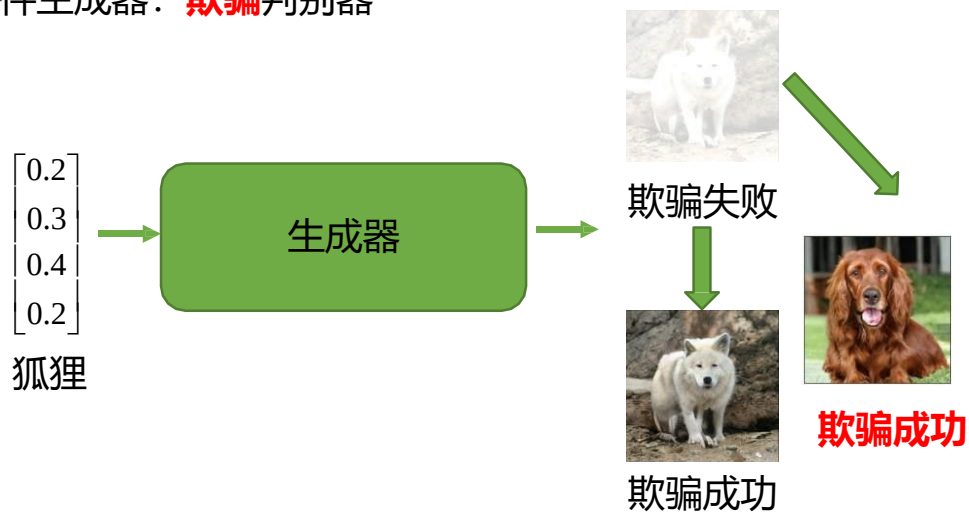
判别器：区分真假样本



生成器：欺骗判别器



条件生成器：欺骗判别器



大家请思考，问题出在哪呢？

改进判别器

判别器：区分真假样本

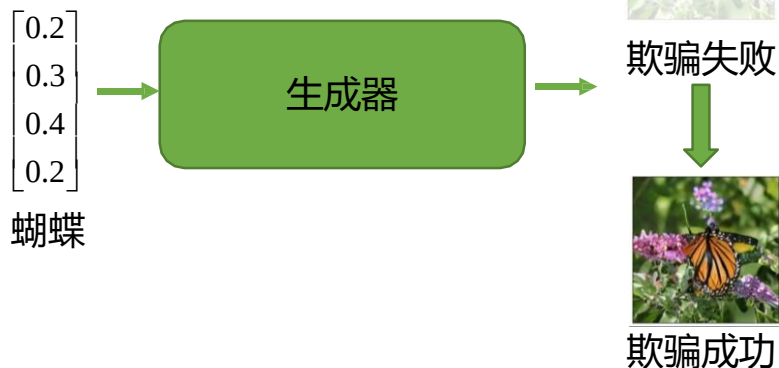


改进判别器

判别器：区分真假样本



生成器：欺骗判别器

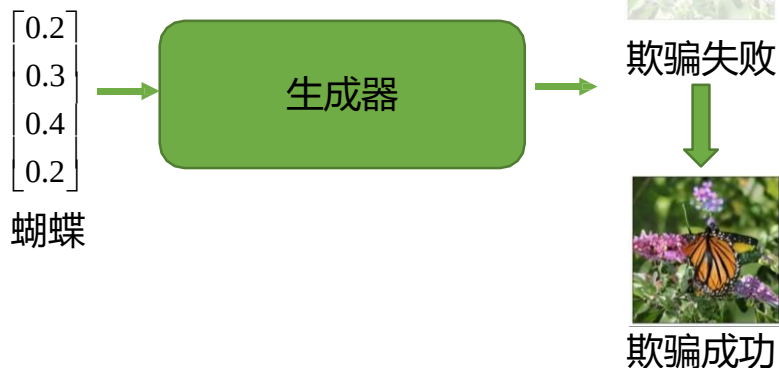


改进判别器

判别器：区分真假样本



生成器：欺骗判别器

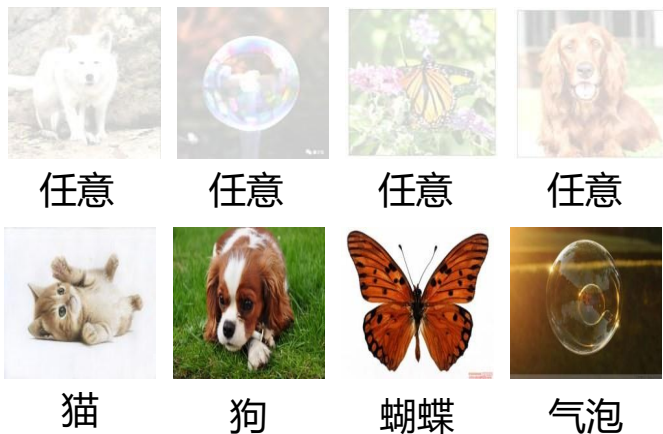


欺骗失败

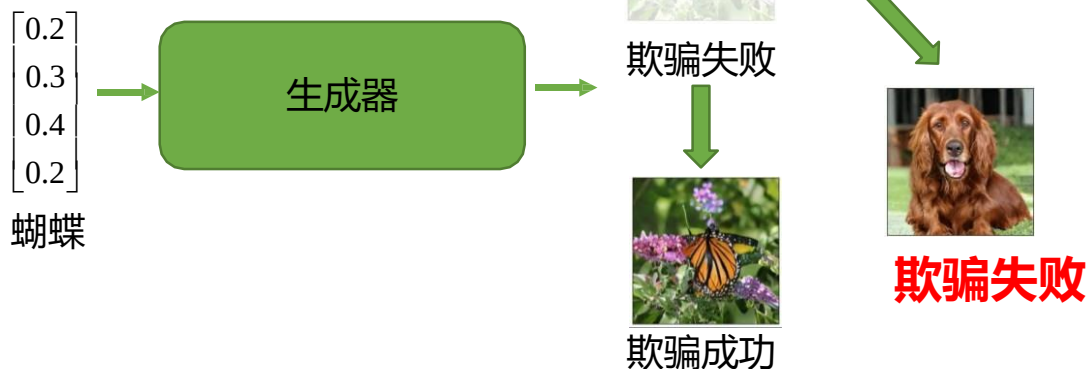
cGAN

改进判别器

判别器：区分真假样本



生成器：欺骗判别器

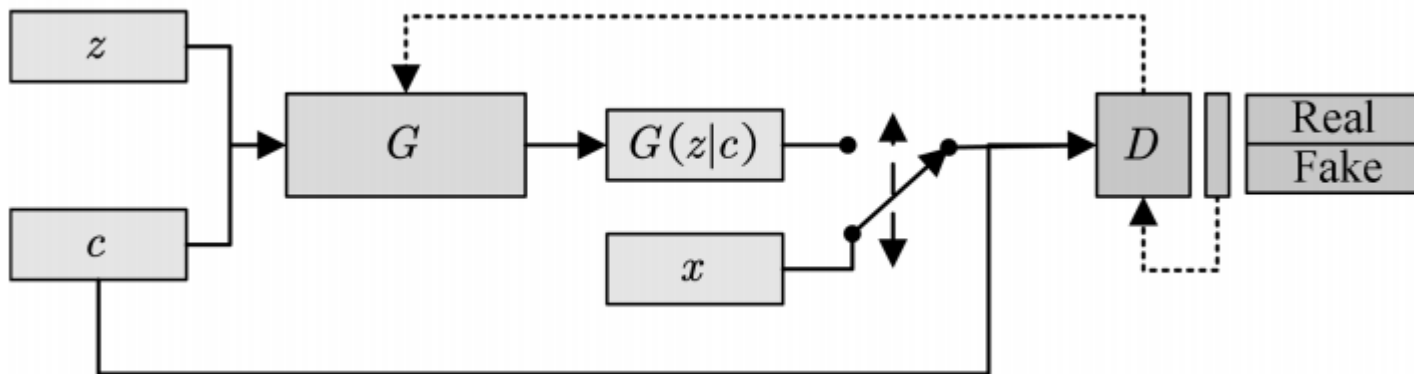




cGAN

cGAN

网络结构

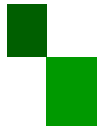


Improved CGAN 目标函数:

$$\min_G \max_D V(G, D) = E_{(x,y) \sim p_{data}} [\log D(x|y) + \log(1 - D(x|\hat{y}))] + E_{z \sim p_z} [\log(1 - D(G(z|y)))]$$

CGAN 目标函数:

$$\min_G \max_D V(G, D) = E_{(x,y) \sim p_{data}} [\log D(x|y)] + E_{z \sim p_z} [\log(1 - D(G(z|y)))]$$



GAN

cGAN

优化判别器



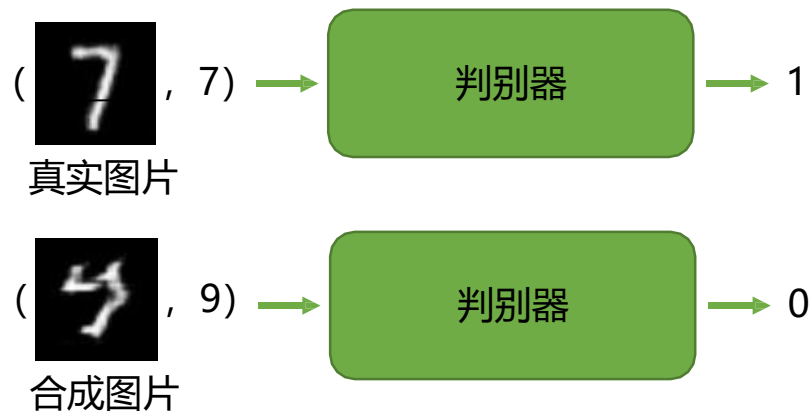
优化生成器



GAN

cGAN

优化判别器

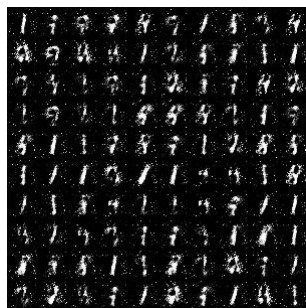
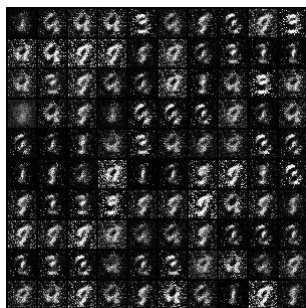


优化生成器

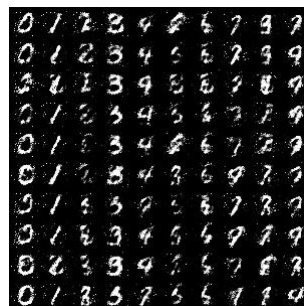
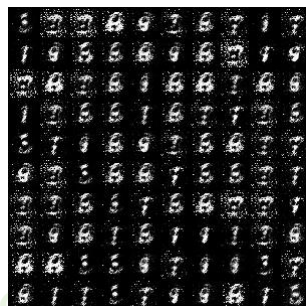
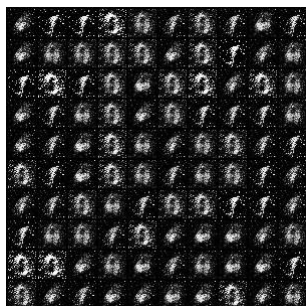


实验结果

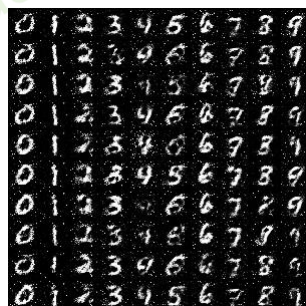
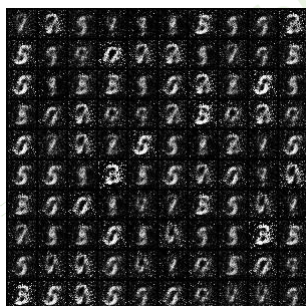
GAN



