

CV Week 1

1. OpenCV 用 BGR order purity brightness
2. RGB HSV (human perception) Grayscale
 color type $\rightarrow 100$ 鲜艳
 Hue Saturation Value
 (human perception)
3. Top left 读取 矩阵, 先行后列

绿色在转化时权重最大

Week 2 image filters

1. box filter. ① mean filter $\frac{1}{9} \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix}$ 全 1 $h[m,n] = \sum_{k,l} g(k,l) f[m+k,n+l]$
 可以 smooth, remove sharp features

- ② sharpening $\begin{bmatrix} 0 & 1 & 0 \\ 1 & 5 & 1 \\ 0 & 1 & 0 \end{bmatrix} - \frac{1}{9} \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix}$ 锐化中心

linear filters 性质 (一个 filter 是线性, iff $T[\alpha f_1 + \beta f_2] = \alpha T[f_1] + \beta T[f_2]$)
 α, β 是标量常数. f_1, f_2 是任意图像

- ① 线性: $\text{filter}(f_1 + f_2) = \text{filter}(f_1) + \text{filter}(f_2)$
- ② 平移无关性: shift invariance $\text{filter}(\text{shift}(f)) = \text{shift}(\text{filter}(f))$ 输入移动像素 输出也移动像素

任何线性, 平移无关的计算可被卷积表示

- ③ 齐次性: 缩放输入等比缩放输出
- ④ 交换律: $f * h = h * f$ (仅限卷积定义, 互相关不对称)
- ⑤ 结合律 $(f * h_1) * h_2 = f * (h_1 * h_2)$
- ⑥ 分配律: $f * (h_1 + h_2) = f * h_1 + f * h_2$
- ⑦ Identity $a * e = a$

2. Gaussian filter $G_6 = \frac{1}{2\pi 6^2} e^{-\frac{(x^2+y^2)}{2 \cdot 6^2}}$

$$= \left(\frac{1}{\sqrt{2\pi} 6} e^{-\frac{x^2}{2 \cdot 6^2}} \right)$$

$$\left(\frac{1}{\sqrt{2\pi} 6} e^{-\frac{y^2}{2 \cdot 6^2}} \right)$$

- ① 移除高频, smooth

- ② 和自己卷还是高斯. $6 \rightarrow \sqrt{2} 6$

- ③ separable kernel 可以被分解为两个 1-d 高斯变量
 好. 假设输入为 $M \times N$. 一个像素要做 mn 次乘法.
 $MNmn \rightarrow MN(m+n)$

filter 多大. half-width $\geq 3\sigma$ 左右, 右是左

边界咋办? ① clip filter(black) 全填黑 ② wrap around ③ copy edge 把边 pixel extend ④ 镜像

3. Sobel filter $\Rightarrow$ 识别横, 纵边

4. laplacian filter unit-Gaussian $\approx$ Laplacian of Gaussian 可以提取高频细节

5. 为什么 filtering 重要

- ① 空间相关 ② locality 重要信息在小 region ③ Stationarity ^{平稳性} ④ Hierarchical structure ^{层次结构} ⑤ similarity ⑥ scale invariance

6. Denoising

- ① Gaussian: 大方差可 smooth + 压噪, 但伤害画质
 ② Median filtering 抑制 salt-and-pepper (impulse noise)
 select the median intensity: Robustness to outliers
 还有 weighted median, clipped mean 等

7. Template matching

a. 直接把 eye patch 当作 filter cons: 易被高亮响应, 而非对应 patch

b. 把 eye patch 去均值

c. Sum of Square Distance (SSD)
$$h[m, n] = \sum_{k, l} (f[k, l] - im[m+k, n+l])^2$$

 可进一步拆成更快的 linear filter

$$= \sum_{k, l} f[k, l]^2 - 2 \sum_{k, l} im[m+k, n+l] f[k, l] + \sum_{k, l} im[m+k, n+l]^2$$

 cons: 对 intensity shift 敏感

d. normalized cross correlation

(cosine similarity)
$$h[m, n] = \frac{\sum_{k, l} (f[k, l] - \bar{f})(im[m+k, n+l] - \bar{im}_{m, n})}{\sqrt{\sum_{k, l} (f[k, l] - \bar{f})^2 \sum_{k, l} (im[m+k, n+l] - \bar{im}_{m, n})^2}}$$

$$h[m, n] = \frac{\sum_{k, l} (f[k, l] - \bar{f})(im[m+k, n+l] - \bar{im}_{m, n})}{\sqrt{\sum_{k, l} (f[k, l] - \bar{f})^2 \sum_{k, l} (im[m+k, n+l] - \bar{im}_{m, n})^2}}$$

Comparison: zero-mean filter 最快, 但效果一般

SSD: 第二快, 但对 overall intensity 敏感

Normalized cross-correlation: 最慢, 和局部均值与对比无关

8. Low-pass filter

图 $\xrightarrow{\text{Gaussian filter}}$ $\xrightarrow{\text{Sample}}$ low-res

构建多分辨率层级
 单尺度全局大搜索 $\rightarrow$ 局部多尺度搜索

Gaussian Pyramid.

Laplacian Pyramid.

为什么先 gaussian? 直接 downsample 引发 Aliasing. 高频细节 $\rightarrow$ 低频伪影
 但 Gaussian Pyramid 的 downsample 是有损压缩

引入 Laplacian Pyramid 存预测偏差 $L_i = G_i - \text{Expand}(G_{i+1})$

可逆性

$\uparrow$
 上采样 + 高斯平滑

$$I_{\text{hybrid}} = \text{Low Freq}(A) + \text{High Freq}(B)$$

对 A 做 Gaussian filter

对 B 做 Laplacian Gaussian filter
 或 B-Gaussian(B)

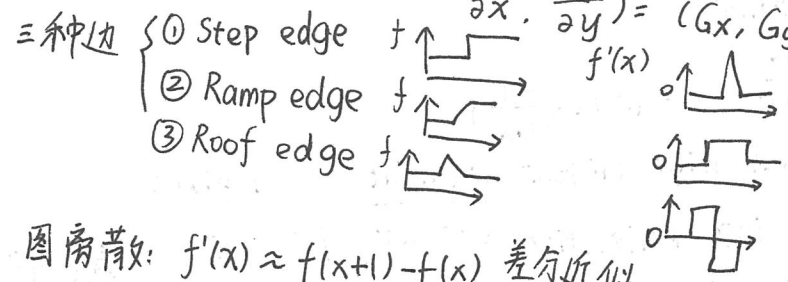
CV Week 3 Edges & Contours 边缘检测与轮廓分析

1. 尝试 - : High frequency $\approx$ Edges?

- ① Disconnected : 高频响应在弱边界处断裂. 无法生成闭合环
- ② Noise Amplification : 随机像素波动被高频算子放大. 产生大量伪边
- ③ Thick Edges : 边缘带宽 5~10 pixels. ~~丢失~~ not precise 但对 noise 更 sensitive
 edge 是 brightness 突然变化
 - 阶导在边缘处达到峰值, = 阶导在 edge 处穿过 0

2. Gradient-based : $\nabla I = (\frac{\partial I}{\partial x}, \frac{\partial I}{\partial y}) = (G_x, G_y)$

Gradient Magnitude $|\nabla I| = \sqrt{G_x^2 + G_y^2}$



Gradient Direction $\theta = \arctan2(G_y, G_x)$

图像离散: $f'(x) \approx f(x+1) - f(x)$ 差分近似
 只用 2 pixel, 噪音大 没有 smoothing

gradient $\perp$ edge 通用

2.1 Sobel Operator

Direction info: X 和 Y

2.1.1 Sobel X Kernel - 检测 纵边

Sobel X = $\begin{bmatrix} -1 & 0 & 1 \\ -2 & 0 & 2 \\ 1 & 0 & 1 \end{bmatrix}$ $f(x+1) - f(x-1)$
 水平变化 $\Leftrightarrow$ 纵边

2.1.2 Sobel Y Kernel - 检测 横边

Sobel Y = $\begin{bmatrix} 1 & -2 & 1 \\ 0 & 0 & 0 \\ -1 & 2 & -1 \end{bmatrix}$ $f(y+1) - f(y-1)$

这样 3x3 邻域 + Built-in smoothing (center weight = 2) + 对 噪声 不敏感

2.1.3 Design Insight

Sobel X = $\begin{bmatrix} -1 & 0 & 1 \\ -2 & 0 & 2 \\ 1 & 0 & 1 \end{bmatrix} = \begin{bmatrix} 1 \\ 2 \\ 1 \end{bmatrix} \otimes \begin{bmatrix} -1 & 0 & 1 \end{bmatrix}$

note 用 cv2.CV_64F

CV-8U 不能负 0~255

Smooth vertically

$\nwarrow$ differentiate horizontally

2.1.4 $|\nabla I|$ 大的 $\sqrt{G_x^2 + G_y^2}$ all edge directions combined

2.1.5 Limitations : ① Thick edges . 非 pixel-level precision

② Still Noise sensitive : 比 simple diff 好

③ No thresholding : 不能说哪个是边. 没有自动评判标准
 Do G 比 LoG 快
 $G(x,y,k) - G(x,y,k-1) \approx (k-1) \cdot \nabla^2 G$
 比 smooth 再拉 $\downarrow$ 高斯差分
 $LoG \approx DoG$ approximation

我们需要用 entire pipeline 把 gradient $\Rightarrow$ clean, precise 边

2.2 其他算子

$\begin{bmatrix} 0 & 0 & 0 \\ 1 & -2 & 1 \\ 0 & 0 & 0 \end{bmatrix} + \begin{bmatrix} 0 & 1 & 0 \\ 0 & -2 & 0 \\ 0 & 1 & 0 \end{bmatrix} = \begin{bmatrix} 0 & 1 & 0 \\ 1 & -4 & 1 \\ 0 & 1 & 0 \end{bmatrix}$

$[1, -2, 1]$ 与 $\begin{bmatrix} 1 \\ 2 \\ 1 \end{bmatrix}$ 相加

2.2.1 Prewitt Operator

有方向

Prewitt X = $\begin{bmatrix} -1 & 0 & 1 \\ -1 & 0 & 1 \\ -1 & 0 & 1 \end{bmatrix}$

Prewitt Y = $\begin{bmatrix} -1 & -1 & -1 \\ 0 & 0 & 0 \\ 1 & 1 & 1 \end{bmatrix}$

no center emphasis

2.2.2 Laplacian: 2nd : 找 $\nabla^2 I = 0$

$\begin{bmatrix} 0 & -1 & 0 \\ 1 & 4 & -1 \\ 0 & -1 & 0 \end{bmatrix}$ or $\begin{bmatrix} -1 & -1 & -1 \\ 1 & 8 & 1 \\ -1 & -1 & -1 \end{bmatrix}$

比 Sobel 对 noise 还敏感 !!!

