

## 基于空间卷积神经网络模型的图像显著性检测

高 东 东, 张 新 生

(西安建筑科技大学 管理学院, 西安 710055)

**摘 要:** 针对现有显著性检测方法鲁棒检测效果较差这一问题, 提出一种新的基于空间卷积神经网络的显著性检测算法。利用去均值、归一化的预处理方法获取目标候选区。一方面通过引入卷积变换网络, 建立提取显著物体上下文信息的全局模型, 得到相应的目标检测信息显著图; 另一方面构建特征子网络结构输出 6 维变换矩阵, 经过空间变换模块改造输入图像, 获取边缘信息。将空间变换网络输出的局部置信度融入到全局显著信息图, 求取特征表达最大值, 实现显著性与非显著性划分, 完成显著性检测任务。实验结果表明, 该算法不仅在同等条件下显著检测的 AUC 值得到了提高, 并且生成的显著性图聚焦点突显, 鲁棒检测效果得到明显改善。

**关键词:** 显著性检测; 特征融合; 卷积神经网络; 空间变换网络; 显著图

**中文引用格式:** 高东东, 张新生. 基于空间卷积神经网络模型的图像显著性检测[J]. 计算机工程, 2018, 44(5): 240-245.

**英文引用格式:** GAO Dongdong, ZHANG Xinsheng. Image Saliency Detection Based on Spatial Convolutional Neural Network Model[J]. Computer Engineering, 2018, 44(5): 240-245.

## Image Saliency Detection Based on Spatial Convolutional Neural Network Model

GAO Dongdong, ZHANG Xinsheng

(School of Management, Xi'an University of Architecture and Technology, Xi'an 710055, China)

**[Abstract]** To overcome the drawback of bad robustness detection of the existing saliency detection methods, this paper presents a novel saliency detection based on spatial convolutional neural network. The candidate object areas are conformed via preprocessing methods of the removing mean and normalization. By introducing the Convolutional Neural Network (CNN), a global model that captures the context information of saliency objects is constructed, and the corresponding target detection saliency map is obtained. On the other hand, the feature sub-network structure is established to output the 6-dimensional transformation matrix, which is used to transform input image through the spatial deformation module to obtain edge information of salient object. The output of the spatial transformer network local confidence coefficient is introduced into the global information saliency map to seek the maximum value of the feature expression. The discrimination of saliency and non-saliency is realized to accomplish the saliency detection task. Experimental results show that the proposed algorithm improves AUC accuracy under the same condition, generates a saliency map of focus highlighting and achieves impressive robust detection results.

**[Key words]** saliency detection; feature fusing; Convolutional Neural Network (CNN); Spatial Transformer Network (STN); saliency map

**DOI:** 10.19678/j.issn.1000-3428.0048490

### 0 概述

显著性检测已成为计算机视觉领域的重要研究课题, 它的目的就是捕捉吸引人类注意的像素或区域。随着信息技术的快速发展, 如何设计及利用计算机模型处理这些以爆炸式速度增长的图像及视频信息, 对于弥补人与电脑在视觉理解中的差距具有重要的研究意义和应用价值。

快速且以数据驱动的自底向上的显著性检测方式成为目前研究的主流<sup>[1]</sup>。其中, Itti 模型<sup>[2]</sup>最早提出利用图像底层特征的“中心-周边差”进行显著性检测。文献[3]则在文献[2]的基础上提出基于图模型的显著性检测算法 (GBVS), 通过计算不同特征图的马尔科夫链平衡分布获取显著图。文献[4]使用颜色与亮度低级特征计算图像子区域像素与其邻域的像素平均特征向量之间的距离获得显著值,

**基金项目:** 陕西省自然科学基金 (2016JM6023)。

**作者简介:** 高东东 (1990—), 男, 硕士研究生, 主研方向为机器学习、图像显著性检测; 张新生, 教授。

**收稿日期:** 2017-04-24      **修回日期:** 2017-05-27      **E-mail:** 18729530885@163.com

该方法快速且易于实现。文献[5]提出基于局部特征对比度的显著性检测,学习局部信息进行显著性评估,该方法在显著区域的边缘处产生较高的显著值。上述方法都是基于手工设计特征提取图像特征信息,它不能有效地捕捉显著目标的深层特征,也没有均匀地突出显著区域。