CGAN



Improved
CGAN



1 Epoch

5 Epoch

10 Epoch

20 Epoch

50 Epoch



Pytorch实现cGAN

准备工作

定义网络结构：判别器和生成器的网络结构

定义目标函数：根据预测标签和真实标签，计算损失

定义优化器：优化器，自动计算梯度并优化网络权重

数据加载：加载数据集

判别器GAN

判别器cGAN

生成器GAN

生成器cGAN

其他准备工作

训练直到收敛

加载数据

提取一个batch的真实数据

生成一个batch的虚假数据

计算损失并优化模型

计算判别器损失，并优化判别器

计算生成器损失，并优化生成器

训练过程

目录

CONTENTS

01

GAN

Generative Adversarial Nets, 生成式对抗网络

02

cGAN

Conditional GAN, 条件生成式对抗网络

03

DCGAN

Deep Convolutional GAN, 深度卷积生成式对抗网络

04

WGAN/WGAN-GP

Wasserstein GAN with Weight Clipping/ Gradient Penalty

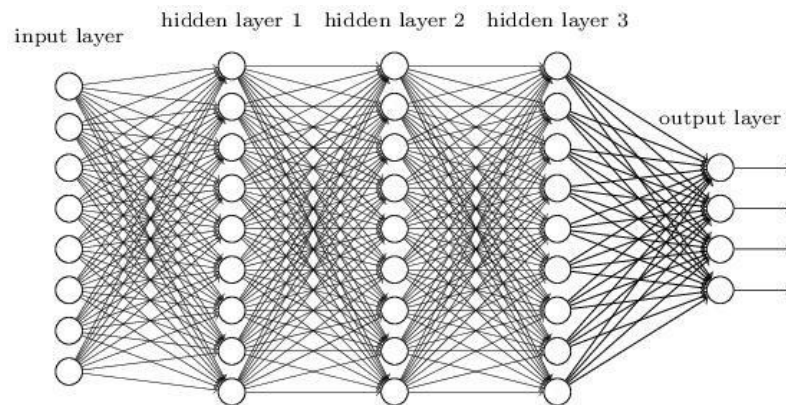
DCGAN

DCGAN

原始GAN，使用全连接网络作为判别器和生成器

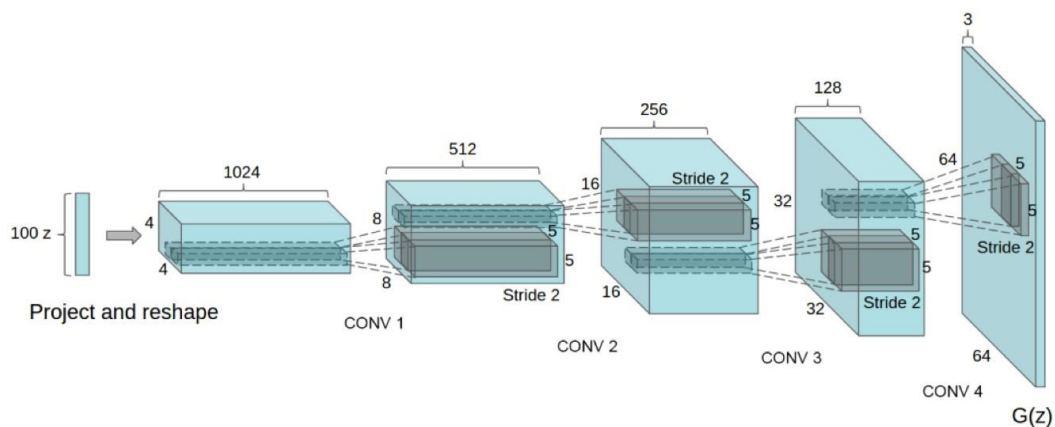
不利于建模图像信息

参数量大，需要大量的计算资源，难以优化



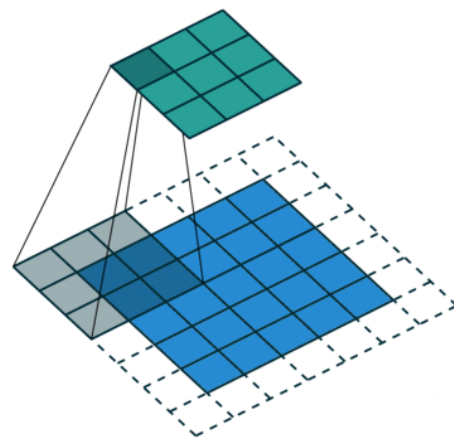
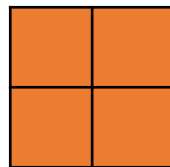
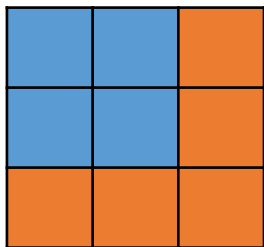
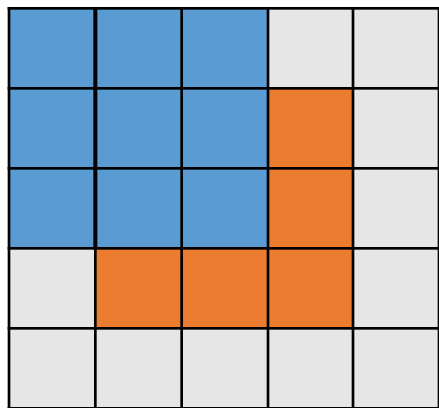
DCGAN，使用卷积神经网络作为判别器和生成器

通过大量的工程实践，经验性地提出一系列的网络结构和优化策略，来有效的建模图像数据



DCGAN

判别器：滑动卷积



灰色：Padding

橘色：Feature Map

蓝色：Kernel

橘色：Feature Map

蓝色：Pooling

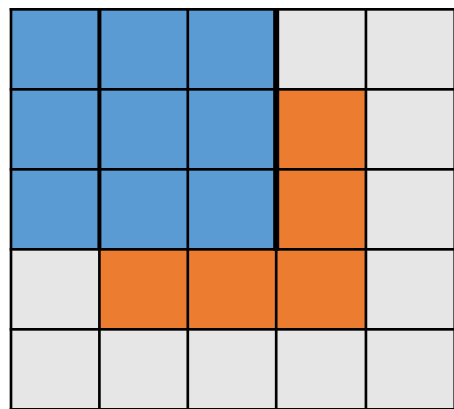
橘色：Feature Map

通过Pooling下采样，Pooling是不可学习的，这可能造成GAN训练困难

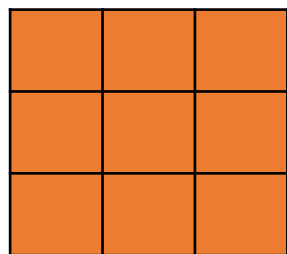
使用滑动卷积 (Stride Convolution)
所有参数都可学习
实验发现效果更好

DCGAN

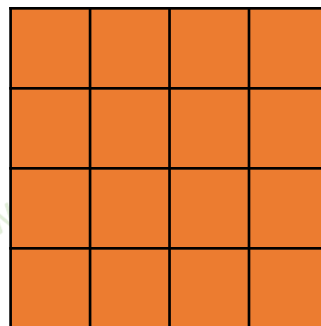
生成器：滑动反卷积



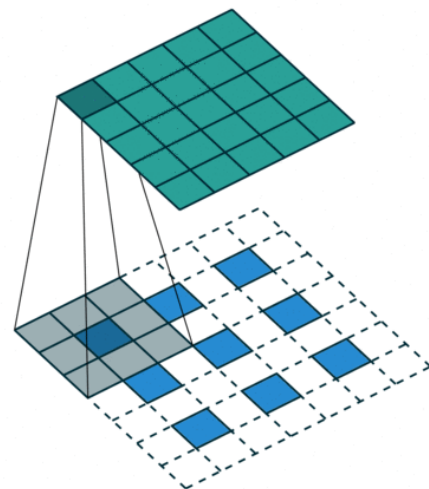
灰色：Padding
橘色：Feature Map
蓝色：Kernel



通过插值
上采样



比如：线性插值



通过插值法上采样，
插值方法是固定的，不可学习的，
这可能给训练造成困难

使用滑动反卷积所有参数都可学习
实验发现效果更好



DCGAN

批归一化

IID独立同分布假设，就是假设训练数据和测试数据是满足相同分布的，这是通过训练数据获得的模型能够在测试集获得好的效果的一个基本保障

BatchNorm就是在深度神经网络训练过程中使得每一层神经网络的输入保持相同分布的

优势：

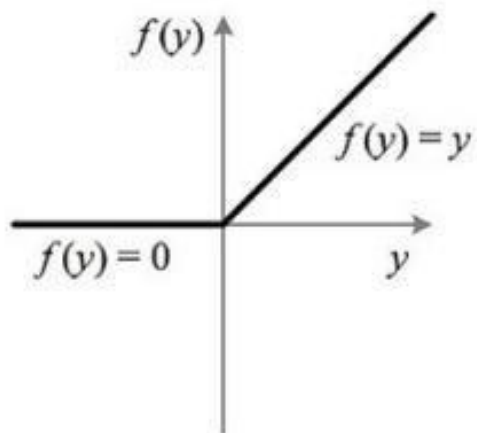
加速神经网络收敛

减小神经网络参数对于初始化的依赖

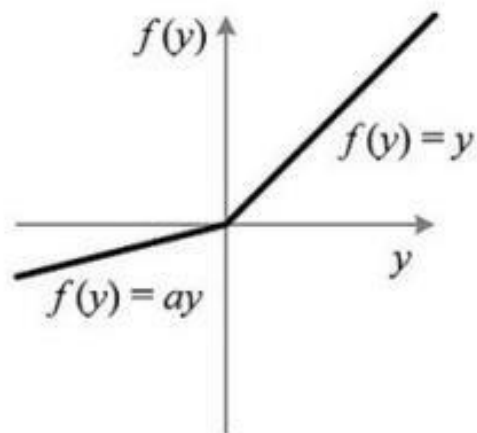
实验发现效果更好~

DCGAN

激活函数



$$ReLU(x) = \begin{cases} x & \text{if } x > 0 \\ 0 & \text{if } x \leq 0 \end{cases}$$



$$PReLU(x_i) = \begin{cases} x_i & \text{if } x_i > 0 \\ a_i x_i & \text{if } x_i \leq 0 \end{cases}$$

i 表示不同的通道

Relu可以缓解梯度消失

PRelu是Relu的改进版本，当 $x < 0$ 时仍然可以提供梯度



DCGAN

DCGAN

网络结构（判别器）

- 使用滑动卷积（strided convolution）
- 除了输入层，全部使用批归一化
- 使用 Leaky ReLu 激活函数
- 除了最后一层，不使用全连接层

网络结构（生成器）

- 使用 fractional strided convolution
- 除了输出层，全部使用批归一化
- 使用 ReLu 激活函数，最后一层使用tanh激活函数

滑动卷积、滑动翻卷机：使得判别器和生成器可以学习自己的上采样和下采样策略

批归一化（Batch Normalization）：训练更稳定

Tanh激活函数：更快的学习到真实数据的颜色空间

训练策略

- 数据预处理：所有输入数据归一化到 $[-1, 1]$
- 激活函数：Leaky Relu的斜率设置为0.2
- 初始化：使用均值为0，标准差为0.02的正态分布初始化网络参数
- 优化器：使用Adam优化器，学习率为0.0002， $\beta_1=0.5$ ， $\beta_2=0.999$