LoG: Gaussian + Laplacian

Smooth first $\rightarrow$ clean edges

用于 detects all edges equally (isotropic)
 LoG (with 高斯核)

3. Canny Edge Detection 40年历久弥新. 先定义 optimal. 再设计

3.1 3 Goals

- ① low error rates: 真边✓ 假边✗
- ② precise localization: 精确确定位置
- ③ Single Response: 一条边 → 一个 pixel

Gaussian Smoothing + dual threshold
NMS
NMS Non-Maximum Suppression

3.2 5 步

$$\frac{d}{dx}(f+n) = \frac{df}{dx} + \frac{dn}{dx}$$

白噪声均值为0, 但方差被放大

① Gaussian Smoothing: 去噪. 保结构

② Gradient: Sobel $M = \sqrt{G_x^2 + G_y^2}$ $\theta = \arctan(\frac{G_y}{G_x})$

③ NMS: keep only ridge (pixel with largest gradient) 1 pixel
沿梯度方向比较 (和边垂直) 量化到最近角度
将连续 gradient 方向 quantized into 4 directions: $0^\circ, 45^\circ, 90^\circ, 135^\circ$ 额外计算
对每个 pixel, 沿该方向比较前后两个相邻 pixel 的 gradient magnitude 满足一个方向即可
若 $\text{current} \geq \text{both} \rightarrow \text{keep}$; else: suppress 舍弃
→ one pixel

④ Dual Threshold 单个 threshold $\begin{cases} T_{\text{大}}: \text{noise 抑制但不连续} \\ T_{\text{小}}: \text{有 noise 但连续} \end{cases}$ → 双 threshold

$M \geq T_{\text{high}}$: 强边 (一定是真边 "种子")

$T_{\text{low}} \leq M \leq T_{\text{high}}$: 弱边 (noise / 边的延伸)

$M < T_{\text{low}}$: 舍弃

强度不是唯一, connectivity decides

$T_{\text{high}}: T_{\text{low}}$ 2:1 to 3:1

① 够宽才行. $(T_{\text{L}}, T_{\text{H}})$ too narrow, 真 contour weak segments mostly fall below T_{L} and being discarded outright. 后面无法重连
→ 强弱边

⑤ Edge Tracking by Hysteresis

① 从强边出发, 沿 8 邻域搜弱边 可用队列 queue 维护

② 连到的弱边 promote 成强边 能在此上搜 8 邻域

③ 孤立弱边为 noise

② 太宽不行. $\Rightarrow T_{\text{L}}$ 很低, all noise pixels 进入 dense noise regions 在 step 5 仍可能成边

cv2. Canny (img, threshold1=50, threshold2=150) built-in 用固定 5x5 G Blur

3.3 去 ① 很多 noise → false edges ② 去 ③ thick bands ③ 单 threshold 去边/noise
去 ⑤ weaker edges 可能是真边被去掉了

3.4 选什么. ① general: Canny

② 需要 gradient 方向: Sobel

③ Isotropic: LoG

④ Real time: Sobel (single-step)

调参 平滑尺度
 $\sigma, T_{\text{L}}, T_{\text{H}}$

4. Edge → Contour

4.1 区别: 边是 pixel labels, 2D image, 不连续, 不能数 object, 不能看形状
 轮廓是 connected curves, 点序列, 连续, 可以数, 能看形状

4.2 pipeline Original $\xrightarrow{\text{Canny}}$ Edge Image $\xrightarrow{\text{find Contour}}$

4.3 find Contours Contour List → Object Boundary

- ① 扫图找第一个 edge 起点
- ② 沿 edge 逐像素追踪, 直到返回原点, 作为一个 contours
- ③ 继续扫, 提取所有

可以只提最外层 | 所有 (list) | 所有 (Tree)
 RETR_EXTERNAL RETR_LIST RETR_TREE

4.4 性质: S, C, 矩, Centroid, bounding box
 压缩: approx - none 不压, approx - simple 只保端点
 空间矩 $M_{ij} = \sum_{x,y} x^i y^j I(x,y)$ $S = M_{00}$; 质心 $\bar{x}, \bar{y} = (\frac{M_{10}}{M_{00}}, \frac{M_{01}}{M_{00}})$

$M_{pq} = \sum_{(x,y) \in \text{contour}} x^p y^q$ 次方

M_{00} zeroth moment = S 像素总数
 M_{10} - 阶矩 x 方向, 所有像素 x 坐标之和
 M_{01} y 方向, 所有像素 y 坐标之和

5. 分析工具

5.1 Circularity: 多圆 $= \frac{4\pi \times \text{Area}}{\text{Perimeter}^2}$ 图是 1, 越小越不圆

aspect ratio: $\frac{w}{h}$

$x, y, w, h = cv2.boundingrect(c)$ extent = $\frac{\text{Contour Area}}{\text{Bounding box Area}}$

面积可用于过滤, 把面积低于 threshold 的去掉

0.785 circle $\frac{\pi r^2}{4r^2}$
 square = 1.0

完整流程 转灰度 → 高斯模糊 → canny → 形态学闭运算
 → 查找过滤 contour
 convex hull = $\frac{A}{\text{convexHull Area}}$ solidity
 morphological closing = dilation then erosion
 expands region, closing small gaps in the edges
 shrinks them back to approximately original size
 1-3 pixels get bridged broken curves → closed loops

1. The first part of the paper is devoted to a general discussion of the problem of the origin of life. It is shown that the problem is not only a scientific one, but also a philosophical one. The author discusses the various theories of the origin of life, and shows that the most plausible one is the theory of spontaneous generation.

2. The second part of the paper is devoted to a detailed discussion of the theory of spontaneous generation. The author shows that this theory is based on the assumption that life is a necessary consequence of the laws of chemistry. He then discusses the various experiments which have been conducted to test this theory, and shows that the results are in favor of spontaneous generation.

3. The third part of the paper is devoted to a discussion of the theory of evolution. The author shows that this theory is based on the assumption that life is a result of a gradual process of change. He then discusses the various theories of evolution, and shows that the most plausible one is the theory of natural selection.

4. The fourth part of the paper is devoted to a discussion of the theory of the origin of the human race. The author shows that this theory is based on the assumption that the human race is a result of a gradual process of change. He then discusses the various theories of the origin of the human race, and shows that the most plausible one is the theory of evolution.

5. The fifth part of the paper is devoted to a discussion of the theory of the origin of the universe. The author shows that this theory is based on the assumption that the universe is a result of a gradual process of change. He then discusses the various theories of the origin of the universe, and shows that the most plausible one is the theory of evolution.

6. The sixth part of the paper is devoted to a discussion of the theory of the origin of the earth. The author shows that this theory is based on the assumption that the earth is a result of a gradual process of change. He then discusses the various theories of the origin of the earth, and shows that the most plausible one is the theory of evolution.

7. The seventh part of the paper is devoted to a discussion of the theory of the origin of the solar system. The author shows that this theory is based on the assumption that the solar system is a result of a gradual process of change. He then discusses the various theories of the origin of the solar system, and shows that the most plausible one is the theory of evolution.

8. The eighth part of the paper is devoted to a discussion of the theory of the origin of the life on earth. The author shows that this theory is based on the assumption that life is a result of a gradual process of change. He then discusses the various theories of the origin of life on earth, and shows that the most plausible one is the theory of evolution.

9. The ninth part of the paper is devoted to a discussion of the theory of the origin of the human mind. The author shows that this theory is based on the assumption that the human mind is a result of a gradual process of change. He then discusses the various theories of the origin of the human mind, and shows that the most plausible one is the theory of evolution.

10. The tenth part of the paper is devoted to a discussion of the theory of the origin of the human soul. The author shows that this theory is based on the assumption that the human soul is a result of a gradual process of change. He then discusses the various theories of the origin of the human soul, and shows that the most plausible one is the theory of evolution.

CV Week 4 Match

1. 希望找 $(x_i, y_i) \leftrightarrow (x'_i, y'_i)$

1.1 能用 edge to match? 平行于边的变化无法察觉. (无法区分边上的点)

1.2 好 feature 两个 hard-req

① repeatability 不同视角. 光照仍能检测. | or 找不到

② distinctiveness: 不易与别的附近区域 mix or wrong match

↓

Corner 优于 edge. 在二维方向有明显变化. 完整位置约束

good Features To Track

2. Harris Corner: 两边都变为 corner

2.1 Window shift energy: 设窗口中心在 (x, y) , shift by (u, v) , Intensity change

$$E(u, v) = \sum_{(i, j) \in W} [I(i+u, j+v) - I(i, j)]^2$$

对所有方向 (u, v) 都小: flat.