然而,深度学习方案的提出显著地提高了系统的检测性能。其中,文献[6]提出利用双层深度玻尔兹曼机(DBM)判别显著性区域,增强了基于低级特征学习的能力,但此方法需要经过大量学习训练,复杂度较高。卷积神经网络(Convolutional Neural Network, CNN)被广泛用于学习图像数据中的丰富特征信息<sup>[7-9]</sup>,在显著区域检测<sup>[10-11]</sup>与对象轮廓检测<sup>[12-13]</sup>中都取得较为理想的检测结果。另一方面,一些方法也引入或改进了新的网络架构。文献[14]设计了一个能同时预测人眼固定点及分割突出目标的CNN模型,并结合预先训练的VGG网络和SALICON数据集定位出显著区域。文献[15]提出基于超像素卷积神经网络的显著性检测方法,从颜色唯一性和颜色分布2个超像素序列来计算显著性与非显著性区域。该方法突出了显著物体的内部信息,但在显著边缘的聚焦点不清晰,尤其处理非均匀背景或较复杂场景的图像检测性能较差。

基于以上分析,为了能够更准确地检测出人眼感兴趣的显著点,本文采用完全数据驱动的方式,以CNN模型为基础模型,添加空间变换思想,学习图片中应该关注的区域并且通过仿射变换放大该区域。再使用Adam优化器来训练网络,提高模型收敛速度和检测的准确度。

## 1 CNN与STN算法

### 1.1 卷积神经网络

CNN是首个真正成功训练多层神经网络的学习算法,它通过描述数据的后验概率进而实现网络结构的优化。其基本结构由输入层(Input)、卷积层(Convolution)、池化层(Pooling)、全连接层(Fully Connected)及输出层(Output)组成。

卷积层利用卷积核进行特征抽取,而池化层旨在特征的计算。在每一个卷积层中,先将 $t-1$ 层输出的特征图与可学习的卷积核进行卷积操作,再通过激活函数进行非线性变换,则可输出该层特征图 $z'_j$ ,如式(1)所示。

$$z'_j = f\left(\sum_{i \in N_j} k'_{i,j} * z_i^{t-1} + b'_j\right) \quad (1)$$

其中, $z'_j$ 和 $z_i^{t-1}$ 是当前层的特征图与上一层特征图, $t$ 代表层数, $k'_{i,j}$ 表示从上一层第 $i$ 个特征图到该层第 $j$ 个特征图的卷积核, $*$ 表示2维卷积操作, $b'_j$ 代表神经元偏置, $f(x) = \max(0, x)$ 为网络激活函数。

下采样层的主要目的就是降维,即减少卷积层的特征维数,对池化层中每个大小为 $n \times n$ 的池进行“最大值(max pooling)”或“均值(mean pooling)操作,进一步获得抽样特征,输出过程可表示为:

$$z'_j = f(\beta'_j \text{down}(z_i^{t-1}) + b'_j) \quad (2)$$

其中, $\beta'_j$ 为网络权重, $\text{down}(\cdot)$ 表示下采样函数。为了最大程度地保留图像的纹理信息,本文选用最大值池化进行采样。

### 1.2 空间变换神经网络

为了更好地处理显著性检测任务,本文在CNN模型上插入空间变换网络(Spatial Transformer Network, STN)<sup>[16]</sup>。它应用在深度卷积网络之前,与深度网络一起进行端对端训练,增加模型特征提取的旋转不变性。即以一种动态的方式对输入图像进行任意变形(扭曲、拉伸等),这种方式的组成步骤如下:

**步骤1** 由预处理后的图像集 $U$ ,通过优化后的VGG网络结构输出变换参数 $\theta$ ,即 $2 \times 3$ 维的空间变换矩阵。

**步骤2** 由上述参数 $\theta$ ,经过仿射变换实现逆向坐标映射,得到输入与输出图像之间的采样网格 $T_\theta$ 。

**步骤3** 将采样网格 $T_\theta$ 输出的结果通过双线性插值技术进行处理,得输出变换图像 $V$ 。

设 $(x_i^s, y_i^s)$ 对应输入 $U$ 中的像素坐标, $(x_i^t, y_i^t)$ 对应输出 $V$ 中的像素坐标,步骤2中空间变换参数 $\theta$ 通过仿射变换进行逆向坐标映射,仿射变换的原理为:

$$\begin{pmatrix} x_i^s \\ y_i^s \\ 1 \end{pmatrix} = T_\theta(G_i) = A_\theta \begin{pmatrix} x_i^t \\ y_i^t \\ 1 \end{pmatrix} = \begin{pmatrix} \theta_{11} & \theta_{12} & \theta_{13} \\ \theta_{21} & \theta_{22} & \theta_{23} \end{pmatrix} \begin{pmatrix} x_i^t \\ y_i^t \\ 1 \end{pmatrix} \quad (3)$$

其中, $A_\theta$ 表示仿射变换,即对网格进行仿射变换,再把变换后的网格逆向映射到原图,针对原图中对应坐标点的像素值填充变换后的网格,得到真实像素值,即变换后图像 $V_i^c$ 的表达式为:

$$V_i^c = \sum_{n=1}^H \sum_{m=1}^W U_{nm}^c \max(0, 1, |x_i^s - m|) \max(0, 1, |y_i^s - n|) \quad (4)$$

经过以上理论推导,即实现了对输入图像的空间变换操作,增大了本文模型获取具有高激活值的显著区域的概率。

## 2 空间卷积网络的显著性检测模型

显著性检测就是在图像区域内找出特征最明显的子区,即检测的关键集中在特征学习上,基于CNN与STN具有的强大的特征提取能力,本文提出应用空间卷积特征学习的显著性检测模型。图1给出了本文模型的整体框架图,其中包括图像预处理、全局与局部显著性估计、显著性融合及显著训练。

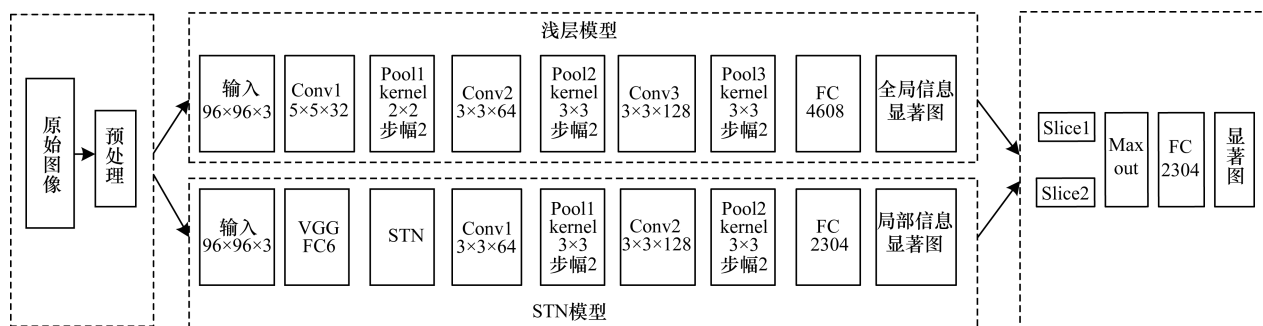


图 1 显著性检测模型框架图

## 2.1 图像预处理

给定待处理图像数据集  $I = \{I_{\text{train}}, I_{\text{val}}, I_{\text{test}}\}$ , 图像预处理过程包括以下步骤:

**步骤 1** 将  $I$  全部缩放到  $96 \times 96$  大小, 以适应本文模型的输入。

**步骤 2** 移除图像的平均亮度值, 以快速提取图像内容。具体操作为先计算  $I_{\text{train}}$  的平均像素值, 再将  $I$  中亮度值都减去这个平均值 ( $I - \bar{I}_{\text{train}}$ )。