DCGAN

实验结果

DCGAN



1 Epoch



2 Epoch



5 Epoch

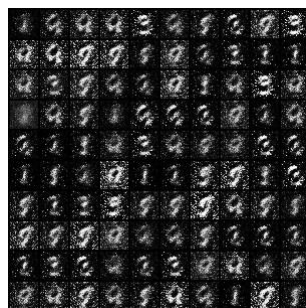


10 Epoch

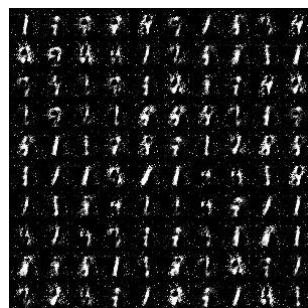


20 Epoch

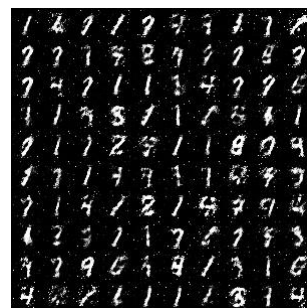
GAN



1 Epoch



5 Epoch



10 Epoch



20 Epoch



50 Epoch

DCGAN

Pytorch实现DCGAN

准备工作

定义网络结构：判别器和生成器的网络结构

定义目标函数：根据预测标签和真实标签，计算损失

定义优化器：优化器，自动计算梯度并优化网络权重

数据加载：加载数据集

```
class Discriminator(nn.Module):
    """判别器网络结构"""
    def __init__(self):
        super(Discriminator, self).__init__()

        # 定义卷积层
        conv = []
        # 第一个卷积层，通道32，3x3卷积核
        conv.append(nn.Conv2d(1, 32, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第二个卷积层，通道64，3x3卷积核
        conv.append(nn.Conv2d(32, 64, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第三个卷积层，通道128，3x3卷积核
        conv.append(nn.Conv2d(64, 128, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第四个卷积层，通道256，3x3卷积核
        conv.append(nn.Conv2d(128, 256, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第五个卷积层，通道512，3x3卷积核
        conv.append(nn.Conv2d(256, 512, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第六个卷积层，通道1024，3x3卷积核
        conv.append(nn.Conv2d(512, 1024, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第七个卷积层，通道2048，3x3卷积核
        conv.append(nn.Conv2d(1024, 2048, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第八个卷积层，通道4096，3x3卷积核
        conv.append(nn.Conv2d(2048, 4096, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第九个卷积层，通道8192，3x3卷积核
        conv.append(nn.Conv2d(4096, 8192, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第十个卷积层，通道16384，3x3卷积核
        conv.append(nn.Conv2d(8192, 16384, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第十一个卷积层，通道32768，3x3卷积核
        conv.append(nn.Conv2d(16384, 32768, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第十二个卷积层，通道65536，3x3卷积核
        conv.append(nn.Conv2d(32768, 65536, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第十三个卷积层，通道131072，3x3卷积核
        conv.append(nn.Conv2d(65536, 131072, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第十四个卷积层，通道262144，3x3卷积核
        conv.append(nn.Conv2d(131072, 262144, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第十五个卷积层，通道524288，3x3卷积核
        conv.append(nn.Conv2d(262144, 524288, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第十六个卷积层，通道1048576，3x3卷积核
        conv.append(nn.Conv2d(524288, 1048576, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第十七个卷积层，通道2097152，3x3卷积核
        conv.append(nn.Conv2d(1048576, 2097152, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第十八个卷积层，通道4194304，3x3卷积核
        conv.append(nn.Conv2d(2097152, 4194304, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第十九个卷积层，通道8388608，3x3卷积核
        conv.append(nn.Conv2d(4194304, 8388608, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第二十个卷积层，通道16777216，3x3卷积核
        conv.append(nn.Conv2d(8388608, 16777216, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第二十一卷积层，通道33554432，3x3卷积核
        conv.append(nn.Conv2d(16777216, 33554432, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第二十二个卷积层，通道67108864，3x3卷积核
        conv.append(nn.Conv2d(33554432, 67108864, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第二十三个卷积层，通道134217728，3x3卷积核
        conv.append(nn.Conv2d(67108864, 134217728, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第二十四个卷积层，通道268435456，3x3卷积核
        conv.append(nn.Conv2d(134217728, 268435456, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第二十五个卷积层，通道536870912，3x3卷积核
        conv.append(nn.Conv2d(268435456, 536870912, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第二十六个卷积层，通道1073741824，3x3卷积核
        conv.append(nn.Conv2d(536870912, 1073741824, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第二十七个卷积层，通道2147483648，3x3卷积核
        conv.append(nn.Conv2d(1073741824, 2147483648, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第二十八个卷积层，通道4294967296，3x3卷积核
        conv.append(nn.Conv2d(2147483648, 4294967296, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第二十九个卷积层，通道8589934592，3x3卷积核
        conv.append(nn.Conv2d(4294967296, 8589934592, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第三十个卷积层，通道17179869184，3x3卷积核
        conv.append(nn.Conv2d(8589934592, 17179869184, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第三十一个卷积层，通道34359738368，3x3卷积核
        conv.append(nn.Conv2d(17179869184, 34359738368, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第三十二个卷积层，通道68719476736，3x3卷积核
        conv.append(nn.Conv2d(34359738368, 68719476736, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第三十三个卷积层，通道137438953472，3x3卷积核
        conv.append(nn.Conv2d(68719476736, 137438953472, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第三十四个卷积层，通道274877906944，3x3卷积核
        conv.append(nn.Conv2d(137438953472, 274877906944, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第三十五个卷积层，通道549755813888，3x3卷积核
        conv.append(nn.Conv2d(274877906944, 549755813888, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第三十六个卷积层，通道1099511627776，3x3卷积核
        conv.append(nn.Conv2d(549755813888, 1099511627776, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第三十七个卷积层，通道2199023255552，3x3卷积核
        conv.append(nn.Conv2d(1099511627776, 2199023255552, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第三十八个卷积层，通道4398046511104，3x3卷积核
        conv.append(nn.Conv2d(2199023255552, 4398046511104, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第三十九个卷积层，通道8796093022208，3x3卷积核
        conv.append(nn.Conv2d(4398046511104, 8796093022208, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第四十个卷积层，通道17592186044416，3x3卷积核
        conv.append(nn.Conv2d(8796093022208, 17592186044416, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第四十一个卷积层，通道35184372088832，3x3卷积核
        conv.append(nn.Conv2d(17592186044416, 35184372088832, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第四十二个卷积层，通道70368744177664，3x3卷积核
        conv.append(nn.Conv2d(35184372088832, 70368744177664, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第四十三个卷积层，通道140737488355328，3x3卷积核
        conv.append(nn.Conv2d(70368744177664, 140737488355328, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第四十四个卷积层，通道281474976710656，3x3卷积核
        conv.append(nn.Conv2d(140737488355328, 281474976710656, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第四十五个卷积层，通道562949953421312，3x3卷积核
        conv.append(nn.Conv2d(281474976710656, 562949953421312, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第四十六个卷积层，通道1125899906842624，3x3卷积核
        conv.append(nn.Conv2d(562949953421312, 1125899906842624, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第四十七个卷积层，通道2251799813685248，3x3卷积核
        conv.append(nn.Conv2d(1125899906842624, 2251799813685248, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第四十八个卷积层，通道4503599627370496，3x3卷积核
        conv.append(nn.Conv2d(2251799813685248, 4503599627370496, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第四十九个卷积层，通道9007199254740992，3x3卷积核
        conv.append(nn.Conv2d(4503599627370496, 9007199254740992, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第五十个卷积层，通道18014398509481984，3x3卷积核
        conv.append(nn.Conv2d(9007199254740992, 18014398509481984, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第五十一个卷积层，通道36028797018963968，3x3卷积核
        conv.append(nn.Conv2d(18014398509481984, 36028797018963968, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第五十二个卷积层，通道72057594037927936，3x3卷积核
        conv.append(nn.Conv2d(36028797018963968, 72057594037927936, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第五十三个卷积层，通道144115188075855872，3x3卷积核
        conv.append(nn.Conv2d(72057594037927936, 144115188075855872, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第五十四个卷积层，通道288230376151711744，3x3卷积核
        conv.append(nn.Conv2d(144115188075855872, 288230376151711744, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第五十五个卷积层，通道576460752303423488，3x3卷积核
        conv.append(nn.Conv2d(288230376151711744, 576460752303423488, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第五十六个卷积层，通道1152921504606846976，3x3卷积核
        conv.append(nn.Conv2d(576460752303423488, 1152921504606846976, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第五十七个卷积层，通道2305843009213693952，3x3卷积核
        conv.append(nn.Conv2d(1152921504606846976, 2305843009213693952, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第五十八个卷积层，通道4611686018427387904，3x3卷积核
        conv.append(nn.Conv2d(2305843009213693952, 4611686018427387904, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第五十九个卷积层，通道9223372036854775808，3x3卷积核
        conv.append(nn.Conv2d(4611686018427387904, 9223372036854775808, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第六十个卷积层，通道18446744073709551616，3x3卷积核
        conv.append(nn.Conv2d(9223372036854775808, 18446744073709551616, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第六十一个卷积层，通道36893488147419103232，3x3卷积核
        conv.append(nn.Conv2d(18446744073709551616, 36893488147419103232, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第六十二个卷积层，通道73786976294838206464，3x3卷积核
        conv.append(nn.Conv2d(36893488147419103232, 73786976294838206464, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第六十三个卷积层，通道147573952589676412928，3x3卷积核
        conv.append(nn.Conv2d(73786976294838206464, 147573952589676412928, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第六十四个卷积层，通道295147905179352825856，3x3卷积核
        conv.append(nn.Conv2d(147573952589676412928, 295147905179352825856, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第六十五个卷积层，通道590295810358705651712，3x3卷积核
        conv.append(nn.Conv2d(295147905179352825856, 590295810358705651712, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第六十六个卷积层，通道1180591620717411303424，3x3卷积核
        conv.append(nn.Conv2d(590295810358705651712, 1180591620717411303424, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第六十七个卷积层，通道2361183241434822606848，3x3卷积核
        conv.append(nn.Conv2d(1180591620717411303424, 2361183241434822606848, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第六十八个卷积层，通道4722366482869645213696，3x3卷积核
        conv.append(nn.Conv2d(2361183241434822606848, 4722366482869645213696, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第六十九个卷积层，通道9444732965739290427392，3x3卷积核
        conv.append(nn.Conv2d(4722366482869645213696, 9444732965739290427392, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第七十个卷积层，通道18889465931478580854784，3x3卷积核
        conv.append(nn.Conv2d(9444732965739290427392, 18889465931478580854784, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第七十一个卷积层，通道37778931862957161709568，3x3卷积核
        conv.append(nn.Conv2d(18889465931478580854784, 37778931862957161709568, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第七十二个卷积层，通道75557863725914323419136，3x3卷积核
        conv.append(nn.Conv2d(37778931862957161709568, 75557863725914323419136, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第七十三个卷积层，通道151115727451828646838272，3x3卷积核
        conv.append(nn.Conv2d(75557863725914323419136, 151115727451828646838272, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第七十四个卷积层，通道302231454903657293676544，3x3卷积核
        conv.append(nn.Conv2d(151115727451828646838272, 302231454903657293676544, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第七十五个卷积层，通道604462909807314587353088，3x3卷积核
        conv.append(nn.Conv2d(302231454903657293676544, 604462909807314587353088, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第七十六个卷积层，通道1208925819614629174706176，3x3卷积核
        conv.append(nn.Conv2d(604462909807314587353088, 1208925819614629174706176, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第七十七个卷积层，通道2417851639229258349412352，3x3卷积核
        conv.append(nn.Conv2d(1208925819614629174706176, 2417851639229258349412352, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第七十八个卷积层，通道4835703278458516698824704，3x3卷积核
        conv.append(nn.Conv2d(2417851639229258349412352, 4835703278458516698824704, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第七十九个卷积层，通道9671406556917033397649408，3x3卷积核
        conv.append(nn.Conv2d(4835703278458516698824704, 9671406556917033397649408, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第八十个卷积层，通道19342813113834066795298816，3x3卷积核
        conv.append(nn.Conv2d(9671406556917033397649408, 19342813113834066795298816, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第八十一个卷积层，通道38685626227668133590597632，3x3卷积核
        conv.append(nn.Conv2d(19342813113834066795298816, 38685626227668133590597632, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第八十二个卷积层，通道77371252455336267181195264，3x3卷积核
        conv.append(nn.Conv2d(38685626227668133590597632, 77371252455336267181195264, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第八十三个卷积层，通道154742504910672534362390528，3x3卷积核
        conv.append(nn.Conv2d(77371252455336267181195264, 154742504910672534362390528, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第八十四个卷积层，通道309485009821345068724781056，3x3卷积核
        conv.append(nn.Conv2d(154742504910672534362390528, 309485009821345068724781056, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第八十五个卷积层，通道618970019642690137449562112，3x3卷积核
        conv.append(nn.Conv2d(309485009821345068724781056, 618970019642690137449562112, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第八十六个卷积层，通道1237940039285380274899124224，3x3卷积核
        conv.append(nn.Conv2d(618970019642690137449562112, 1237940039285380274899124224, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第八十七个卷积层，通道2475880078570760549798248448，3x3卷积核
        conv.append(nn.Conv2d(1237940039285380274899124224, 2475880078570760549798248448, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第八十八个卷积层，通道4951760157141521099596496896，3x3卷积核
        conv.append(nn.Conv2d(2475880078570760549798248448, 4951760157141521099596496896, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第八十九个卷积层，通道9903520314283042199192993792，3x3卷积核
        conv.append(nn.Conv2d(4951760157141521099596496896, 9903520314283042199192993792, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第九十个卷积层，通道19807040628566084398385987584，3x3卷积核
        conv.append(nn.Conv2d(9903520314283042199192993792, 19807040628566084398385987584, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第九十一个卷积层，通道39614081257132168796771975168，3x3卷积核
        conv.append(nn.Conv2d(19807040628566084398385987584, 39614081257132168796771975168, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第九十二个卷积层，通道79228162514264337593543950336，3x3卷积核
        conv.append(nn.Conv2d(39614081257132168796771975168, 79228162514264337593543950336, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第九十三个卷积层，通道158456325028528675187087900672，3x3卷积核
        conv.append(nn.Conv2d(79228162514264337593543950336, 158456325028528675187087900672, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第九十四个卷积层，通道316912650057057350374175801344，3x3卷积核
        conv.append(nn.Conv2d(158456325028528675187087900672, 316912650057057350374175801344, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第九十五个卷积层，通道633825300114114700748351602688，3x3卷积核
        conv.append(nn.Conv2d(316912650057057350374175801344, 633825300114114700748351602688, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第九十六个卷积层，通道1267650600228229401496703205376，3x3卷积核
        conv.append(nn.Conv2d(633825300114114700748351602688, 1267650600228229401496703205376, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第九十七个卷积层，通道2535301200456458802993406410752，3x3卷积核
        conv.append(nn.Conv2d(1267650600228229401496703205376, 2535301200456458802993406410752, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第九十八个卷积层，通道5070602400912917605986812821504，3x3卷积核
        conv.append(nn.Conv2d(2535301200456458802993406410752, 5070602400912917605986812821504, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第九十九个卷积层，通道10141204801825835211973625643008，3x3卷积核
        conv.append(nn.Conv2d(5070602400912917605986812821504, 10141204801825835211973625643008, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第一百个卷积层，通道20282409603651670423947251286016，3x3卷积核
        conv.append(nn.Conv2d(10141204801825835211973625643008, 20282409603651670423947251286016, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第一百零一个卷积层，通道40564819207303340847894502572032，3x3卷积核
        conv.append(nn.Conv2d(20282409603651670423947251286016, 40564819207303340847894502572032, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第一百零二个卷积层，通道81129638414606681695789005144064，3x3卷积核
        conv.append(nn.Conv2d(40564819207303340847894502572032, 81129638414606681695789005144064, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第一百零三个卷积层，通道162259276829213363391578010288128，3x3卷积核
        conv.append(nn.Conv2d(81129638414606681695789005144064, 162259276829213363391578010288128, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第一百零四个卷积层，通道324518553658426726783156020576256，3x3卷积核
        conv.append(nn.Conv2d(162259276829213363391578010288128, 324518553658426726783156020576256, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第一百零五个卷积层，通道649037107316853453566312041152512，3x3卷积核
        conv.append(nn.Conv2d(324518553658426726783156020576256, 649037107316853453566312041152512, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第一百零六个卷积层，通道1298074214633706907132624082305024，3x3卷积核
        conv.append(nn.Conv2d(649037107316853453566312041152512, 1298074214633706907132624082305024, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第一百零七个卷积层，通道2596148429267413814265248164610048，3x3卷积核
        conv.append(nn.Conv2d(1298074214633706907132624082305024, 2596148429267413814265248164610048, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第一百零八个卷积层，通道5192296858534827628530496329220096，3x3卷积核
        conv.append(nn.Conv2d(2596148429267413814265248164610048, 5192296858534827628530496329220096, kernel=3, stride=1, padding=1))
        conv.append(nn.LeakyReLU(0.2, inplace=True))
        # 第一百零九个卷积层，通道10384593717069655257060992658440192，3x3卷积核
        conv.append(nn.Conv2d(51922968
```