某些大, 某些小: edge. 所有大 corner

对小 (u, v) 有 $I(i+u, j+v) \approx I(i, j) + I_x u + I_y v$

$$E(u, v) \approx \sum_{(i, j) \in W} (I_x u + I_y v)^2 = [u \ v] \underbrace{\sum_{(i, j) \in W} \begin{bmatrix} I_x^2 & I_x I_y \\ I_x I_y & I_y^2 \end{bmatrix}}_M \begin{bmatrix} u \\ v \end{bmatrix}$$

$\nearrow$ x 梯度 $\nearrow$ 相关性

$$M = \begin{bmatrix} \sum_w I_x^2 & \sum_w I_x I_y \\ \sum_w I_x I_y & \sum_w I_y^2 \end{bmatrix} = \begin{bmatrix} A & C \\ C & B \end{bmatrix}$$

$\nwarrow$ y 梯度

也就是 λ_1 大, x 上梯度大, 有纵边

λ_2 大, y 上梯度大, 有横边

都大, 有 corner

M 可以拆成 $M = Q \Lambda Q^T$ $\Lambda = \text{diag}(\lambda_1, \lambda_2)$

$$E(u, v) = [u \ v] M \begin{bmatrix} u \\ v \end{bmatrix} = [u \ v] Q \Lambda Q^T \begin{bmatrix} u \\ v \end{bmatrix} = \Delta \cdot [u' \ v'] \Lambda \begin{bmatrix} u' \\ v' \end{bmatrix} = \lambda_1 (u')^2 + \lambda_2 (v')^2$$

2.2 怎么求 λ_1, λ_2

$$\det(M) = \lambda_1 \lambda_2$$

$$\det(M - \lambda I) = 0 \Rightarrow \lambda^2 - (a+c)\lambda + (ac-b^2) = 0$$

$$\lambda_1, \lambda_2 = \frac{(a+c) \pm \sqrt{(a+c)^2 - 4(ac-b^2)}}{2}$$

2.3 定义 Harris Response

$$R = \det(M) - k(\text{trace}(M))^2$$

$$k = 0.04 \sim 0.06$$

$\begin{cases} \text{Corner } (\lambda_1, \lambda_2 \text{ 大}): R \gg 0 \\ \text{Flat}: R \approx 0 \\ \text{Edge } (\lambda_2 = 0): R \approx -k\lambda_1^2 < 0 \end{cases}$

$$\lambda_1 \lambda_2 = ac - b^2 = \det(M)$$

$$\lambda_1 + \lambda_2 = a + c = \text{trace}(M)$$

$$AB - C^2 - k(A+B)^2$$

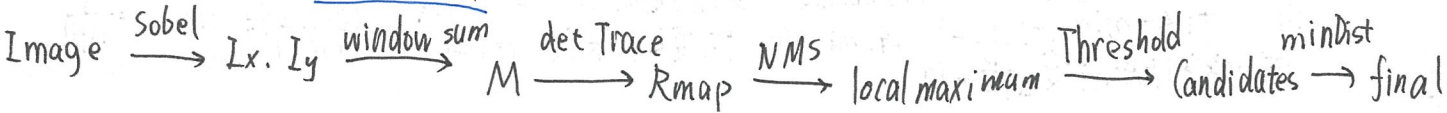
3. Harris 后处理: Harris 得到的是 Harris Response Map

3.1 NMS: 局部邻域内只保留最大的 $\Rightarrow$ exactly one pixel $\rightarrow \begin{bmatrix} 0 & 0 & 0 \\ 0 & 8 & 0 \\ 0 & 0 & 0 \end{bmatrix}$

3.2 Threshold: $R > R_{\max} \cdot \text{quality level}$

↑ 整图最大 ↑ 保留点质量比例

3.3 min Distance: 按 R 从大到小排序. 每选一个, 排除其 minDistance 范围内候选点. $\Rightarrow$ 让点均匀分布



Harris 三大问题 : 无尺度. 无方向, 无 descriptor

缺点	具体	SIFT解决
scale	远近变化窗口不同	DoG scale space: 6
orientation	旋转 patch 对不齐	主方向估计. 描述前旋转对齐
descriptor	只有位置. 无法验证是否角点	128-D gradient histogram descriptor

4. SIFT

4.1 scale invariance : DoG

$$L(x, y, \sigma) = G(x, y, \sigma) * I(x, y)$$

尺度6下平滑图

标准差为6高斯核

建 Gaussian Pyramid (blur at $\uparrow 6$). 相邻层相减

$$D(x, y, \sigma) = L(x, y, k\sigma) - L(x, y, \sigma)$$

k是相邻尺度间倍数.

DoG response

$$G(x, y, k\sigma) - G(x, y, \sigma) \approx (k-1)\sigma^2 \nabla^2 G$$

approximates scale-normalized Laplacian of Gaussian

SIFT在DoG体中找三维极值

$$D(x, y, \sigma) > D(x', y', \sigma') \quad \forall (x', y', \sigma') \in N_{26}(x, y, \sigma)$$

N_{26} 当前尺度同层8邻居

$$D(x, y, \sigma) < D(x', y', \sigma') \quad \forall (x', y', \sigma') \in N_{26}(x, y, \sigma)$$

上下尺度各9邻居. 总计26邻居

一个包含5个可用尺度层次 要5+2张DoG $\Rightarrow$ 5+3张高斯模糊图像

极大极小 $\rightarrow$ 候选 keypoint. 并记录其 σ

4.2 Lowe's keypoint sieve : 过滤低质量点, 以上一步keypoint为中心, 按照6个角点

上一步90% junk

① Low Contrast. $|D(\hat{x})| < k_{ao3} \quad \downarrow \text{contrast threshold} \quad \Rightarrow \text{Discard}$

② Edge Response (Hessian filter)

令 Hessian 为 $\begin{bmatrix} D_{xx} & D_{xy} \\ D_{xy} & D_{yy} \end{bmatrix} = H$ 二阶导 $\Rightarrow$ 特征值 α, β 且 $\alpha \geq \beta$ 令 $r = \frac{\alpha}{\beta}$

若 $\frac{\text{tr}(H)^2}{\det(H)} < \frac{(r+1)^2}{r^2}$ 作为保留条件. ratio 太高, 则 r 大 $\Rightarrow$ 为边, 舍弃

和 Harris 很像, 不过这里用曲率

4.3 Rotation Invariance 以上一步keypoint为中心 按照6确定圆形邻域大小

在一个大圆区域
①为36个方向

36-bin 梯度直方图. 取峰值方向为 dominant orientation; 旋转坐标系对齐

$$\theta(x, y) = \text{atan2}(L(x, y+1) - L(x, y-1), L(x+1, y) - L(x-1, y)) \quad \theta^* = \arg \max_b h(b)$$

SIFT会返回此

extremas 多的keypoint

$$\tilde{x} = R(-\theta^*)(x - x_0)$$

- 阶导 若次要峰值达到主峰值 $\geq 80\%$ $\Rightarrow$ keypoint is duplicate

4.4 Lighting Invariance (raw pixel 敏感, pixel difference 也 sensitive, gradient direction robust; 每个keypoint已有 position, scale, orientation $\Rightarrow$ 需 compact numerical vector encoding local appearance)

SIFT descriptor 标准构造

以关键点的 pos, σ , orientation 为基准, 将 patch 规范化到同一坐标系下, 算方向直方图, 生成

① 尺度归一, 方向对齐后的 key point 周围取 16×16 邻域 ② 划分为 4×4 16个 cell descriptor

③ 每个 cell 统计 8-bin 梯度方向直方图 ④ 拼接成 $16 \times 8 = 128$ R^{128} ⑤ L2 归一化 $f \leftarrow \frac{f}{\|f\|_2}$

⑥ clip $f_i \leftarrow \min(f_i, 0.2)$ 再归一化 恢复单位长度 防止单个主导梯度 dominates descriptor

② 减少梯度光照突变的影响

5. Feature Mapping: 从 descriptor 相似到几何一致

5.1 Brute-force matching

直接找最近 l_2 -norm

$$D_A = \{d_i^A\}_{i=1}^N \quad D_B = \{d_j^B\}_{j=1}^M \quad \text{对每个 } d_i^A \text{ 寻找最近邻 } j^* = \arg\min_j \|d_i^A - d_j^B\|_2$$

$O(MN)$ 会带来 false matches. 存在 scalability wall

5.2 KNN + Lowe Ratio Test 对

找 $k=2$ 最近邻 比较 ratio of d repetitive textures 效果最好

$$\frac{d_1}{d_2} < 0.75 \rightarrow \text{keep} \quad \begin{cases} *d_1 \ll d_2 \text{ 好} \\ d_1 \approx d_2 \rightarrow \text{discard} \\ \text{差别很大} \rightarrow \text{discard} \end{cases}$$

解决局部 descriptor 唯一性问题

但不检查多匹配是否满足同一个几何变换

5.3 RANSAC (Random Sample Consensus)

正确匹配需满足一个全局几何变换

Homography

$$x' \sim Hx \quad x = \begin{bmatrix} x \\ y \\ 1 \end{bmatrix} \quad x' = \begin{bmatrix} x' \\ y' \\ 1 \end{bmatrix} \quad \text{8个自由度}$$

① 随机采4对非共线匹配点

② 求候选 H

③ 对所有匹配求 rejection error

④ error < threshold 为 inliers

⑤ 重复 N 次

⑥ 留 inlier 最多 H

⑦ 用所有 inliers 重新拟合 H

minimum for a homography

$$s \begin{bmatrix} x_2' \\ y_2' \\ 1 \end{bmatrix} = H \begin{bmatrix} x_1' \\ y_1' \\ 1 \end{bmatrix}$$

$$N = \frac{\log(1-p)}{\log(1-w^s)} \quad \begin{matrix} \text{采样点数} \\ s=4 \end{matrix}$$

什么时候 hold

- ① scene is planar
 - ② rotates around its optical centre
- 若有显著 depth change 和 相机平移
视差会破坏模型