**步骤 3** 将  $I$  中像素除以 255, 归一化到  $[0, 1]$  中, 以加快模型收敛速度。

## 2.2 模型结构

检测框架主体由全局 (shallow model) 与局部 (STN model) 块组成。其中, 全局模型包含 3 个卷积层和 3 个池化层, 在每一个卷积层和全连接层的输出过程当中, 都经过 ReLU 非线性化处理。其连接模型使用:

$\text{Conv1} \rightarrow \text{MaxPool1} \rightarrow \text{Conv2} \rightarrow \text{MaxPool2} \rightarrow \text{Conv3} \rightarrow \text{MaxPool3} \rightarrow \text{FC}$

同时, 局部块也采用了相同的设计思路, 每层的详细参数说明如图 2 所示。

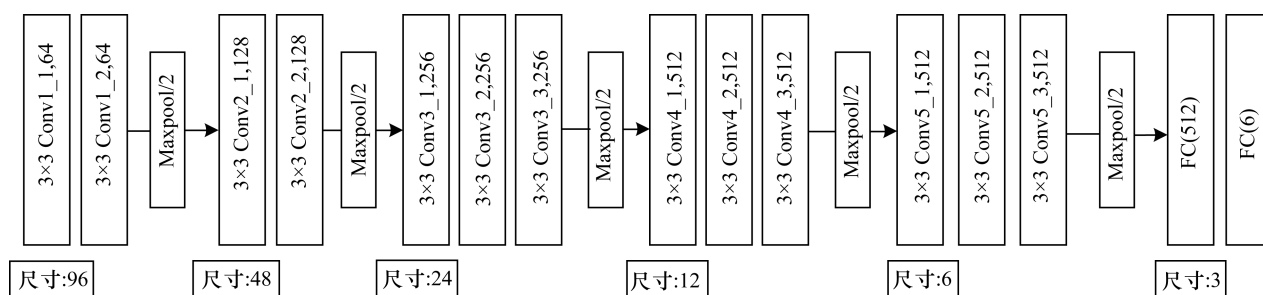


图 2 优化后的 VGG 结构

结构设计遵循如下步骤:

**步骤 1** 使用式 (1) 和式 (2), 组合 3 个卷积层与池化层, 以构建 shallow model。

**步骤 2** 通过采用及优化 VGG 模型以学习出 6 个变换参数  $\theta$ , 用于下一步的空间变换。

**步骤 3** 依据步骤 2 的空间参数, 经过式 (3) 与式 (4), 实现对输入图像的空间聚焦变换。

**步骤 4** 根据步骤 3 得到变换后的图像, 再嵌入到 CNN, 进一步组建 STN model。

**步骤 5** 融合步骤 1 及步骤 4 中所得全局显著图与局部显著图, 以检测出图像特征表示的最大值。

**步骤 6** 在最后一层全连接层前, 增加一层 maxout 激活函数, 以过滤出高显著性特征表达, 完成显著性检测。

在步骤 2 的输出中, 由于本实验最终目的并不是分类, 即对 VGG-16 结构进行优化以适应本文的显著性检测任务。具体做法为: 去掉最后一层原本

用于输出类别的全连接层, 将第 1 个 FC 设置为 512, 第 2 个 FC 大小为 6。图 2 显示了优化后 VGG 结构的详细参数设置。

## 2.3 模型训练与优化策略选择

本文采用 Adam 法代替传统的梯度下降法来训练模型, 即每一次迭代学习率都随前面的梯度值进行自适应的调整。相比传统的梯度下降, Adam 法计算速度更加高效, 且更适用于大规模的数据集训练。采用这种参数更新策略使得参数比较平稳, 也加快了模型的收敛速度。同时, 训练过程中将 iSUN 数据集作为训练数据, 为了尽量利用有限数据进行训练, 使用数据预处理与数据提升技术来增加训练样本。具体做法是在数据预处理阶段, 通过一系列的随机变换将训练数据进行提升, 使模型捕获不到任何两张完全相同的图像, 进而缓解过拟合现象, 增强模型的泛化能力。最后, 所有样本均进行分批量的训练, 每一个批量为 32 个样本, 动量为 0.9。鉴于硬件的内存限制,