目录

CONTENTS

01

GAN

Generative Adversarial Nets, 生成式对抗网络

02

cGAN

Conditional GAN, 条件生成式对抗网络

03

DCGAN

Deep Convolutional GAN, 深度卷积生成式对抗网络

04

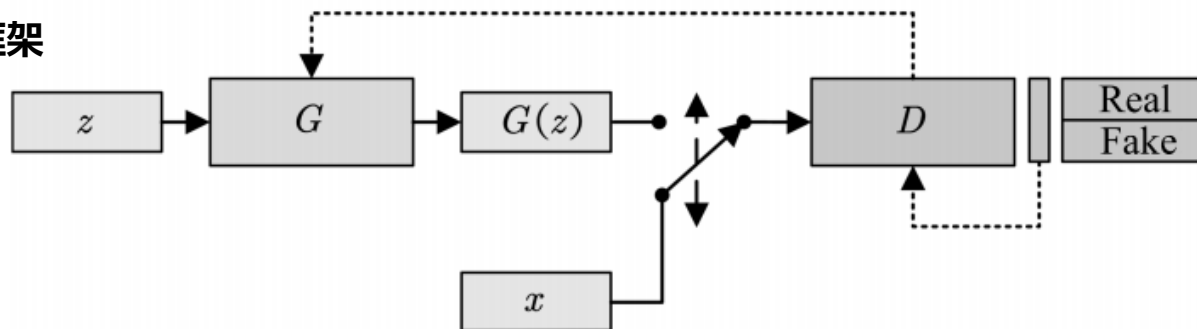
WGAN/WGAN-GP

Wasserstein GAN with Weight Clipping / Gradient Penalty

WGAN

回顾原始 GAN

GAN框架



生成式对抗网络的目标函数

$$\max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

$$\min_G V(G, D) = E_{z \sim p_z} [\log(1 - D(G(z)))]$$

等价于

$$\min_G \max_D V(G, D) = E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$$

博弈角度： 判别器D 区分真假样本；生成器G 合成真实样本；（**罪犯和商贩的例子**）

理论分析： 优化GAN等于最小化JS散度。

$$D^*(x) = \frac{p_{data}(x)}{p_{data}(x) + p_g(x)}$$

$$\min_G V(G, D^*) = 2JS(p_{data} || p_g) + 2\log \frac{1}{2}$$

WGAN

GAN、cGAN



GAN



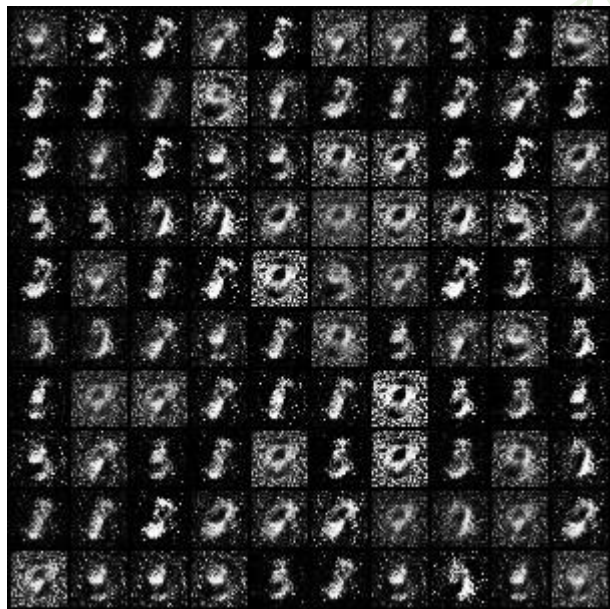
CGAN



Improved CGAN

WGAN

原始 GAN 存在的问题



训练困难：生成器无法生成想要的数



模式崩塌：生成器无法学习到完整的数据分布

WGAN

JS散度

生成式对抗网络: $\min_G V(G, D^*) = 2JS(p_{data}||p_g) + 2\log\frac{1}{2}$

JS散度: $JS(P_1||P_2) = \frac{1}{2}KL(P_1||\frac{P_1+P_2}{2}) + \frac{1}{2}KL(P_2||\frac{P_1+P_2}{2})$
 $JS(P_1||P_2) = \frac{1}{2} \int_x P_1(x) \log(\frac{P_1(x)}{\frac{1}{2}[P_1(x)+P_2(x)]}) dx + \frac{1}{2} \int_x P_2(x) \log(\frac{P_2(x)}{\frac{1}{2}[P_1(x)+P_2(x)]}) dx$

分析任意一个点x对JS散度的贡献:

当 $P_1(x) = 0$ and $P_2(x) = 0$ 时:

$JS(P_1||P_2)=0$, 对计算JS散度无贡献

当 $P_1(x) \neq 0$ and $P_2(x) = 0$ 或 $P_1(x) = 0$ and $P_2(x) \neq 0$ 时:

贡献等于常数 $\log 2$, 但是梯度等于零

当 $P_1(x) \neq 0$ and $P_2(x) \neq 0$ 时:

对计算JS散度有贡献且不为常数, 因此梯度不为零。

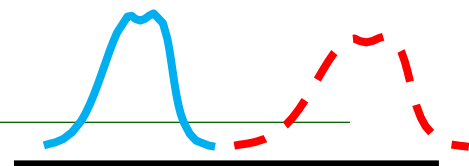
WGAN

不重叠或重叠部分可忽略概率

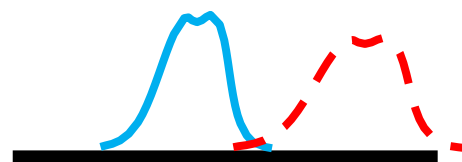
分析 $P_{\text{data}}(x)$ 与 $P_g(x)$ 不重叠或重叠部分可忽略的可能性有多大?

答: 非常大!

原因: $P_{\text{data}}(x)$ 与 $P_g(x)$ 的支撑集 (support) 是**高维空间中的低维流形** (manifold) 时, $P_{\text{data}}(x)$ 与 $P_g(x)$ **重叠部分测度 (measure) 为0的概率为1**。

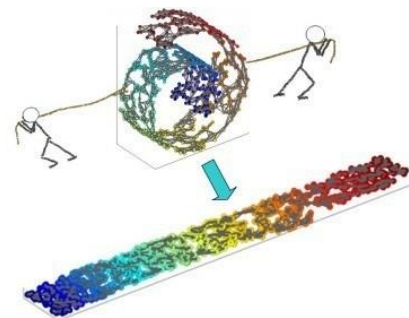
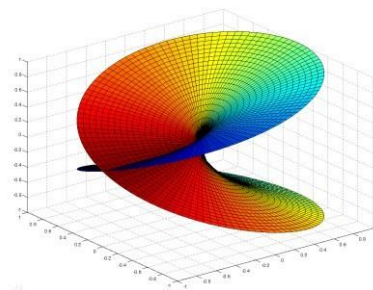
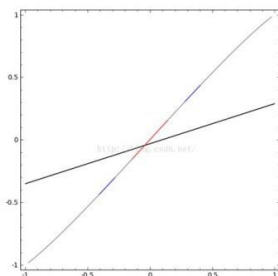
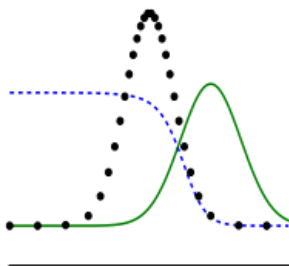
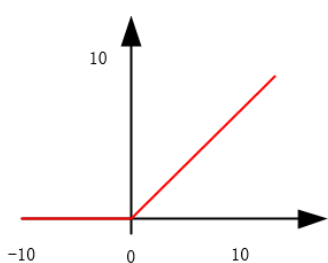


$$P_1(x) = 0 \text{ and } P_2(x) \neq 0$$
$$P_1(x) \neq 0 \text{ and } P_2(x) = 0$$



$$P_1(x) \neq 0 \text{ and } P_2(x) \neq 0$$

支撑集 (support) : 函数的非零部分子集。ReLU函数的支撑集就是零到正无穷, 概率分布的支撑集就是所有概率密度非零部分的集合



流形 (manifold) : 是高维空间中曲线、曲面概念的拓广。二维空间中的曲线是一个一维流形; 三维空间中的一个曲面是一个二维流形。因为它的本质维度 (intrinsic dimension) 只有2, 一个点在这个二维流形上移动只有两个方向的自由度。

测度 (measure) : 高维空间中长度、面积、体积概念的拓广。

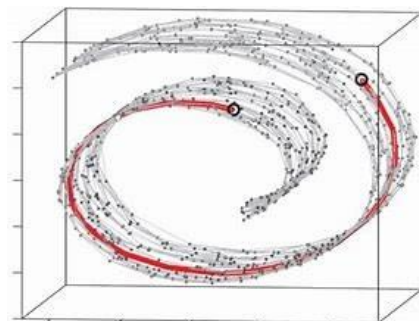
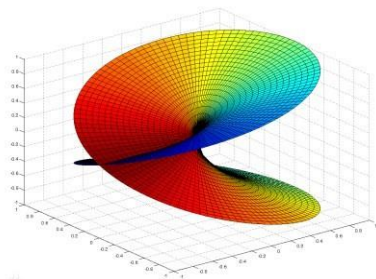
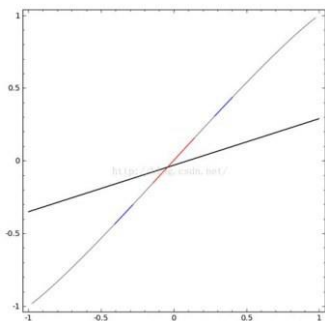
WGAN

不重叠或重叠部分可忽略概率

答2: $P_1(x)$ 与 $P_2(x)$ 的支撑集 (support) 是**高维空间中的低维流形** (manifold) 时, $P_1(x)$ 与 $P_2(x)$ **重叠部分测度 (measure) 为0的概率为1**。

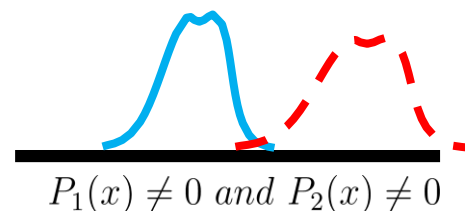
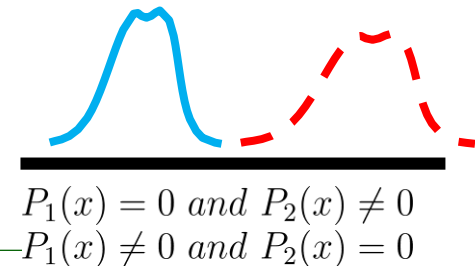
P_{data} 与 P_g 重叠部分测度为0的概率为1: **成立!**

“撑不满”就会导致真实分布与生成分布难以“碰到面”, 这很容易在二维空间中理解。一方面, 二维平面中随机取两条曲线, 它们之间刚好存在重叠线段的概率为0。另一方面, 虽然它们很大可能会存在交叉点, 但是相比于两条曲线而言, 交叉点比曲线低一个维度, 长度 (测度) 为0, 可以忽略。



当 P_{data} 与 P_g 的支撑集是高维空间中的低维流形时: **成立!**

如果生成图片大小是32, 那么它在一个 $32*32=1024$ 维的高维空间中。生成图片的输入一般来自低维 (比如100) 随机噪声, 即使经过非线性映射, 它的维度依然被限制在100维。所以生成图片是高维 (1024) 空间中的低维 (100) 流形!