为什么用直方图不是 raw pixels

① illumination robustness. Uniform brightness change doesn't affect ∇ and its dir normalization removes absolute magnitude

② Spatial Tolerance: 4×4 cell grids gives tolerance to small shift

③ Compactness: 128 floats $\approx$ 512 bytes per point. 省空间

cumulative drift via bundle adjustment 多张一起优化
wide-angle distortion via cylindrical projection
brightness differences via multi-band blending 多频段

CV Week 5

小图 $\Rightarrow$ no room for multi-scale feature hierarchies, texture statistics, or spatial pyramid

1. Pre-requisite

1.1 scalar $|a+b| \leq |a| + |b|$

1.2 vector $\|x\| = \sqrt{\sum x_i^2}$

1.2 Matrix $C = AB$

红特征 $\left\{ \begin{array}{l} \text{combinational explosion: 组合数太多, 类内变化大} \\ \text{domain expertise bottleneck: design} \\ \text{brittleness: 只能用在指定 conditions} \end{array} \right.$

Why DL

- ① 数据
- ② 算力
- ③ 架构

$$C_{ik} = \sum_j A_{ij} B_{jk}$$

$$\|A\|_{\text{Frob}} = \left[\sum_{i,j} A_{ij}^2 \right]^{\frac{1}{2}}$$

$$\|aA\|_{\text{Frob}} = |a| \cdot \|A\|_{\text{Frob}}$$

$$\|A+B\|_{\text{Frob}} \leq \|A\|_{\text{Frob}} + \|B\|_{\text{Frob}}$$

2. KNN

2.1 1-NN: 训练时不真正学参数, 记住全部 data, labels, test 时找最近的. 复制 label

$$i = \underset{i}{\operatorname{argmin}} d(x, x_i) \quad \hat{y} = y_i^* \quad d \text{ 是 distance function 常用 } L_1$$

2.2 从 1-NN $\rightarrow$ KNN

1-NN 对 noise sensitive 若最近点刚好异常. wrong.

$N_k(x)$ = the set of k nearest training samples to x

$$\hat{y} = \underset{c}{\operatorname{argmax}} \sum_{i \in N_k(x)} \mathbb{I}(y_i = c) \quad \text{多数投票}$$

2.3 超参选择, k 和 $d(\cdot, \cdot)$

不能直接选, 应 train set 训模型

但 DL 用得少

Cross-Validation

训练分多个 fold.

一个 val, 剩余 train

然后平均验证结果

2.4 基本不用

$\left\{ \begin{array}{l} \text{① pixel distance 无语义信息} \\ \text{② 慢 } O(ND) \end{array} \right.$ $\xrightarrow{\text{dimension}}$ $\xrightarrow{\text{训练样本}}$

val set 选超多

test set 最后评一次

3. Linear Classifier

4. SVM

4.1 Loss 要求正负类别比错误类别至少多一个 margin $\angle$ set to 1

$$S = f(x_i; W) \in \mathbb{R}^c \quad c \text{ 是类别数}$$

$$L_i = \sum_{j \neq y_i} \begin{cases} 0, & \text{if } S_{y_i} > S_{j+1} \\ S_j - S_{y_i} + 1, & \text{otherwise} \end{cases}$$

S_j 是模型对 j 的打分

S_{y_i} 是模型对正负类别打分

对每个错误类别, 希望 $S_{y_i} > S_{j+1}$

$$L = \frac{1}{N} \sum L_i \quad \text{margin}$$

4.2 optimal W 并不唯一

$f(x; W) = Wx$ s.t. $L=0$ 可令 $W = nW$ $n \uparrow$ margin $\downarrow$ overfit

$$n(W_j^T x_i - W_{y_i}^T x_i + \frac{1}{n}) \quad \text{等效 margin } 1 \rightarrow \frac{1}{n}$$

引入 regularization

$$L = \frac{1}{N} \sum_{i=1}^N L_i(x_i; W, y_i) + \lambda R(W) \quad \begin{cases} L_1 & R(W) = \|W\|_1 = \sum_k |W_k| \\ L_2 & R(W) = \sum_k W_k^2 \end{cases}$$

小 W 对应大 margin, 泛化更好; 限 W 能让输出平滑

4.3 SVM score $\rightarrow$ prob: softmax $\frac{e^k}{\sum e^i}$

$$\text{Cross entropy } L = -\frac{1}{N} \sum p_i \log(q_i) = -\frac{1}{N} \sum \log \frac{e^{S_{y_i}}}{\sum_j e^{S_j}}$$

当 $S_{y_i} \rightarrow \infty$ 整个 $\rightarrow 1$ $L \rightarrow 0$ 当 $S_{y_i} \rightarrow 0$ 那 $L \rightarrow \infty$

$$L_i = -S_{y_i} + \log \sum_j e^{S_j}$$

若初始化, 所有 score 等 $L_i = -\log \frac{1}{C} = \log C$

4.4 Comparison

SVM Loss $\rightarrow$ score

Softmax cross entropy 输出 prob

5. Optimization

5.1 Random Search 低效

上游 $\frac{\partial L}{\partial \text{output}}$

局部 $\frac{\partial \text{output}}{\partial \text{input}}$

往下游传 upstream x local

5.2 GD

5.2.1 SVM Loss的 ∇ $L_i = \sum_{j \neq y_i} \max(0, W_j^T X_i - W_{y_i}^T X_i + 1)$

对某错误类别 j $m_j = W_j^T X_i - W_{y_i}^T X_i + 1$

$$\frac{\partial L_i}{\partial W_j} = X_i \quad \text{对正确类边权重} \quad \frac{\partial L_i}{\partial W_{y_i}} = -X_i$$

若有多类别超出margin, $\frac{\partial L_i}{\partial W_{y_i}}$ 会累加

$$\frac{\partial f(x)}{\partial x} = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

$$z = (\langle x, w \rangle - y)^2 \quad \text{令 } a = \langle x, w \rangle \quad b = a - y \quad z = b^2$$

$$\frac{\partial z}{\partial b} = 2b$$

$$\frac{\partial b}{\partial a} = 1$$

$$\frac{\partial z}{\partial a} = 2b$$

$$\frac{\partial a}{\partial w} = x$$

$$\frac{\partial z}{\partial w} = 2bx = 2(a-y)x$$

$$= 2(\langle x, w \rangle - y)x$$

决策边界

$$S = W_1 X_1 + W_2 X_2 + b = 0$$

i 和 j 的决策边界在分数相等处

$$S_i = S_j \Rightarrow (W_i - W_j) X + (b_i - b_j) = 0$$

每个区域是凸的, 任何凸区域内连线仍在凸区域内

XOR

若有 $W_1 X_1 + W_2 X_2 + b = 0$ 分开

$b \leq 0$ 代入 $(0, 0)$ 类别为 0

$$W_1 + W_2 + b \leq 0 \quad W_2 + b > 0 \quad W_1 + b > 0$$



$$P_i = \frac{e^{S_i/T}}{\sum_{j=1}^C e^{S_j/T}}$$

$T < 1$ 低温 指数随得分差异, 分布更尖锐

$T \rightarrow 0$ max 独热

$T > 1$ 高温 压缩得分间差异

$$\| -\eta \lambda \| < 1 \Rightarrow \eta < \frac{2}{\lambda} \quad \eta \text{ 超过这个, 每一步会增加 } |\theta| \text{ 远离最小值}$$

批量 GD

$$O(N \times D \times C)$$

$N \uparrow$ 输入维度 $C \uparrow$ class

单样本 ∇ 是真实梯度有噪声无偏估计

$$E[\nabla L_i(\theta)] = \frac{1}{N} \sum_{i=1}^N \nabla L_i(\theta) = \nabla L(\theta)$$

$$CE = -\log(P_y) = -\log\left(\frac{e^{z_j}}{\sum_{i=1}^C e^{z_i}}\right) = -z_j + \log\left(\sum_{i=1}^C e^{z_i}\right) = -z_j + m + \log\left(\sum_{i=1}^C e^{z_i - m}\right)$$

$$\text{令 } m = \max_j z_j$$

保证指数输入小于 0

CV Week 6 MLP → CNN

Linear Classifier 只能学习线性边界. 无法处理 XOR 等非线性

MLP 引入 activation function → 非线性 提升表达能力. 但对图像全连接参数量巨大

CNN 利用 image 局部性, 平移共享, 空间结构. 大幅减少参数并学局部视觉特征

1. 激活函数

$g'(x) \leq 0.25$ 每次乘 0.25

且总在 $[0, 1]$, 随深度加大, 问题放大, 消失

Leaky ReLU $\max(0.1x, x)$ ↓ Maxout $\max(w_1^T x + b_1, w_2^T x + b_2)$ ELU = $\begin{cases} x, & x \geq 0 \\ \alpha(e^x - 1), & x < 0 \end{cases}$

问题 ① sigmoid { ① 过大过小 saturated

$$\frac{\partial L}{\partial w_j} = \frac{\partial L}{\partial z} \cdot a_j \leftarrow \text{总正}$$

② 非 0-centered 若神经元输出总为正, 局部梯度也为正. 所有权重符号受 upstream gradient 的统一控制, 要么全增, 要么全减, 无法针对更复杂.

不适合隐藏层

$$g'(x) = g(x)(1 - g(x)) \geq 0$$

② tanh $(-1, 1)$ zero-centered. 比 sigmoid 适合当中间层. 但仍会 saturate

③ ReLU pros: ① 正区间不 saturate

- ② easy - compute
- ③ 实践比 sigmoid / tanh 收敛快
- ④ Sparsity 约一半神经元关闭, 有计算优势

cons: ① output 非 zero-centered

② 可能出现 dead neuron 输出总负. 无法更新 { ① 大lr, 把b推到很大负值, $w^T x + b < 0$ ② 不巧的初始化 - 一开始就死

leaky ReLU $\begin{cases} x, & x > 0 \\ \alpha x, & x \leq 0 \end{cases}$ 可以监控 $\max(x, \alpha x)$ 减小lr 使用小的正偏置 $b=0.01$ 0.01 死了还学

PRELU α 可学

④ ELU $f(x) = \begin{cases} x, & x > 0 \\ \alpha(\exp(x) - 1), & x \leq 0 \end{cases}$

cons: exp 难算

- pros ① 正区间不 saturate
- ② 负区间接近 zero mean
- ③ 负区间 saturate 增加对 noise 的 robustness.