这是可以使用的最大的批量数。Adam 中的参数设置: $\alpha=0.001, \beta_1=0.9, \beta_2=0.999, \delta=10^{-8}$ 。

为了防止模型训练过拟合,本文选择了一些优化策略:1) L2 正则化对网络参数进行约束;2) 网络损失函数采用均方误差 (Mean Squared Error, MSE),以滤除输出中的高空间频率;3) 在 VGG 网络中的第一个全连接层后接入 dropout 层。其本质是在训练过程中的权重调整时,在隐藏层以一定的概率丢弃一些神经元(丢失率设为 0.5)。同时,为使 VGG 网络更好地适应显著性检测任务,把预训练好的权重值加载到优化后的 VGG 结构中。最后,对新加入网络层的权重也进行初始化,并随整个网络结构一起训练。

### 3 实验结果与分析

实验使用 Python2.7、NumPy 和深度学习库 Keras 来实现,采用相对复杂的公开数据集进行性能的检测。与现有显著性检测方法进行比较,并从不同的角度分析实验结果。

#### 3.1 数据集

本实验选择了在 LSUN 挑战中提出的 2 个数据集,因为它们与通常使用的数据集相比具有一定的复杂性。这些数据集是:1) SALICON<sup>[17]</sup>。它通过鼠标跟踪点击突出点进行构建,且所有图像来自流行的 MS COCO 图像数据库,包含 80 个物体类别,共有 10 000 个训练图像、5 000 个验证图像、5 000 个测试图像;同时,具有丰富的语境信息和分辨率为  $640 \times 480$  像素的图像。2) iSUN。它通过使用网络摄像头并从亚马逊土耳其机器人 (Amazon Mechanical Turk) 的眼动追踪进行收集,且所有的图像选择于自然场景 SUN<sup>[18]</sup> 图像数据库,共有 6 000 个训练图像、926 个验证图像、2 000 个测试图像。在选用的数据集中,有些图片的背景特征特别复杂,这也增大了显著性检测的难度。

#### 3.2 评价指标

为了定量评估本文模型的性能及有效性,选用相似性 (Similarity)、AUC (Area Under Curve) 及相关系数 (Correlation Coefficient, CC) 作为评价指标。其中, AUC 值为 ROC 曲线与  $x$  轴之间距离的积分,能直观地衡量检测结果的优劣程度,取值范围为  $0 \sim 1$ ,且 1 为最佳状态。CC 表示模型预测的显著性区域与人眼关注区域之间的线性关系,由 2 个映射对应的标准偏差之间的协方差计算得出,该取值范围为  $0 \sim 1$  之间,且结果值越大说明两者之间具有较好的线性关系,显著性检测准确率也就最高。

### 3.3 结果分析

在本实验中,所有数据分为 2 个子集:一个用于训练(75%);另一个用于测试(25%)。55% 的数据(训练集)用于训练模型。20% 的数据(验证集)用于测试模型训练结果。测试数据(25%)用于获取本文所提出模型的显著性值。图 3 展示了训练过程的数据分布情况。同时,将本文算法与其他参与者进行比较,包括 LCYLab<sup>[19]</sup>、基线 Itti<sup>[2]</sup>、改进 Rare 2012<sup>[20]</sup>、基线 GBVS<sup>[3]</sup>、Xidian<sup>[21]</sup>、基线 BMS<sup>[22]</sup>、WHU IIP<sup>[19]</sup>。