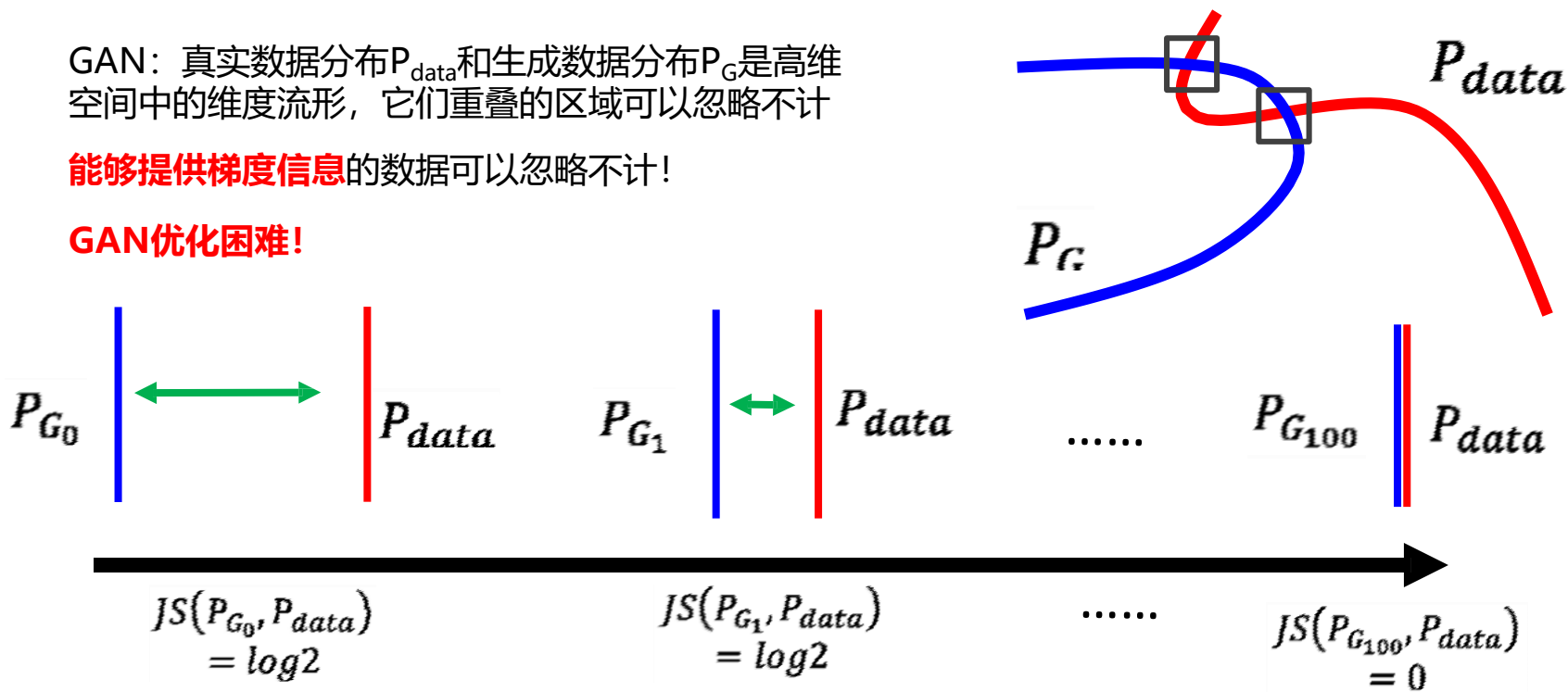
WGAN

总结JS散度

GAN: 真实数据分布 P_{data} 和生成数据分布 P_G 是高维空间中的维度流形, 它们重叠的区域可以忽略不计

能够提供梯度信息的数据可以忽略不计!

GAN优化困难!



结论:

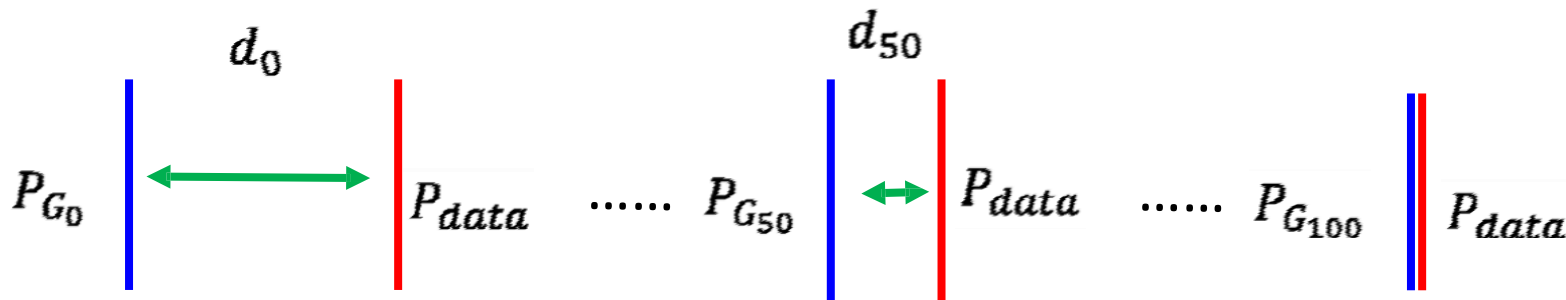
P_G 和 P_{data} 是高维空间的低维流形。判别器 (近似) 最优, 最小化生成器的目标函数等价于最小化真实分布 P_{data} 与生成分布 P_g 之间的JS散度, 而由于**真实分布 P_{data} 与生成分布 P_g 几乎不可能有不可忽略的重叠**, 所以无论它们相距多远, JS散度都是常数。最终**导致生成器的梯度 (近似) 为0**, 造成梯度消失, 造成**GAN难以优化**。

WGAN

Wasserstein 距离

KL散度、JS散度：当两个分布距离较远、或重叠部分可忽略不计的时候，两个分布之间的JS散度为常数，无法反应分布的远近，也就无法进行优化。

解决方案：找到一个距离度量，即使在两个分布距离较远、或重叠部分可忽略不计的时候，也能充分反应两个分布的远近。



✗

$$JS(P_{G_0}, P_{data}) = \log 2$$

$$JS(P_{G_{50}}, P_{data}) = \log 2$$

$$JS(P_{G_{100}}, P_{data}) = 0$$

✓

$$W(P_{G_0}, P_{data}) = d_0$$

$$W(P_{G_{50}}, P_{data}) = d_{50}$$

$$W(P_{G_{100}}, P_{data}) = 0$$

WGAN

Wasserstein 距离

Wasserstein 距离: 即使两个分布没有重叠或可忽略不计, Wasserstein距离仍然能够反映它们的远近。

$$W(P_1, P_2) = \inf_{\gamma \sim \Pi(P_1, P_2)} \mathbb{E}_{(x, y) \sim \gamma} [\|x - y\|]$$

$\Pi(P_1, P_2)$: P_1 和 P_2 组合起来的所有可能的联合分布的集合

$\|x - y\|$: 样本 x 和 y 的距离 L1 距离

$\inf$: 所有可能的下界

分布	C	D
A	0	1
B	1	0

联合分布1

$$\sqrt{2} \quad 2$$

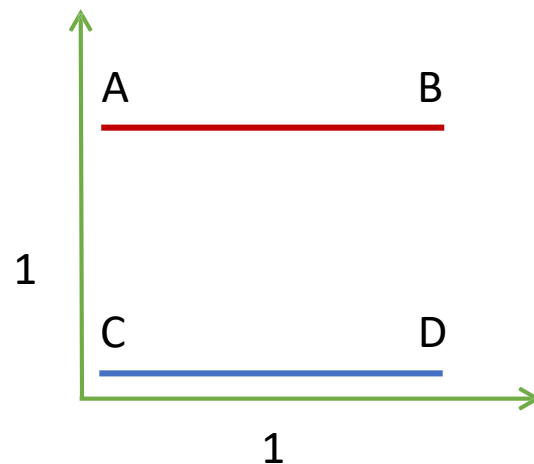
分布	C	D
A	1	0
B	0	1

联合分布2

2

距离	C	D
A	1	$\sqrt{2}$
B	$\sqrt{2}$	1

距离 $\|x - y\|$

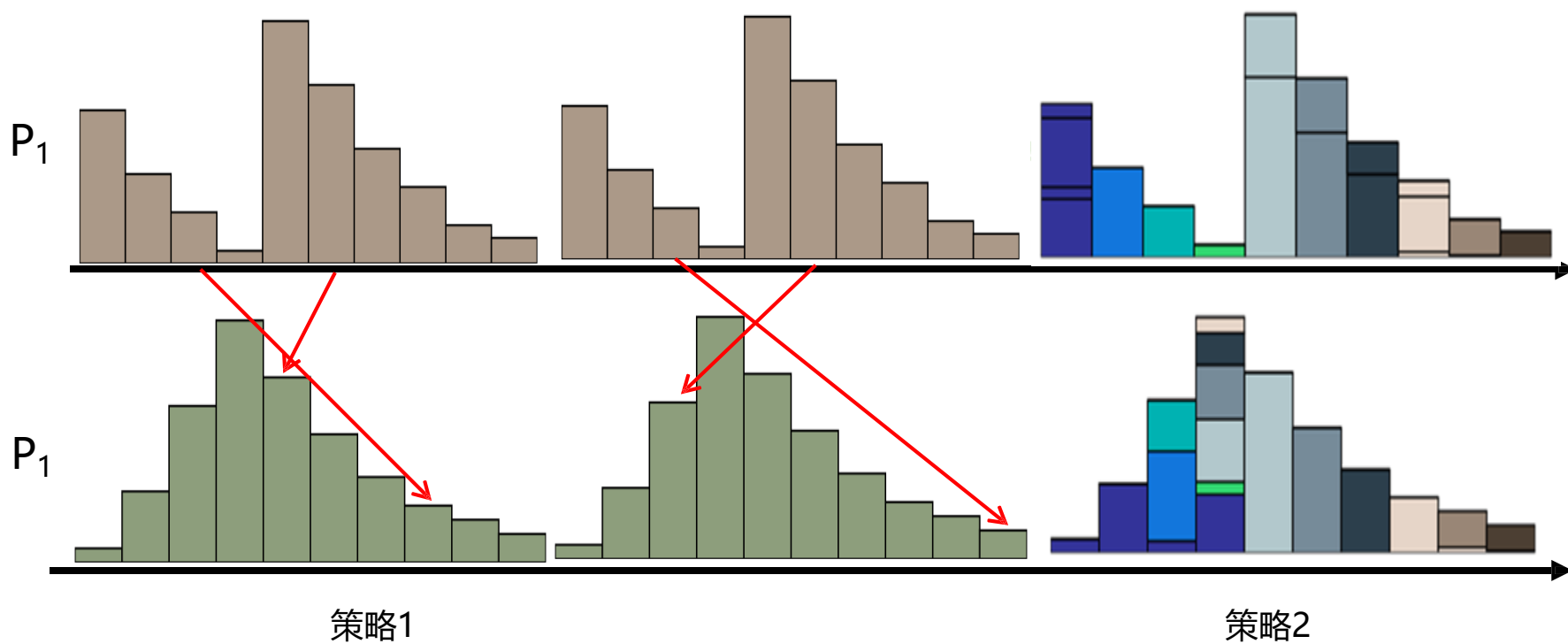


二维空间中的两个离散分布 AB、CD
如何求 AB 和 CD 的 Wasserstein 距离?