对于非常负, 趋近于 $-\alpha$. 产生接近 0 mean 的输出. 解决 ReLU 和 sigmoid 零中心问题是

⑤ Maxout $f(x) = \max(w_1^T x + b_1, w_2^T x + b_2)$ 可以泛化 ReLU 和 Leaky ReLU

- pros { ① liner regime
- ② non-saturating
- ③ does not die

problem: 一个神经元要训两倍参数

建议

- ① hidden 用 ReLU. 但小心 $b \neq 0$
- ② 可以用 Leaky ReLU, Maxout ELU 获得边际提升
- ③ hidden 不用 sigmoid / tanh. 只用来限制 output range

2. CNN

2.1 paras:

$$H \times W \times C + 1 \quad \begin{matrix} \text{filter} \\ \downarrow \\ \text{bias} \end{matrix}$$

C 要求输入 channel - 致

- 层可以 stack k 个 filter 出来 k 层

2.2 输出公式

$$\text{Output size} = \frac{N - F}{S} + 1$$

无padding
输入尺寸 filter size
Stride

$$N=7, F=3 \text{ stride } 1 \text{ 就是 } \frac{7-3}{1} + 1 = 5$$

输出非整数, 不能合法卷积.

$$\text{Receptive field} = L \times (F - 1) + 1$$

不pad的话, size 快速减小

$$Y(i, j) = \sum_{c=0}^{C-1} \sum_{m=0}^{F-1} \sum_{n=0}^{F-1} W(c, m, n) X(c, i+m, j+n) + b$$

有padding $\frac{N+2P-F}{S} + 1$ 相当于原来尺寸变大 translation equivariance

2.3 Receptive field 当层则和 filter size 一致

2.4 Pooling layer

- ① 减少 representation size
- ② 让特征更 manageable
- ③ 对每个 activation map 独立操作
不影响 channel 数

深比赛好

- ① less param
- ② more nonlinearities
- ③ 更易优化

Input $\rightarrow$ [conv $\rightarrow$ ReLU $\rightarrow$ Pool] $\times N \rightarrow$ flatten $\rightarrow$ FC $\rightarrow$ softmax

CV Week 7 Optimization of CNNs

1. Batch Normalization

1.1 Input: $x \in \mathbb{R}^{N \times D}$

对每个 feature $\mu_j = \frac{1}{N} \sum_{i=1}^N x_{i,j}$ $\sigma_j^2 = \frac{1}{N} \sum_{i=1}^N (x_{i,j} - \mu_j)^2$

Xavier $\text{Var}(W) = \frac{1}{n_{in}}$

He $\text{Var}(W) = \frac{2}{n_{in}}$

$$\hat{x}_{i,j} = \frac{x_{i,j} - \mu_j}{\sqrt{\sigma_j^2 + \epsilon}}$$

ϵ 小常数

1N: 对 HxW 归一

GN: 通道组 + 空间

但 zero mean 和 unit variance 可能过于严格

can perfectly undo

引入 learnable scale and shift $\gamma, \beta \in \mathbb{R}^D$

希望恢复 identity mapping 可以学习 $\gamma_j = \sigma_j, \beta_j = \mu_j$

$$y_{i,j} = \gamma_j \hat{x}_{i,j} + \beta_j$$

$$y_{i,j} = \sigma_j \frac{x_{i,j} - \mu_j}{\sigma_j} + \mu_j = x_{i,j}$$

1.2 训练时差距 训练时 μ_j, σ_j^2 来自当前 mini-batch

never hurts

$\mu_j^{\text{test}} = \text{running mean}$ $(\sigma_j^2)^{\text{test}} = \text{running variance}$

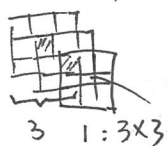
1.3 插在 FC 和卷积层后 在 activation 前

2. 架构

2.1 AlexNet 用较大卷积核, 网络深度有限, FC 层参数量大, 训练依赖工程技巧

2.2 VGG: 用小卷积堆深度.

3个 3x3 stride 1 等价一个 7x7 的 conv



$$3 \times 2 + 1 = 7$$

$$R = 1 + 3 \times (3 - 1) = 7$$

每多一层 kxk, receptive field 增加 k-1

$$R_k = R_{k-1} + (k - 1)$$

$R_0 = 1$ 输入层的一个像素

$$1 \text{ 个 } 7 \times 7 = 49 \text{ Cin Cout} = 49c^2$$

若等

$$3 \text{ 个 } 3 \times 3$$

$$3 \times 9c^2 = 27c^2$$

2.3 Plain deep network: 54层比20层 training error 和 test error 都更差
不是 overfit 因为 54层比 20层 train 还差 是 performance degradation

2.4 ResNet $F(x) = H(x) - x$

$$y = F(x) + x$$

1. Stack Residual Blocks

2. 每个 residual block 两 3x3 conv

3. 周期性 double filters 用 stride 2 downsample spatial size

4. 开头有 stem conv 如 7x7 stride 2

5. 末尾无大 FC layers 只用 global average pooling + FC 1000

$$H \times W \times C$$

$$z_c = \frac{1}{HW} \sum_{i=1}^H \sum_{j=1}^W x_{i,j}, c \Rightarrow z \in \mathbb{R}^C$$

减少参数, 降过拟合

理论支持 variable size input

保留 channel-level semantic representation

3. Optimizer

1. SGD问题
- ① Poor Conditioning: 一方向快, 一方向慢. Steep上来回震荡, flat上进展缓慢

$$\kappa = \frac{\lambda_{\max}(H)}{\lambda_{\min}(H)}$$
 $\rightarrow$ 最flat上曲率方向 $\rightarrow$ 最steep上的曲率方向
 H : Hessian matrix κ 越大, 优化越难
 二阶导
 - ② Local Minimum
 / Saddle point: 无更新方向, 易stuck
 高维更常见
 saddle point: 一些方向曲率为正
 但一阶导为0 一些方向曲率为负
 - ③ Noisy Gradient 随机采样有噪声
 历史梯度方向累积

2. SGD + Momentum

$$v_{t+1} = \rho v_t - \alpha \nabla f(x_t)$$

$$x_{t+1} = x_t + v_{t+1}$$

在 flat direction 累积速度, 加快前进 steep上来回震荡, 互相抵消, 减少 jitter
 通过惯性帮助越过小 local minimal / saddle point

3. Adagrad

$$r_t = r_{t-1} + g_t \odot g_t \quad x_{t+1} = x_t - \alpha \frac{g_t}{\sqrt{r_t} + \epsilon}$$

steep $\rightarrow$ damp $\rightarrow$ 历史梯度平方和 $\rightarrow$ 逐元素开方
 flat $\rightarrow$ accelerated

问题 r_t 单增, step size $\rightarrow 0$

4. RMSProp: 修正 Adagrad lr 衰减问题. 指数滑动平均

$$r_t = \beta r_{t-1} + (1-\beta) g_t \odot g_t \quad x_{t+1} = x_t - \alpha \frac{g_t}{\sqrt{r_t} + \epsilon}$$

5. Adam

- 阶矩 $m_t = \beta_1 m_{t-1} + (1-\beta_1) g_t$
 = 阶矩 $v_t = \beta_2 v_{t-1} + (1-\beta_2) g_t^2 \quad (g_t \odot g_t)$

预热从 10^{-7} 开始
 前 1000-5000 步线性增加
 对 Adam 很有用

$$\hat{m}_t = \frac{m_t}{1-\beta_1^t} \quad \hat{v}_t = \frac{v_t}{1-\beta_2^t} \quad x_{t+1} = x_t - \alpha \frac{\hat{m}_t}{\sqrt{\hat{v}_t} + \epsilon}$$

初期估计偏小

4. Regularization

generalization error = test error - training error

4.1 Early Stopping 监控 val loss/acc 当 val performance $\downarrow$ stop

4.2 Model Ensembles 训多个 $p_k(y|x)$

性能 $\uparrow$ 开销大

$$p(y|x) = \frac{1}{K} \sum_{k=1}^K p_k(y|x) \quad \hat{y} = \arg \max_x p(y|x)$$

4.3 L2 weight decay

惩罚大的

$$+ \lambda \|W\|_2^2$$

梯度 $+ 2\lambda W$

$$L_1 + \lambda \|W\|_1 \Rightarrow \text{sparse weight}$$

$$\text{Elastic net} + \lambda_1 \|W\|_1 + \lambda_2 \|W\|_2^2$$

L_1 稀疏 L_2 平滑

4.4 Dropout

4.4.1 训练时每个 neuron 采样一个 Bernoulli mask p 是留存 prob

dropout后激活 $\downarrow$ 原激活
 $\tilde{h}_i = m_i h_i$

4.4.2 ensemble 解释 每个binary对应一个network. 2^n 种 n 为unit数

4.4.3 test. 所有neuron都启用

$h_i^{\text{test}} = \phi_i$ 缩放, 使期望一致

现代实现常用 inverted dropout 训练除以 p . 测试不用缩放

希望 training 引入 randomness. 测试 ~~并~~ 通过 average out randomness

5. Data Augmentation

$(x, y) \mapsto (T(x), y)$ T 是Augmentation

① horizontal flip

② Random crops and scales

③ Color Jitter

6. Transfer Learning

小数据集 policy: 冻 CNN backbone. 重新初始化分类, 重新训

也可以训一些 layer

CV Week 8 RNN + Transformers

1. RNN

$$h_t = f_w(h_{t-1}, x_t) \quad y_t = g_w(h_t) \quad h_t = \tanh(W_{xh}x_t + W_{hh}h_{t-1} + b_h)$$

$$y_t = W_{yh}h_t + b_y$$

同一组参数 W 在所有 timestamp 共享

1.1 RNN 计算图与任务类型

1.1.1 many $\rightarrow$ many $L = \sum_{t=1}^T L_t(y_t, \hat{y}_t)$ frame-level classification
language modeling