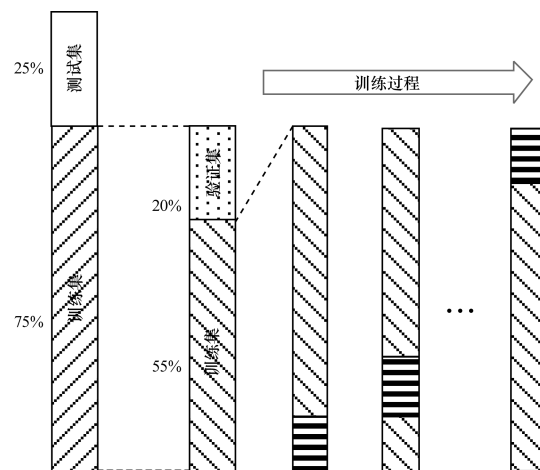


图 3 实验数据集分布

表 1 给出了本文算法在 2 个数据验证集上的检测结果。可以看出,各个指标值都较为理想,在 2 个数据集中呈现出非常有竞争力的结果,也体现了本文算法在数据集之间具有较好的鲁棒性。表 2、表 3 详细列举了本文算法与其他显著性算法在 2 个数据集上所得的实验值。对于所考虑的每个指标,本文算法的检测值与真实显著值的相似度是最高的,且在 2 个数据集上都取得较好的 AUC 值,表明了本文所提模型的优越性能。图 4 直观地比较了显著性预测的平均 AUC。与最先进的方法相比,本文算法的检测精度在 SALICON 数据集上提高了约 6.42%,在 iSUN 数据集上提高了 7.32%。整个模型结构表现出较小的偏差,多角度证明了本文提出的算法具有更强的鲁棒性。

表 1 不同数据验证集上的测试结果

| 系数           | iSUN 数据集 | SALICON 数据集 |
|--------------|----------|-------------|
| CC           | 0.62     | 0.61        |
| AUC Shuffled | 0.68     | 0.77        |
| AUC Borji    | 0.81     | 0.83        |

表 2 iSUN 测试集上的检测结果比较

| 系数           | LCYLab  | 基线 Itti | 改进 Rare2012 | 基线 GBVS | Xidian  | 基线 BMS  | WHU IIP | 本文算法    |
|--------------|---------|---------|-------------|---------|---------|---------|---------|---------|
| Similarity   | 0.547 4 | 0.425 1 | 0.519 9     | 0.479 8 | 0.571 3 | 0.502 6 | 0.559 3 | 0.702 5 |
| CC           | 0.569 9 | 0.372 8 | 0.519 9     | 0.508 7 | 0.616 7 | 0.346 5 | 0.626 3 | 0.805 1 |
| AUC Shuffled | 0.625 9 | 0.602 4 | 0.628 3     | 0.620 8 | 0.648 4 | 0.588 5 | 0.630 7 | 0.673 4 |
| AUC Borji    | 0.792 1 | 0.726 2 | 0.758 2     | 0.791 3 | 0.794 9 | 0.656 0 | 0.796 0 | 0.851 4 |
| AUC Judd     | 0.813 3 | 0.748 9 | 0.784 6     | 0.811 5 | 0.820 7 | 0.691 4 | 0.819 7 | 0.893 9 |

表 3 SALICON 测试集上的检测结果比较

| 系数           | WHU IIP | 基线 GBVS | Xidian  | 基线 BMS  | 改进 Rare 2012 | 基线 Itti | 本文算法    |
|--------------|---------|---------|---------|---------|--------------|---------|---------|
| Similarity   | 0.490 8 | 0.446 0 | 0.461 7 | 0.454 2 | 0.501 7      | 0.377 7 | 0.683 3 |
| CC           | 0.4569  | 0.421 2 | 0.481 1 | 0.426 8 | 0.510 8      | 0.204 6 | 0.823 0 |
| AUC Shuffled | 0.606 4 | 0.630 3 | 0.680 9 | 0.693 5 | 0.664 4      | 0.610 1 | 0.665 0 |
| AUC Borji    | 0.775 9 | 0.781 6 | 0.799 0 | 0.769 9 | 0.804 7      | 0.660 3 | 0.846 3 |
| AUC Judd     | 0.792 3 | 0.789 9 | 0.805 1 | 0.789 9 | 0.814 8      | 0.666 9 | 0.869 3 |