理解: 把 C 全部移动到 A, 把 B 全部移动到 D

WGAN

Wasserstein 距离



WGAN

Wasserstein 距离

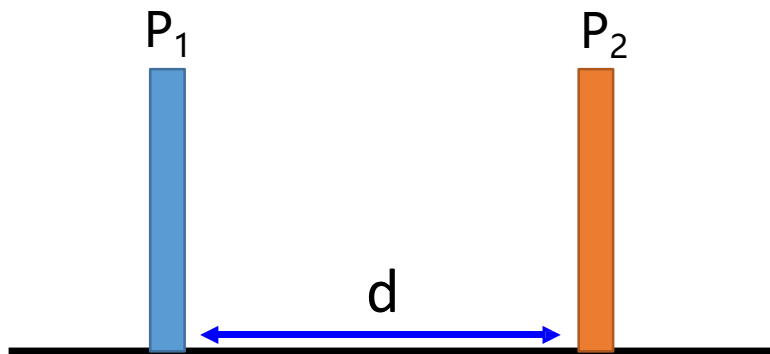
Wasserstein 距离:

$$W(P_1, P_2) = \inf_{\gamma \sim \Pi(P_1, P_2)} \mathbb{E}_{(x, y) \sim \gamma} [\|x - y\|]$$

理解：在**最优的路径规划**下，把土堆从 **P_1** 移动到 **P_2** 所需的路径

Wasserstein 距离又称 **Earth Mover's (推土机) 距离**

可以理解为：把 P_1 这堆“沙土”移动到 P_2 处所需最短距离消耗。



两个离散分布 P_1 、 P_2 的Wasserstein 距离

$$W(P_1, P_2) = d$$



WGAN

Wasserstein 距离

KL散度: 当两个分布重叠时为零, 不重叠是为正无穷

$$KL(P_1||P_2) = \mathbb{E}_{x \sim P_1}[\log \frac{P_1}{P_2}]$$

$$KL(P_1||P_2) = KL(P_1||P_2) = \begin{cases} +\infty & \text{if } \theta \neq 0 \\ 0 & \text{if } \theta = 0 \end{cases}$$

JS散度: 当两个分布重叠时为零, 不重叠时为常数 $\log 2$

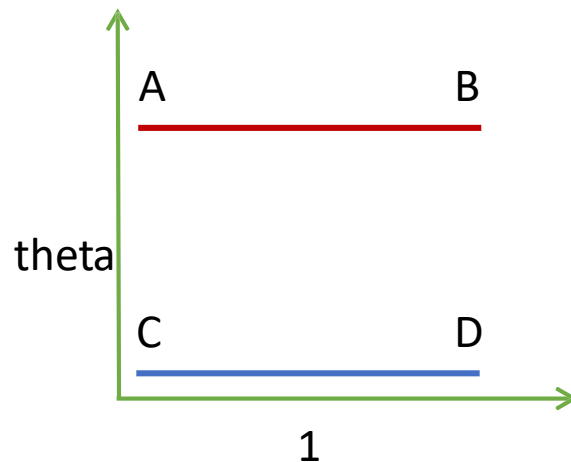
$$JS(P_1||P_2) = \frac{1}{2}KL(P_1||\frac{P_1+P_2}{2}) + \frac{1}{2}KL(P_2||\frac{P_1+P_2}{2})$$

$$JS(P_1||P_2) = \begin{cases} \log 2 & \text{if } \theta \neq 0 \\ 0 & \text{if } \theta = 0 \end{cases}$$

Wasserstein距离W: 即使两个分布没有重叠或可忽略不计, Wasserstein距离仍然能够反映它们的远近。

$$W(P_1, P_2) = \inf_{\gamma \sim \Pi(P_1, P_2)} \mathbb{E}_{(x, y) \sim \gamma} [|x - y|]$$

$$W(P_0, P_1) = |\theta|$$



分布	C	D
A	0.5	0
B	0	0.5

距离	C	D
A	θ	$\sqrt{1+\theta^2}$
B	$\sqrt{1+\theta^2}$	θ

优势: KL散度和JS散度是突变的, 要么最大要么最小, Wasserstein距离却是平滑的。如果我们用梯度下降法优化这个参数, 前两者根本提供不了梯度, Wasserstein距离却可以。类似地, 在高维空间中如果两个分布不重叠或者重叠部分可忽略, 则KL和JS既反映不了远近, 也提供不了梯度, 但是Wasserstein却可以提供有意义的梯度。

WGAN

Wasserstein 距离到损失

Wasserstein 距离:

$$W(P_1, P_2) = \inf_{\gamma \sim \Pi(P_1, P_2)} \mathbb{E}_{(x, y) \sim \gamma} [\|x - y\|]$$

inf (穷举所有可能, 取下界) 项**无法直接求解!**

Wasserstein 损失 (Kantorovich-Rubinstein二元性) :

$$W(P_1, P_2) = \frac{1}{K} \sup_{\|f\|_L \leq K} \mathbb{E}_{x \sim P_1} [f(x)] - \mathbb{E}_{x \sim P_2} [f(x)]$$

sup: 所有可能的上界

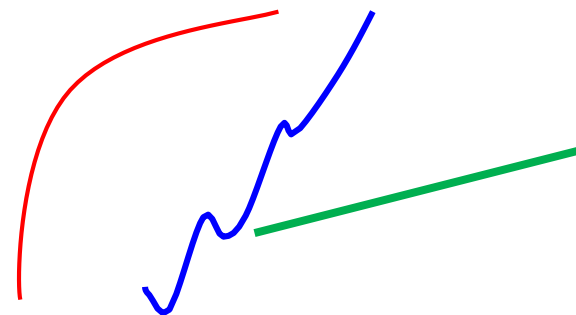
$\|f\|_L \leq K \quad |f(x_1) - f(x_2)| \leq K|x_1 - x_2|$ 函数f的利普希茨 (Lipschitz) 常数 $\|f\|_L$ 小于K

Lipschitz连续: 一个函数f(x), 对于x任意小的波动, f的变化小于一个常数k。此时k成为函数f的Lipschitz常数。Lipschitz连续条件限制了一个连续函数的**最大局部变动幅度**。

f任意函数->神经网络!

期望->采样求和求平均!

SUP->近似等于最大化神经网络f!



WGAN损失:

$$\min_G \max_{D, \|D\|_L \leq K} W(P_{data}, P_g) = \mathbb{E}_{x \sim P_{data}} [D(x)] - \mathbb{E}_{x \sim P_g} [D(x)]$$

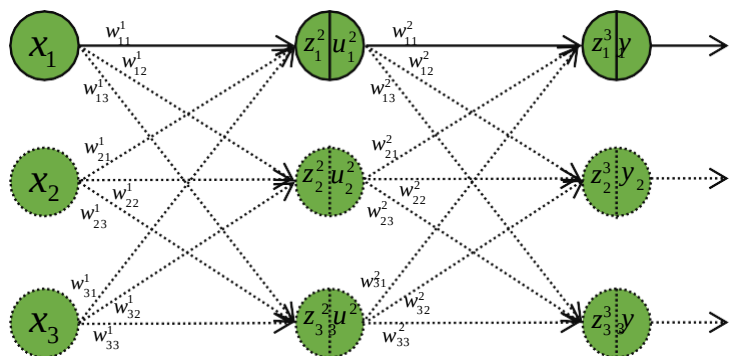
WGAN

利普希茨连续-权重截断

权重截断：在每次优化判别器D之后，把D的权重的绝对值限制在某个范围。

$$w = \begin{cases} w & \text{if } |w| \leq c \\ +c & \text{if } w > +c \\ -c & \text{if } w < -c \end{cases}$$

3层全连接神经网络，每层3个节点



公式表示

$$\begin{aligned} z_1^2 &= w_{11}^1 x_1 + w_{21}^1 x_2 + w_{31}^1 x_3 & u_1^2 &= \sigma(z_1^2) \\ z_2^2 &= w_{12}^1 x_1 + w_{22}^1 x_2 + w_{32}^1 x_3 & u_2^2 &= \sigma(z_2^2) \\ z_3^2 &= w_{13}^1 x_1 + w_{23}^1 x_2 + w_{33}^1 x_3 & u_3^2 &= \sigma(z_3^2) \\ z_1^3 &= w_{11}^2 u_1^2 + w_{21}^2 u_2^2 + w_{31}^2 u_3^2 & y_1 &= \sigma(z_1^3) \\ z_2^3 &= w_{12}^2 u_1^2 + w_{22}^2 u_2^2 + w_{32}^2 u_3^2 & y_2 &= \sigma(z_2^3) \\ z_3^3 &= w_{13}^2 u_1^2 + w_{23}^2 u_2^2 + w_{33}^2 u_3^2 & y_3 &= \sigma(z_3^3) \end{aligned}$$

$$\sigma(x) = \frac{1}{1+e^{-x}}$$

根据链式法则求导发现，导数正比于权重，所以截断权重w可以使网络满足利普希茨连续

$$\frac{dy_1}{dx_1} = \frac{dy_1}{dz_1^3} \cdot \frac{dz_1^3}{du_1^2} \cdot \frac{du_1^2}{dz_1^2} \cdot \frac{dz_1^2}{x_1} = y_1(1 - y_1) \cdot w_{11}^2 \cdot u_1^2(1 - u_1^2) \cdot w_{11}^1$$

$$\frac{dy_2}{dx_2} = \frac{dy_2}{dz_2^3} \cdot \frac{dz_2^3}{du_2^2} \cdot \frac{du_2^2}{dz_2^2} \cdot \frac{dz_2^2}{x_2} = y_2(1 - y_2) \cdot w_{11}^2 \cdot u_2^2(1 - u_2^2) \cdot w_{22}^1$$

...

WGAN

WGAN算法实现

GAN目标函数: $\min_G \max_D V(G, D) E_{x \sim p_{data}} [\log D(x)] + E_{z \sim p_z} [\log(1 - D(G(z)))]$

WGAN目标函数: $\min_G \max_{D, \|D\|_L \leq K} W(P_{data}, P_g) = \mathbb{E}_{x \sim P_{data}} [D(x)] - \mathbb{E}_{x \sim P_g} [D(x)]$

WGAN实现:

1. 判别器最后一层去掉 sigmoid 激活函数: 原始GAN中, 最优的D是一个概率形式, 属于 $[0, 1]$, 因此用sigmoid激活函数可以加速D的训练。WGAN中, D是任意满足利普斯次连续的函数, 因此不需要用Sigmoid激活。

$$\text{GAN } D^*(x) = \frac{p_{data}(x)}{p_{data}(x) + p_g(x)}$$

2. 生成器G 和判别器D 的损失不取对数: 根据Wasserstein距离的定义可知, 不需要取对数。

3. 每次更新判别器D 的参数之后, 权重绝对值截断到不超过一个固定常数c: 使得判别器D 满足利普希茨连续。

4. 不基于动量的优化算法 (包括momentum和Adam), 推荐RMSProp或者SGD: 经验。

WGAN

WGAN算法流程图

Algorithm 1 WGAN, our proposed algorithm. All experiments in the paper used the default values $\alpha = 0.00005$, $c = 0.01$, $m = 64$, $n_{\text{critic}} = 5$.

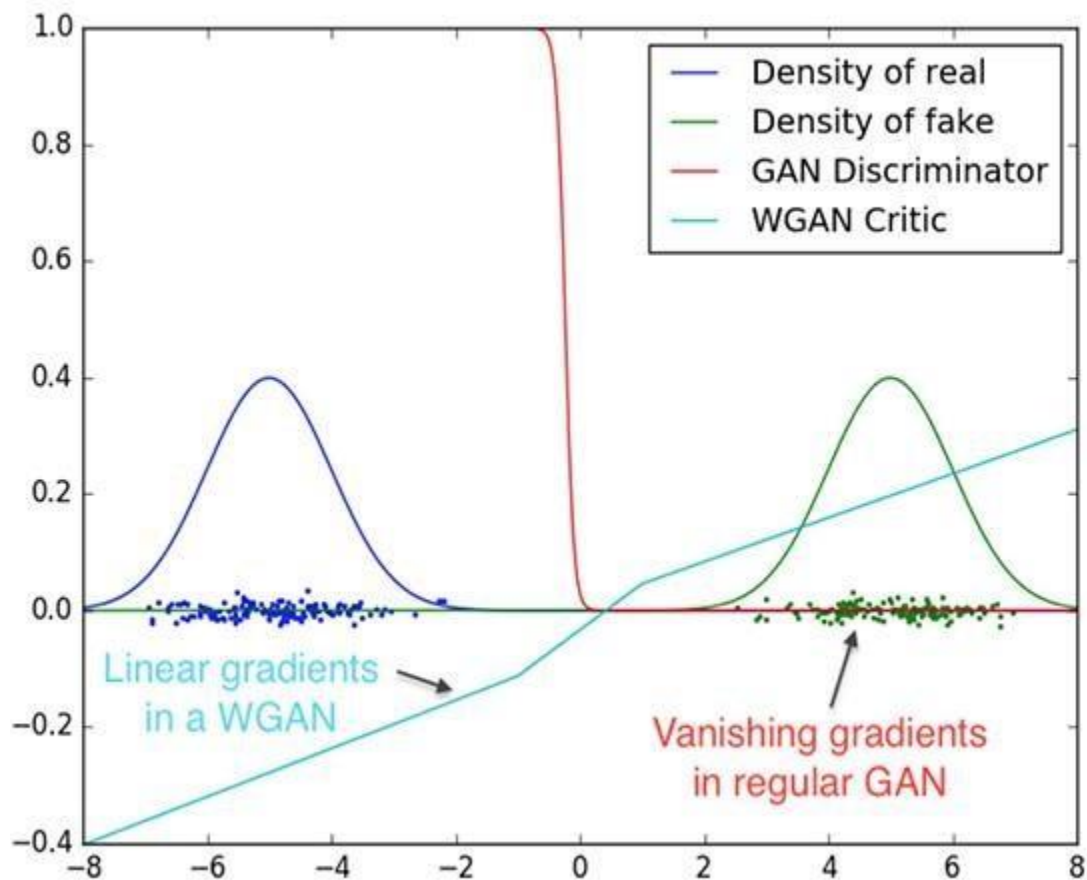
Require: : α , the learning rate. c , the clipping parameter. m , the batch size. n_{critic} , the number of iterations of the critic per generator iteration.