1.1.2 many $\rightarrow$ one $h_T = f_w(h_{T-1}, x_T) \quad y = g(h_T)$ sentiment classification
video classification

1.1.3 one $\rightarrow$ many
image captioning

$$h_t = \tanh(W_{xh}x_t + W_{hh}h_{t-1} + W_{hv}v)$$

把上一步输出塞进来 v 是 image feature

1.2 character-level LM

输入 [Start] $\rightarrow$ 下-词分布 $\rightarrow$ sample $\rightarrow$ 喂回去 $\rightarrow$ 生成

$$o_t = W_{yh}h_t + b_y \quad p(x_{t+1}=k | x_{\leq t}) = \frac{\exp(o_{t,k})}{\sum_j \exp(o_{t,j})}$$

END / 最大长度

2. Seq2Seq: Encoder-Decoder 架构

$[x_1:T] \rightarrow u_t = RNN(x_t, u_{t-1}) \rightarrow y_t = g_v(y_{t-1}, h_{t-1}, c)$ Decoder
或 $h_0 = f_w(z)$ z 可以是 encoder 最终状态 / 所有 encoder states 集合 c 是 context vector 可以是 h_0

2.1 Bottleneck: c 必须承载整个输入序列所有信息. h_0 和 c 要生成后续所有.
但 c 是 fixed-length vector $\Rightarrow$ 信息瓶颈

3. RNN + Attention

3.1 固定 context vector $c \rightarrow$ 每个 timestamp 不同 context

$$y_t = g_v(y_{t-1}, h_t, c) \rightarrow y_t = g_v(y_{t-1}, h_t, c_t)$$

动态选择输入相关区域

3.2 Captioning 中的 attention

① CNN 提取 spatial features $z \in R^{H \times W \times D}$

② 对 decoder 第 t 步, 算 alignment score

$$z_{i,j} \in R^D \rightarrow \text{MLP}$$

$$e_{t,i,j} = f_{att}(h_{t-1}, z_{i,j})$$

当前步与 $z_{i,j}$ 契合度

③ 归一化

$$a_{t,i,j} = \frac{\exp(e_{t,i,j})}{\sum_{\{m,n\} \in H \times W} \exp(e_{t,m,n})}$$

$$\sum_{i,j} a_{t,i,j} = 1$$

④ context vector

$$c_t = \sum_{i,j} a_{t,i,j} z_{i,j}$$

3.3 soft att vs hard attn

soft: $\sum_{i,j} a_{t,i,j} z_{i,j}$ 连续可微, 可 end2end BP

hard att: 离散选某-区域 $i_t \sim \text{categorical}(a_t) \quad c_t = z_{i_t}$ 采样不可微, 要 RL policy gradient

4. NLP 中的 Attention : Translation

Input $x = x_{0:T}$ 输出 $y = y_{0:T'}$ 长度可以不同

Encoder: $z_i = \text{RNN}(x_i, u_{i-1})$

↑ RNN hidden state

Decoder

$$e_{t,i} = f_{\text{att}}(h_{t-1}, z_i)$$

归一化

$$a_{t,i} = \frac{\exp(e_{t,i})}{\sum_{i=1}^T \exp(e_{t,i})}$$

context vector

$$c_t = \sum_{i=1}^T a_{t,i} z_i$$

decoder 输出

$$y_t = g_v(y_{t-1}, h_{t-1}, c_t)$$

不同 target word 对应不同 source word

Attention 自动学习 alignment

且无需人工标注对齐关系

5. General Attention Layer

$x = \{x_i\}_{i=1}^N$ $x_i \in \mathbb{R}^D$ query $h \in \mathbb{R}^D$

alignment $e_i = f_{\text{att}}(h, x_i)$

attention weights $a_i = \frac{\exp(e_i)}{\sum_{j=1}^N \exp(e_j)}$

$c = \sum_{i=1}^N a_i x_i$ 操作本身 permutation

invariant

把 MLP Alignment 换成点积

$e_i = h^T x_i \Rightarrow \text{scaled}$

$$e_i = \frac{h^T x_i}{\sqrt{D}}$$

D 是维度

若 h 和 x_i 各维度近似独立且方差为 1, 那么点积方差大约与 D 成正比. 维度越大, logits 越大, softmax 越易饱和. 梯度变小. 除以 $\sqrt{D}$ 控制数值

5.1 多个 query

$q_i: M$

$$e_{i,j} = \frac{q_j^T x_i}{\sqrt{D}}$$

j 是 query 索引, i 是 input index

$$a_{i,j} = \frac{\exp(e_{i,j})}{\sum \exp(e_{i',j})}$$

$$y_j = \sum_{i=1}^M a_{i,j} x_i$$

6. Transformer

$$X \in \mathbb{R}^{N \times D}$$

~~WQ~~

$$Q \in \mathbb{R}^{M \times D_K}$$

$D_K = D$

$$W_Q \in M \times N$$

$$X W_V$$

$$V \in \mathbb{R}^{N \times D_K}$$

$$\text{Att} = \text{softmax}\left(\frac{QK^T}{\sqrt{D_K}}\right)V$$

7. PE

$$p_j = \text{pos}(j)$$

$$\text{pos}(): \mathbb{R}^N \rightarrow \mathbb{R}^d$$

不能自然泛化到超过 $T_{\max}$ 位置间关系
learned positional embedding 需 model 自学
 $PE \in \mathbb{R}^{T_{\max} \times d}$ $P_t = P[t]$

Sinusoidal

$$PE(\text{pos}, 2i) = \sin\left(\frac{\text{pos}}{10000^{2i/d_{\text{model}}}}\right)$$

$$PE(\text{pos}, 2i+1) = \cos\left(\frac{\text{pos}}{10000^{2i/d_{\text{model}}}}\right)$$

高维低频捕捉
长距离变化

$$\text{MultiHead}(Q, K, V) = \text{concat}(\text{head}_1, \dots, \text{head}_H) W_O$$

CV Lec 9 Detection / Segmentation

1. Semantic Segmentation

Input $I \in \mathbb{R}^{H \times W \times c}$

总共 class 数

$$\Rightarrow Y = \{1, \dots, C\}^{H \times W}$$

$\Rightarrow$ 只关心 what class each pixel belongs to. 忽略同类间实例差异

1.1 直觉方法 sliding window: each pixel 周围取 patch, CNN 判中心 pixel.

$$P_{i,j} = I[i-r:i+r, j-r:j+r] \Rightarrow \text{预测 } \hat{y}_{i,j} = \underset{c}{\operatorname{argmax}} f_{\theta}(P_{i,j})_c$$

$\uparrow$ 在类别中

用了 context, 但未解决 shared computation

1.2 全卷积, 对一张图一次提 feature map, 对每个空间分类

$$F = g_{\theta}(I) \in \mathbb{R}^{D \times H \times W}$$

feature map

$$S = h_{\phi}(F) \in \mathbb{R}^{C \times H \times W}$$

每 pixel 出 prob

$$Y_{i,j} = \underset{c}{\operatorname{argmax}} S_{c,i,j}$$

对 pixel 取最大类

但一直维持 $H \times W$ 计算大, CNN 的 pooling 会降尺寸 $\Rightarrow$ 不满足 pixel-wise prediction.

FCN 折中: encoder $\rightarrow$ downsample $\xrightarrow{\text{decoder}} \uparrow$ upsample

2. Upsampling

2.1 非学习, ① Nearest Neighbour 直接复制值

不可学

② Bed of Nails Upsampling 插 0 配合后续卷积

③

2.2 Max Unpooling, 和 Max pooling 成对出现 下采样记 max, 还记 max 所在位置 即 pooling switches. 上采样把低 reso 值放回原来最大值处, 其余填 0

pros: 部分恢复 pooling 前空间结构

cons: 只恢复最大激活位置, 仍 sparse

2.3 Transposed Convolution: deconvolution: 对每个元素都“撒”出一个乘以该输入值的 filter pattern, 多个 pattern 在输出空间处叠加

对 input $[a, b]$ filter 为 $[x, y, z]$ stride=1

$$\Rightarrow a[x, y, z] + b[-, x, y, z] \Rightarrow [ax, ay+bx, az+by, bz]$$



重叠部分相加 一个值被 3×3 filter 成了 9 个值

stride gives the ratio between movement in output and input

3. U-Net

$$O_{\text{transposed}} = (I-1) \times S - 2P + K$$

$\uparrow$ 步长 $\downarrow$ padding $\downarrow$ 核大小

3.1 FCN 问题: encoder 深层 feature 语义强, 但空间细节弱; 浅层 feature size 大, 但语义弱

3.2 solu: encoder-decoder + skip connection encode 中高 reso 直传 decoder 拼接 copy and crop