Require: : w_0 , initial critic parameters. θ_0 , initial generator's parameters.

```
1: while  $\theta$  has not converged do
2:   for  $t = 0, \dots, n_{\text{critic}}$  do
3:     Sample  $\{x^{(i)}\}_{i=1}^m \sim \mathbb{P}_r$  a batch from the real data.
4:     Sample  $\{z^{(i)}\}_{i=1}^m \sim p(z)$  a batch of prior samples.
5:      $g_w \leftarrow \nabla_w \left[ \frac{1}{m} \sum_{i=1}^m f_w(x^{(i)}) - \frac{1}{m} \sum_{i=1}^m f_w(g_\theta(z^{(i)})) \right]$ 
6:      $w \leftarrow w + \alpha \cdot \text{RMSPProp}(w, g_w)$ 
7:      $w \leftarrow \text{clip}(w, -c, c)$ 
8:   end for
9:   Sample  $\{z^{(i)}\}_{i=1}^m \sim p(z)$  a batch of prior samples.
10:   $g_\theta \leftarrow -\nabla_\theta \frac{1}{m} \sum_{i=1}^m f_w(g_\theta(z^{(i)}))$ 
11:   $\theta \leftarrow \theta - \alpha \cdot \text{RMSPProp}(\theta, g_\theta)$ 
12: end while
```

WGAN

原始GAN v.s. WGAN



原始GAN梯度消失严重，WGAN仍然可以提供梯度

WGAN

实验结果



上图: WGAN

下图: DCGAN

结论: WGAN的视觉效果比DCGAN略胜一筹

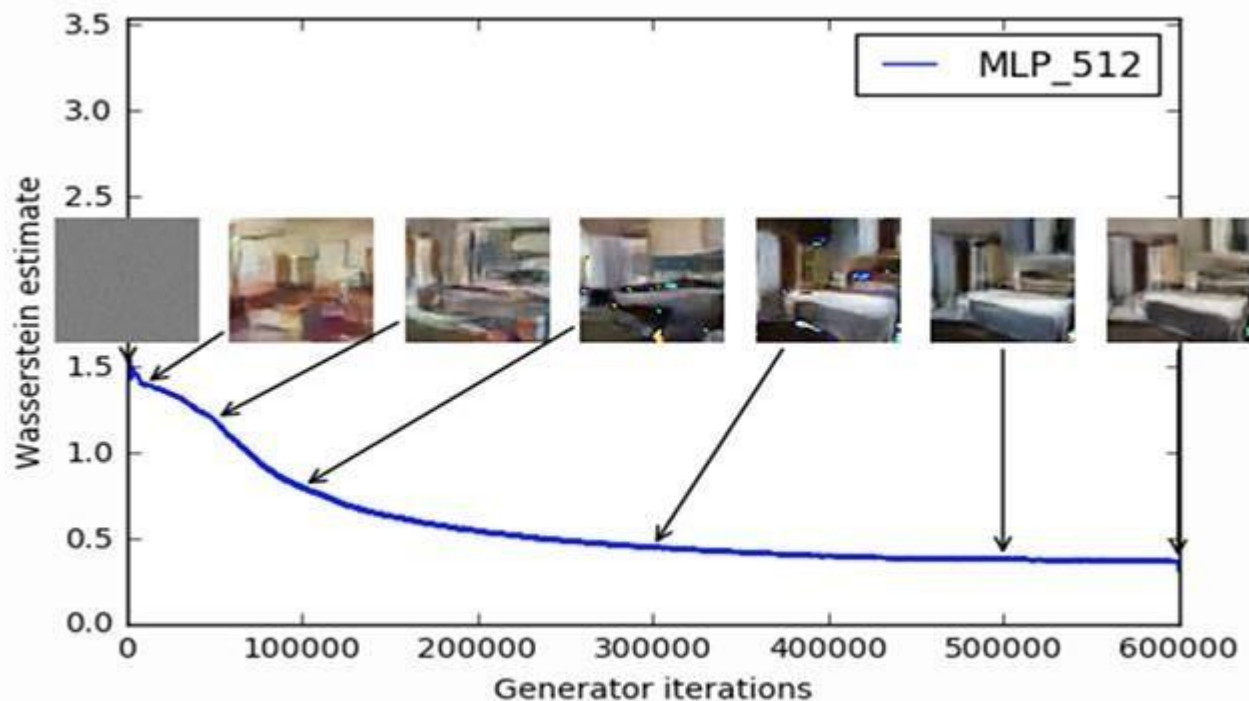
上图: 不用BN的 WGAN

下图: 不用BN的 DCGAN

结论: DCGAN完全崩溃, WGAN仍可以稳定生成

WGAN

W距离指示优化进度



原始GAN的损失函数无法指示：无论两个分布远近，JS散度都是常数。
WGAN损失可以指示训练进度：W距离可以表示两个分布的远近。

目录

CONTENTS

01

GAN

Generative Adversarial Nets, 生成式对抗网络

02

cGAN

Conditional GAN, 条件生成式对抗网络

03

DCGAN

Deep Convolutional GAN, 深度卷积生成式对抗网络

04

WGAN/**WGAN-GP**

Wasserstein GAN with Weight Clipping/ Gradient Penalty

WGAN-GP

回顾WGAN

WGAN目标函数: $\min_G \max_{D, \|D\|_L \leq K} W(P_{data}, P_g) = \mathbb{E}_{x \sim P_{data}}[D(x)] - \mathbb{E}_{x \sim P_g}[D(x)]$

Lipschitz连续: 一个函数 $f(x)$, 对于 x 任意小的波动, f 的变化小于一个常数 k 。此时 k 成为函数 f 的Lipschitz常数。

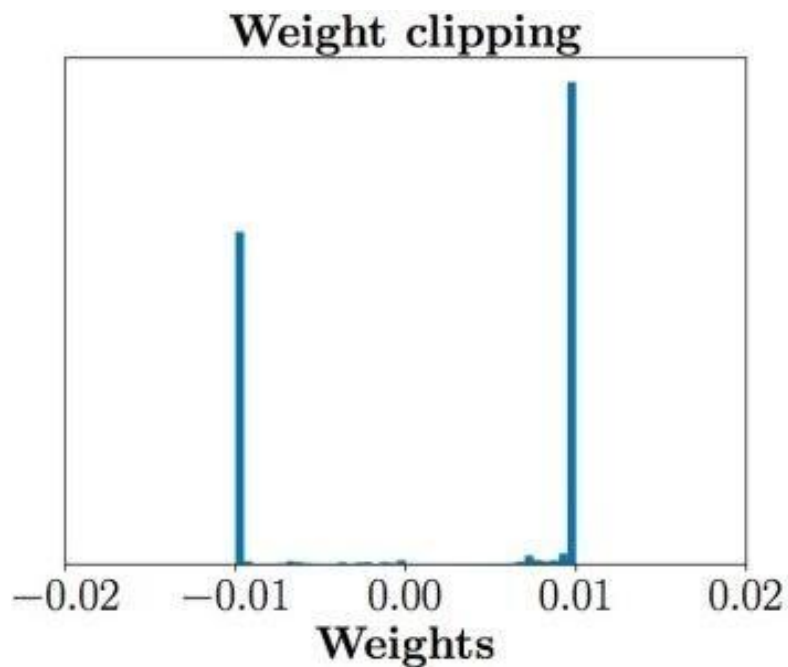
$$|f(x_1) - f(x_2)| \leq K|x_1 - x_2|$$

权重截断: 为了使得判别器 D 满足Lipschitz连续, 强制 D 所有权重的绝对值小于一个常数 c 。

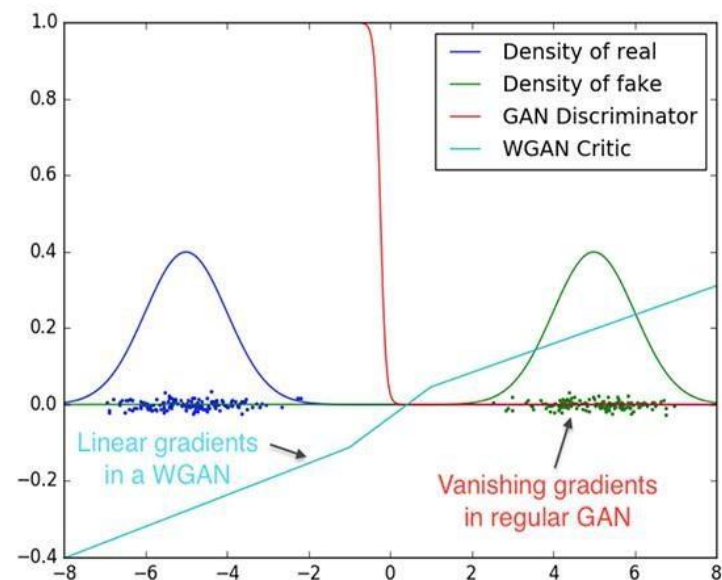
$$w \begin{cases} w & \text{if } |w| \leq c \\ +c & \text{if } w > +c \\ -c & \text{if } w < -c \end{cases}$$

WGAN-GP

缺陷1：权重二值化



WGAN几乎所有参数都是正负0.01：
浪费神经网络的拟合能力。



判别器越好，分界面越抖，梯度绝对值越大，权重绝对值越大，就有越多的权重被截断到边界。

WGAN-GP

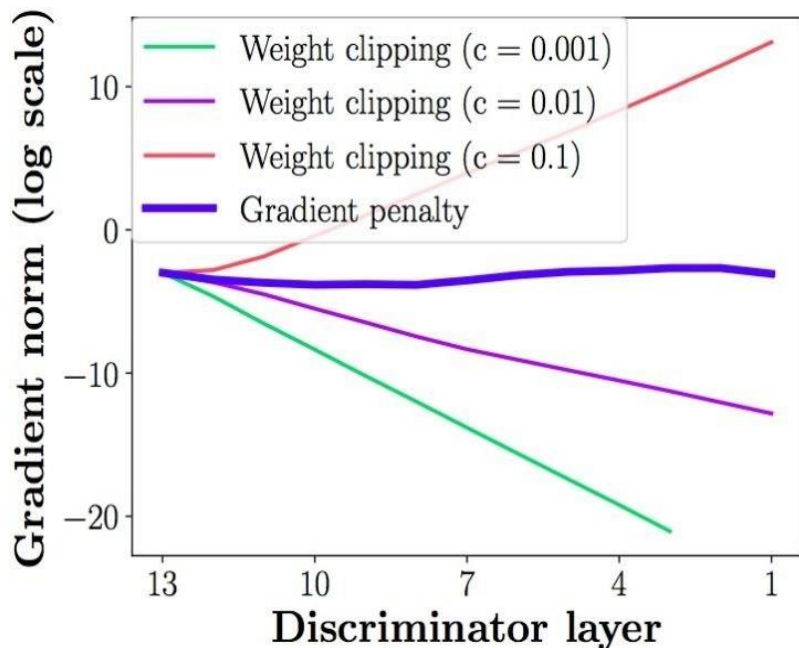
缺陷2：梯度衰减、爆炸

参数截断阈值

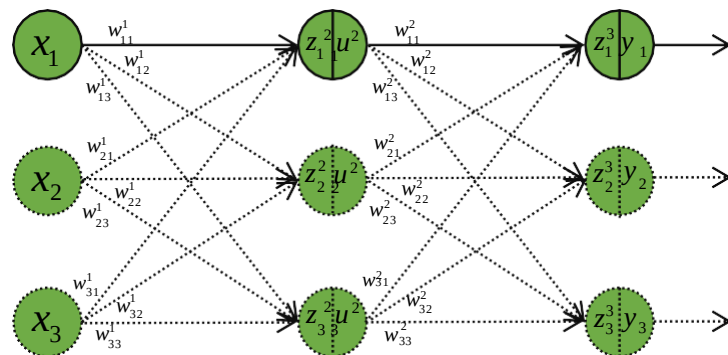
小：根据求导的链式法则，经过多层反向传播，梯度会变得非常小，造成梯度消失

大：根据求导的链式法则，经过多层反向传播，梯度会变得非常大，造成梯度爆炸

适中：实际中适中区域很小，给调参增加了难度



3层全连接神经网络，每层3个节点



$$\frac{dy_1}{dw_{11}^1} = y_1(1 - y_1) \cdot w_{11}^2 \cdot u_1^2(1 - u_1^2) \cdot x_1$$

$$\frac{dy_2}{dw_{11}^2} = y_2(1 - y_2) \cdot u_1^2$$

WGAN-GP

梯度惩罚项

WGAN: 梯度截断太“武断”，对于所有参数采取一刀切的方式。尽管满足了Lipschitz连续，但是却“滥杀无辜”。如何让模型既满足Lipschitz限制，又更加“人性化”地学习参数。

WGAN-GP: 把Lipschitz限制作为一个正则项加到Wasserstein损失上，在优化GAN损失的同时，尽可能满足Lipschitz限制。

WGAN 目标函数

$$\min_G \max_{D, \|D\|_L \leq K} W(P_{data}, P_g) = \mathbb{E}_{x \sim P_{data}}[D(x)] - \mathbb{E}_{x \sim P_g}[D(x)]$$

WGAN使用权重截断来近似满足利普希茨连续条件：

$$w = \begin{cases} w & \text{if } |w| \leq c \\ +c & \text{if } w > +c \\ -c & \text{if } w < -c \end{cases}$$

WGAN-GP使用一个正则项来近似满足利普希茨连续条件：

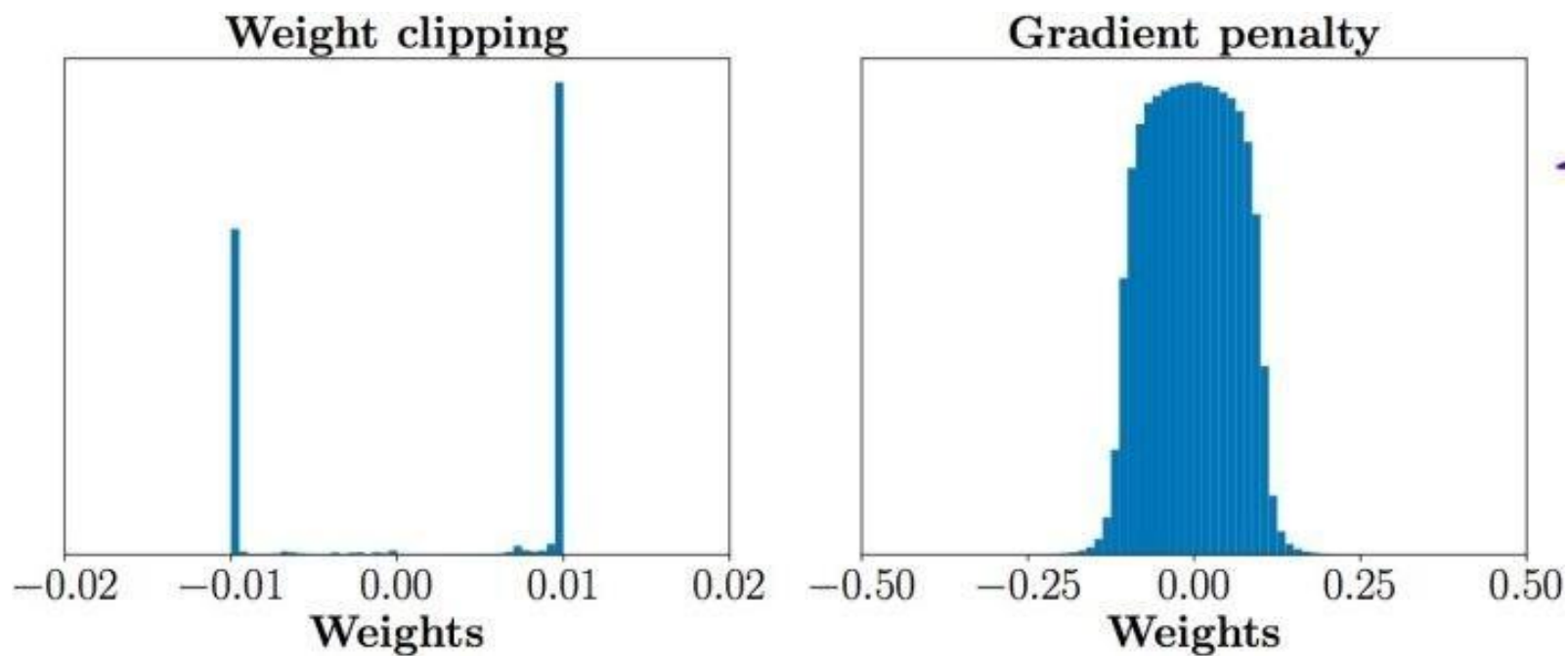
$$\min_G \max_D W(P_{data}, P_g) = \mathbb{E}_{x \sim P_{data}}[D(x)] + \mathbb{E}_{x \sim P_g}[D(x)] + \lambda \mathbb{E}_{x \sim \mathcal{X}} [\|\nabla_x D(x)\|_p - K]^2$$

$$x_{data} \sim P_{data}, x_g \sim P_g, \epsilon \sim Uniform[0, 1]$$

$$\hat{x} = \epsilon x_{data} + (1 - \epsilon) x_g$$

WGAN-GP

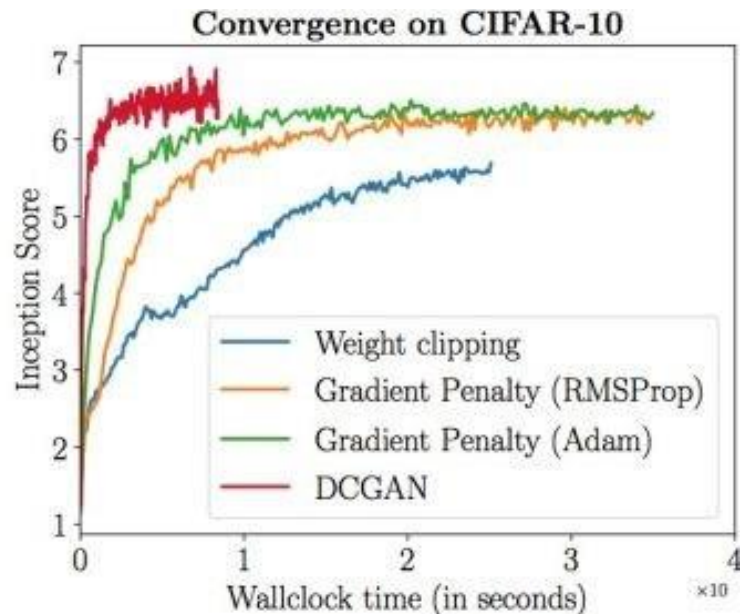
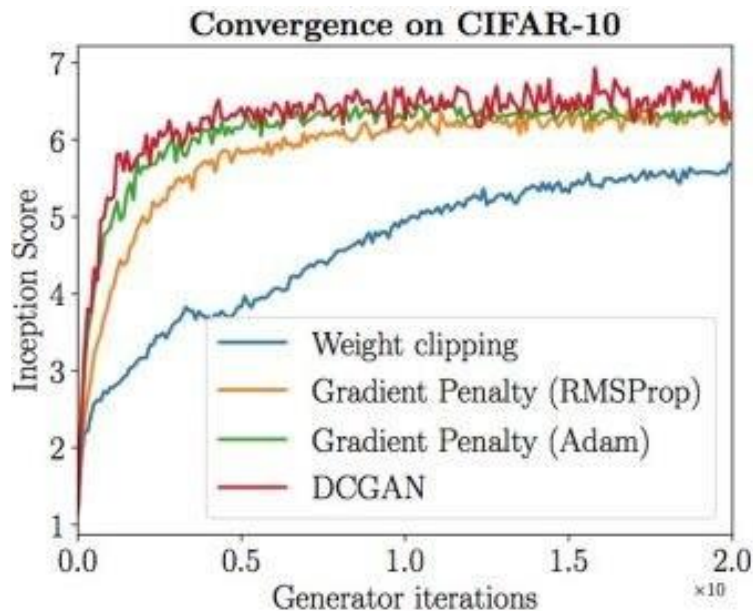
梯度更加均匀



参数分布更加均匀!

WGAN-GP

收敛速度更快

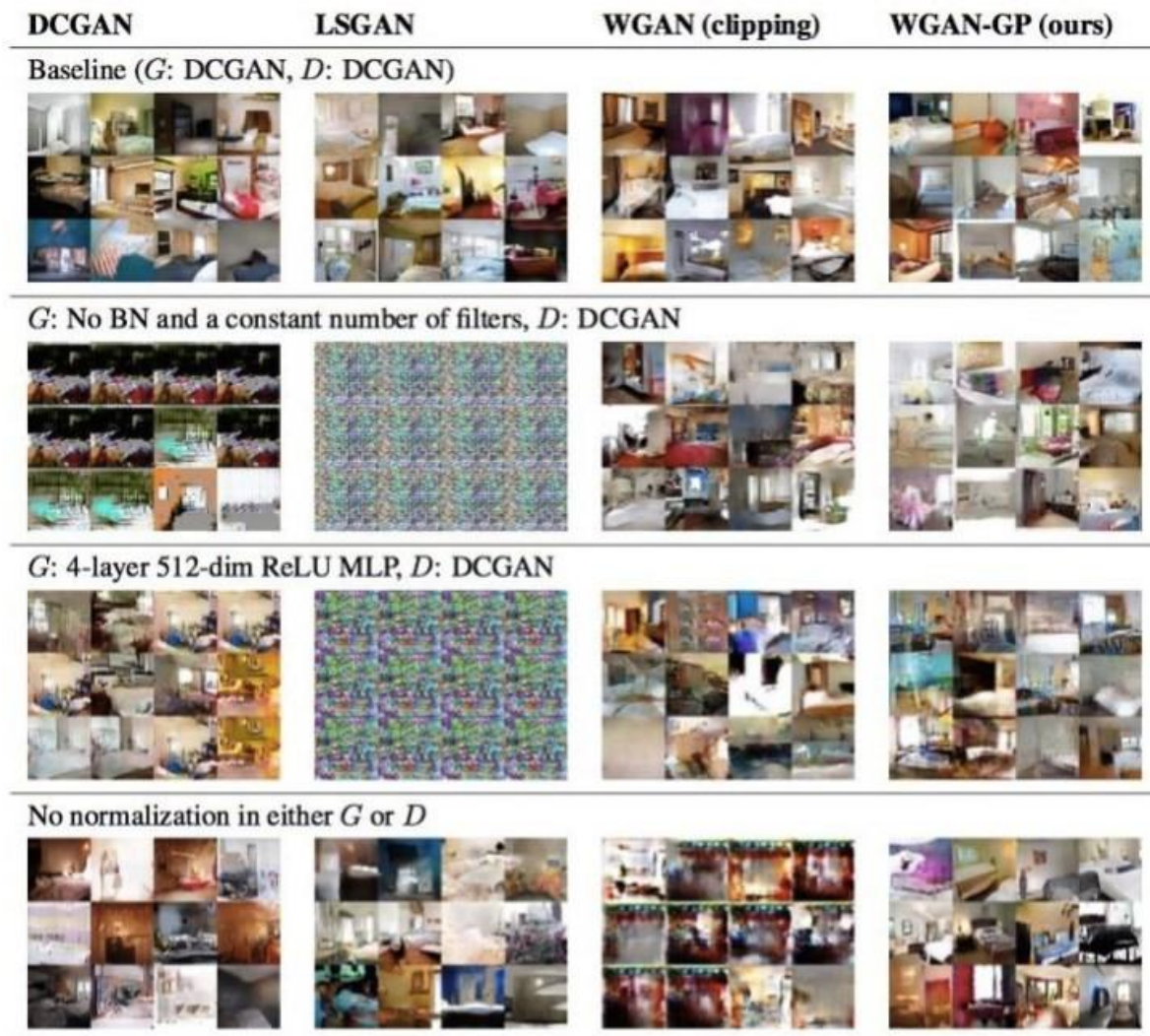


收敛速度显著高于WGAN!
但是仍然慢于DCGAN。

WGAN-GP

对模型更鲁棒

在各种配置下，WGAN-GP均能稳定发挥！
稳定性远远高于DCGAN！



WGAN-GP

对模型更鲁棒

在各种配置下，WGAN-GP均能稳定发挥！
稳定性远远高于DCGAN！

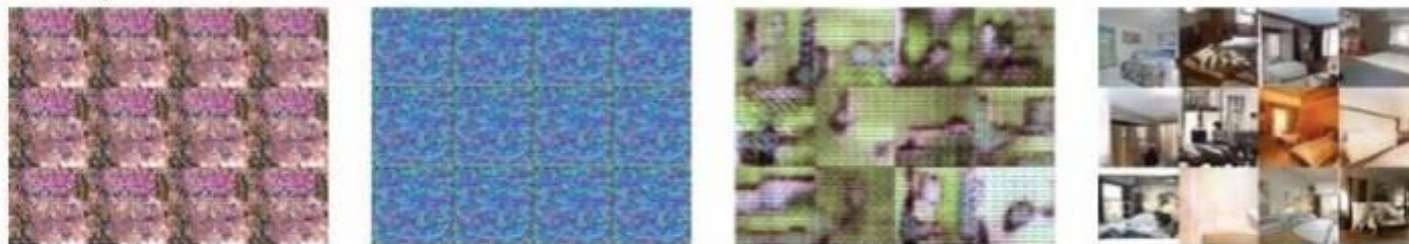
Gated multiplicative nonlinearities everywhere in G and D



$\tanh$ nonlinearities everywhere in G and D



101-layer ResNet G and D





THANKS