$$D'_i = \text{Concat}(D_i, E_i) \quad D_{i-1} = \text{Conv}(D'_i)$$

4. Object Detection:

输出不是像素类别,而是对像集合

类别 $\downarrow$ bounding box
 $\{c_i, b_i\}_{i=1}^N$

box左上或中心
 $b_i = (x_i, y_i, w_i, h_i)$
长宽

单目标检测可看成 classification + localization

网络共享一个 feature vector $z \in \mathbb{R}^{4096}$

分类头 $p = \text{softmax}(W_c z + b_c)$ $\hat{b} = W_b z + b_b$ 定位头

loss $L = L_{cls} + \lambda L_{loc} = -\log p_y + \lambda \|\hat{b} - b^*\|_2^2$

多目标 detection 困难

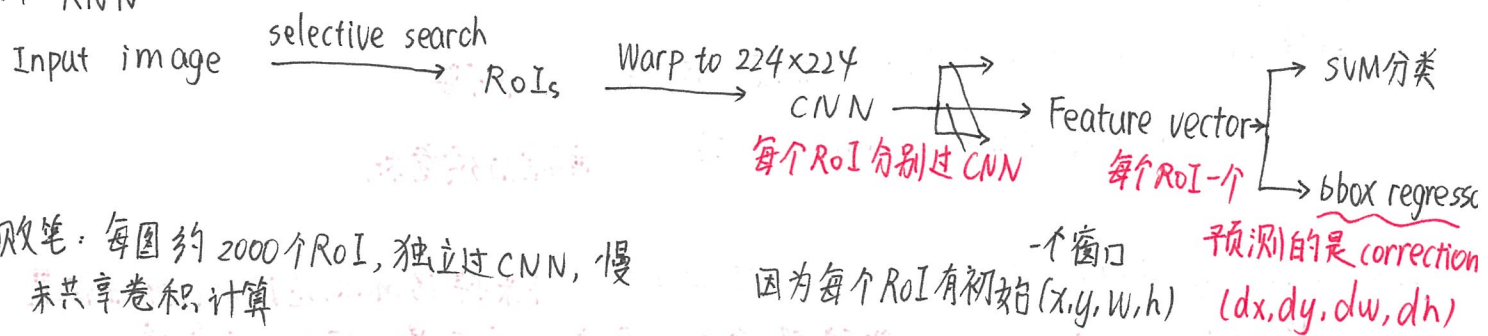
- ① 图 object 类不同, 输出长度不固定
- ② CNN 分类头输出固定长多维向量, 无法直接输出可变长度集合

做法 候选
对图像中大量 crop 分别分类为 object/background

问题: 所有位置, 尺度, 长宽比都枚举, 计算量巨大

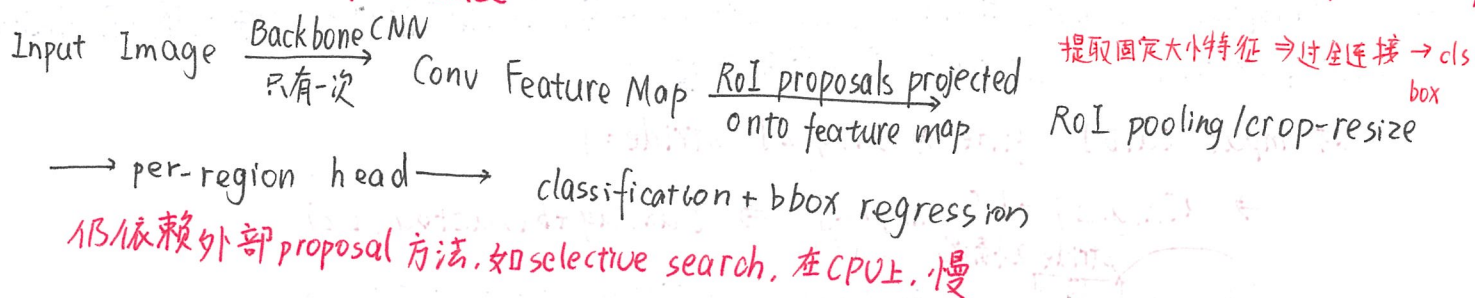
↓ Region Proposals: Selective Search 选举 candidate crops

4.1 RCNN



败笔: 每图约 2000 个 RoI, 独立过 CNN, 慢
未共享卷积计算

4.2 Faster RCNN: 先对 image 做 backbone 得到 feature map, 再把 proposal 映射到 feature map 上做 RoI Pool / Crop + Resize



4.3 Faster R-CNN: 提出 Region Proposal Network (RPN), RPN 与 detector 共享 backbone feature map, 直接在 feature map 每个位置预测 anchor 是否包含 object 以及对应 box correction

RPN: 一个点有 k different anchor boxes of different size/scale. 且做边界微调

让 conv 输出 $K \times H \times W$ (Objectiveness), Score $K \times H \times W$ by "Objectness Score"

用 top ~ 300 as our proposals $\rightarrow$ 输出 box transforms 4KxHxW

全连接

RoI-Pooling 量化取整, 舍去小数部分, 取子窗口内像素最大值, 空间精度低, 存在“未对准”

4.4 Single-Stage Detector: YOLO | SSD | Retina Net

4.4.1 Two-stage: 先生成 proposals 再进行分类回归

Single stage: 直接在 dense grid 上预测 boxes 和 classes

NMS { ① 按置信度对所有检测排序
② 保留置信度最大的
③ 移除与其 IoU 超过 threshold 的结果

4.4.2 YOLO: 把图划成 $S \times S$ 网格, 每个网格预测 B 个 bounding boxes 和类别 prob
 $\Rightarrow S \times S \times (5B + C)$ $B: (dx, dy, dh, dw, confidence)$. C 是类别 (预测), 含 environment

confidence = $P(\text{object}) \cdot \text{IoU}(\hat{b}, b^*)$ 很像 RPN, 但具体到类别!
 $\Rightarrow P(c|\text{object}) P(\text{object}) \text{IoU}(\hat{b}, b^*)$ 每个 cell 还 pred $p(c|\text{object})$

pros: 快, 整张图只 forward 一次

cons: dense prediction 产生大量候选框, 需 NMS
 且早期 YOLO 对小物体, 密集物体不如 two-stage

4.5 DETR: Object Detection with Transformer

不用 anchors. IoU threshold, NMS, 把 object detection 视为 set predictions
 输出固定 N 个 object queries 的预测 $\{(c_i, b_i)\}_{i=1}^N$ 无 object 的 query 被视为 no-object

图 $\rightarrow$ CNN $\rightarrow$ PE $\rightarrow$ encoder $\rightarrow$ decoder $\rightarrow$ FFN $\rightarrow$ prediction heads

Set Prediction + Hungarian Matching: queries (一组可学习向量)

找最优匹配 $\hat{b} = \underset{G \in G_N}{\operatorname{argmin}} \sum_{j=1}^M L_{\text{match}}(c_j, b_j, (p_{\hat{c}_j}, \hat{b}_{\hat{c}_j}))$
 预测出 $\{(p_j, \hat{b}_j)\}_{i=1}^N$ $N \geq M$.
 L_{match} 含分类与 box 代价

找唯一最优一对一匹配. 避免重复预测!

- ① 消除很多手工工程组件 (anchors, IoU thresholds, NMS)
- ② 不用 duplicate prediction 带来的后处理 (NMS 等, 传统方法会重复预测)
- ③ 拥有 global reasoning (self-attention 带来, 任意两 pixel 直接交互)
- ④ end2end. 可以一起训, 不用像传统方法分阶段训 (RPN $\rightarrow$ 分类头)

5. Instance Segmentation

5.1 Mask-RCNN

在 faster R-CNN 上, 除了 RoI feature 上再加一个 mask head $\rightarrow C \times 28 \times 28$ masks

用 RoI Align 解决了空间对齐问题 (用 bilinear interpolation 保留更准确空间对齐, 因为 mask 训练只对 GT 类别对应的 mask 算 binary cross entropy prediction 对 pixel level alignment 很敏感)

6. Visualization: 理解 model 到底学了什么

6.1 filter visualization

6.2 Saliency Map: 看类别分数对输入像素的梯度. 给定类别 c 的 unnormalized score $S_c(I)$

$$G = \frac{\partial S_c}{\partial I} \in 3 \times H \times W$$

$$M_{i,j} = \max_{k \in \{R, G, B\}} \left| \frac{\partial S_c}{\partial I_{k,i,j}} \right|$$

若某像素轻微变化强烈影响 score.
 $\Rightarrow$ 其对类别重要
 取 RGB max 初始化.

6.3 Class Activation Mapping (CAM)

依赖特殊结构: 卷积层后接 GAP, 再接 linear classifier

$$f \in \mathbb{R}^{H \times W \times k} \xrightarrow{\text{GAP}} F \in \mathbb{R}^k \xrightarrow{\text{FC}} \text{SERC}$$

$$F_k = \frac{1}{HW} \sum_{h,w} f_{h,w,k} \quad S_c = \sum_k W_{k,c} F_k = \frac{1}{HW} \sum_k W_{k,c} \sum_{h,w} f_{h,w,k}$$

类别 c 的 CAM map 是 $M_c(i,j) = \sum_{k=1}^k W_{k,c} f_{h,w,k}$

$$= \frac{1}{HW} \sum_{h,w} \sum_k W_{k,c} f_{h,w,k}$$

- Grad-CAM
- ① Any layer, with Activation $A \in \mathbb{R}^{H \times W \times k}$
 - ② 算 S_c 对 A 梯度 $\frac{\partial S_c}{\partial A} \in \mathbb{R}^{H \times W \times k}$
 - ③ GAP $\alpha_k = \frac{1}{HW} \sum_{h,w} \frac{\partial S_c}{\partial A_{h,w,k}}$
 - ④ $M_{h,w}^c = \text{ReLU}(\sum_k \alpha_k A_{h,w,k})$

梯度: 增强 feature channel $\rightarrow$ score $\uparrow$?

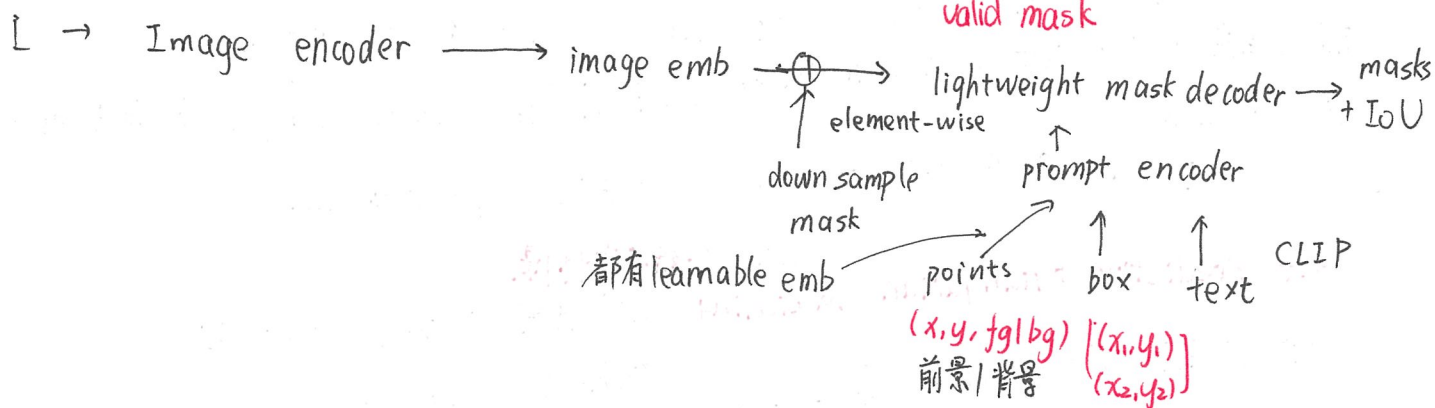
只保留对 c 有正贡献大

7. SAM Image + Prompt $\rightarrow$ Valid Mask

point box text mask

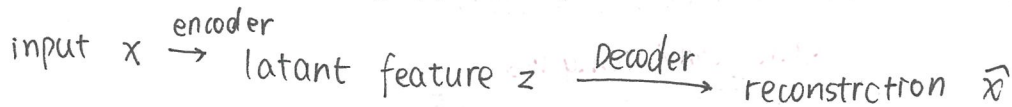
promptable seg

可能 prompt ambiguous, 但仍要返回 valid mask



1. Autoencoder $\rightarrow$ VAE 压缩表征 $\rightarrow$ 生成模型

1.1 Autoencoder: 无监督特征学习



$$L_{AE} = \|x - \hat{x}\|^2$$

$\rightarrow$ 比如说保留训练好的 encoder 来初始化 supervised model 做 pred

降维、特征学习, 辅助预训练. 非好生成模型.

AE 学到的是离散, 不规则, 可能有空洞的 latent space

$x \mapsto z$ 没有保证 z 分布连续光滑可采样

若只生成一样图片那还好 但很难生成类似有差异的图

1.2 VAE

假设 AE 输出确定性 latent vector $z = f_{\phi}(x)$. 但 VAE 输出 latent distribution $q_{\phi}(z|x)$

$$q_{\phi}(z|x) = \mathcal{N}(\mu(x), \Sigma(x)) \text{ 或 } q_{\phi}(z|x) = \mathcal{N}(\mu(x), \text{diag}(\sigma^2(x)))$$

生成时 $z \sim p_z(x)$

$$z \sim p(z)$$

$$\hat{x} = \text{Decoder}_{\theta}(z)$$

$$\text{选 } p(z) = \mathcal{N}(0, I)$$

$$L = \|x - \hat{x}\|^2 + D_{KL}(q(z|x) || p(z))$$

reconstruction error

decoder 能从 sampled latent z 恢复原图

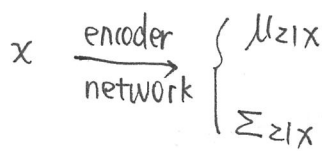
会惩罚 encoder 把不同样本压到孤立 cluster

要求 decoder 输出 latent distribution 接近标准正态

怎么算?

$$z = \mu + \sigma \epsilon$$

$$D_{KL}(q(z|x) || p(z)) = \frac{1}{2} \sum_{i=1}^n (\sum_i^2 + \mu_i^2 - 1 - \log \sum_i^2) = \frac{1}{2} \sum_{j=1}^d (6_j^2 + \mu_j^2 - 1 - \log 6_j^2)$$



Sample

$$z|x \sim \mathcal{N}(\mu_{z|x}, \Sigma_{z|x})$$

Decode

$$p_{\theta}(x|z)$$

$\rightarrow \hat{x}$

$$= \frac{1}{2} \sum_{j=1}^d (6_j^2 + \mu_j^2 - 1 - \log 6_j^2)$$

VAE 易模糊

分布 multi-modal

$p(z) = \mathcal{N}(0, I)$ 不同模式之间被平滑连接, 采样可能落在类间

pixel-wise reconstruction loss 本身有 averaging effect. 若多输出合理, l_2

loss 倾向于预测平均图像

$$\hat{x} = \underset{\hat{x}}{\text{argmin}} E[\|x - \hat{x}\|^2] \text{ 趋向条件均值}$$

VAE: limited diversity

2. Generative Adversarial Networks (GANs)

不要求重建某样本, 而是学从随机噪声到真实数据分布的 mapping

$$z \sim p_z(z)$$

$$\hat{x} = G(z)$$

不知对应 $\Rightarrow$ 不用 reconstruction loss

$$\text{目标 } G(z) \sim p_{\text{data}}(x)$$

D : Discriminator

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

训练: ① fix G , train D

② fix D , train G to fool D

对 D , 希望 $D(x) \rightarrow 1$ $D(G(z)) \rightarrow 0$

对 G , 希望 $D(G(z)) \rightarrow 1$

训练不稳定, 生成图像更锐利

- GAN 问题
- ① Bad convergence
 - ② Mode collapse
 - ③ Discriminator too strong/too weak
 - ④ No explicit likelihood

DL L8 LLM

2.1 Cycle GAN & Style GAN

Cycle GAN: 解决 unpaired image-to-image translation

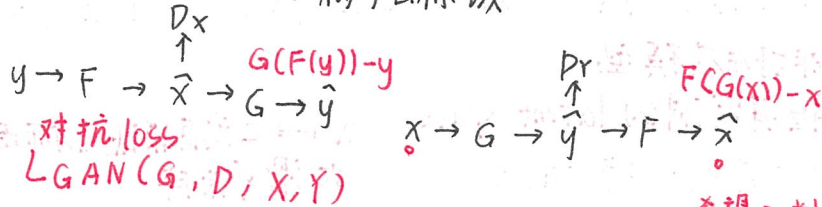
$$G: X \rightarrow Y$$

$$F: Y \rightarrow X$$

D: 判断是否属于目标域

$$L_{cyc}(G, F) = E_{x \sim P_{data}(x)} [\|F(G(x)) - x\|_1] + E_{y \sim P_{data}(y)} [\|G(F(y)) - y\|_1]$$

Cycle consistency loss



对 G, minimize

$$E_{x \sim P_{data}(x)} [(D(G(x)) - 1)^2]$$

对 D, minimize

$$E_{y \sim P_{data}(y)} [(D(y) - 1)^2] + E_{x \sim P_{data}(x)} [D(G(x))^2]$$

Style GAN: 把 latent code 映射到 style space

希望 D 判 G(x) 为真

AdaIN 模块

3. Diffusion

3.1 forward

$$x_t = \alpha_t x_{t-1} + \beta_t \epsilon$$

$$\epsilon \sim N(0, I)$$

$$q(x_t | x_{t-1}) = N(x_t; \sqrt{1 - \beta_t} x_{t-1}, \beta_t I)$$

$$x_t \rightarrow N(0, I)$$

3.2 Reverse

$$x_T \sim N(0, I)$$

$$x_{t-1} = \frac{1}{\sqrt{\alpha_t}} \left(x_t - \frac{1 - \alpha_t}{\sqrt{1 - \alpha_t}} \epsilon_\theta(x_t, t) \right) + \beta_t \epsilon$$

3.3 训练 采样

$$x_0 \sim q(x_0)$$

$$t \sim \text{Uniform}(\{1, \dots, T\})$$

$$\epsilon \sim N(0, I)$$

$$\nabla_\theta = \|\epsilon - \epsilon_\theta(\sqrt{\alpha_t} x_0 + \sqrt{1 - \alpha_t} \epsilon, t)\|^2$$

直到收敛

$$x_T \sim N(0, I)$$

for $t = T, \dots, 1$ do

$$z \sim N(0, I) \text{ if } t > 1, \text{ else } z = 0$$

$$x_{t-1} = \frac{1}{\sqrt{\alpha_t}} \left(x_t - \frac{1 - \alpha_t}{\sqrt{1 - \alpha_t}} \epsilon_\theta(x_t, t) \right) + \beta_t z$$

end for

return x_0

3.4 Stable diffusion

直接在像素空间生成计算大 $H \times W \times 3$

Image x $\xrightarrow{\text{Autoencoder Encoder}}$ latent z

latent z $\xrightarrow{\text{Diffusion U-Net}}$ $\hat{z}$

AE decoder

Generated image

Conditioning (Guiding the output): text prompts, semantic maps, reference images

U-Net 还用了 cross-attention 去看 text prompt

3.4 ControlNet 很难精确描述空间结构

3.4.1 text prompt limitation: ① pose ② 透视 ③ 边缘 ④ 深度

可输入 Openpose, Canny Edge, Depth Map, Semantic Map

锁住 pretrained diffusion model
加一个可训练控制分支

$y_c = y + \text{ControlNet}(x)$ zero conv 是 1×1 conv, 权重 bias 初始化为 0

初始训练时 $\text{ControlNet}(x) = 0$

$$y_c = y$$

复制一份 encoder, 在旁边加一份可训练副本
用零卷积逐步注入原始模型

离输入近的
逐渐获得权重