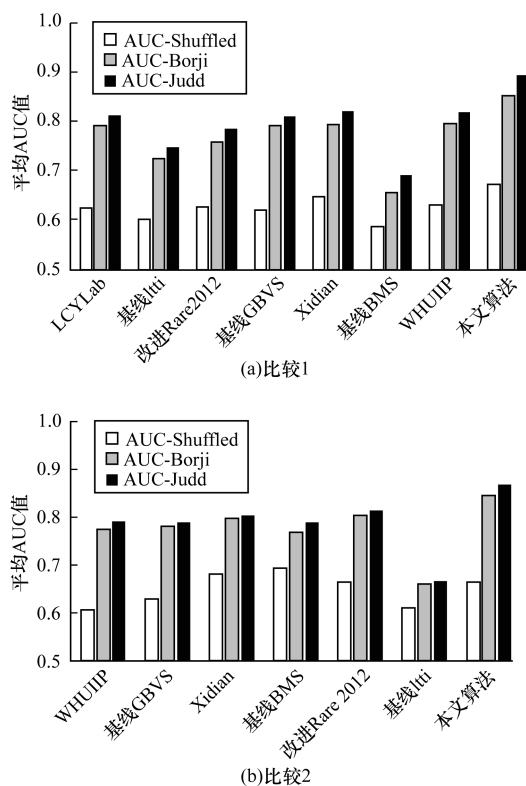


图 4 7 种显著性检测算法的平均 AUC 比较

可视化效果如图 5 所示,在每组中,左边第 1 列为原始图像,第 2 列为 Ground truth,第 3 列为本文算法预测的显著图。可以看到,本文算法提供更高的空间分辨率,显著区域的聚焦点突显,也与视觉注意点吻合,满足人眼视觉的观赏性。此外,无论是在复杂场景(草地),还是较大显著目标区域(木屋)或较小显著目标区域(足球场),本文算法都能取得不错的效果,表明提出的模型具有良好的稳定性。然而,在效果图中有一些块状片,影响了部分可视化的愉悦度,产生该现象的原因是对于显著性与非显著区

域特征的提取不是很精确。总体来说,从定量化评价与整体视觉效果都一致证明本文算法的预测结果要优于其他算法。这也与本文任务相一致,将显著性检测转换为回归问题,利用深度学习设计端对端模型,进而提高检测准确率的能力。

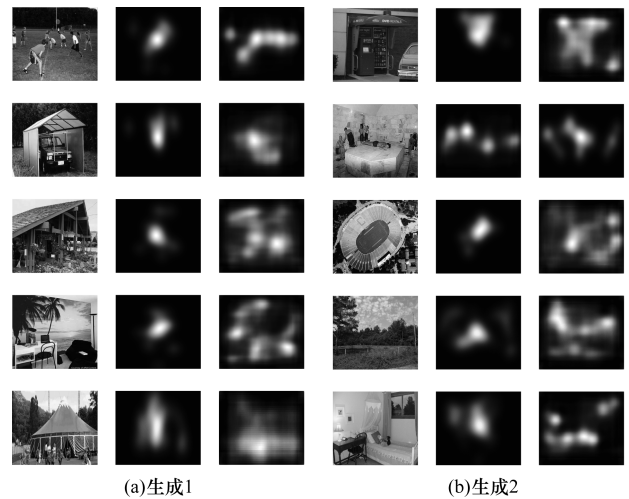


图 5 由本文算法在实验数据集上生成的显著图

## 4 结束语

本文提出一种组合卷积神经网络与空间变换神经网络进行显著性检测的新型端对端算法。首先对源图像进行预处理,接着对全局与局部特征采用不同的提取策略并进行融合。利用 CNN 深度捕捉上下文信息从而丰富图像的全局特征,用于突出显示显著区域内部的优势;将空间变换技术加入到 CNN 层以学习局部信息,有助于控制背景噪声。融合全局显著图与局部显著图后获得最终显著结果。实验结果表明,本文算法相较于其他基于图像的显著性检测模型在相识度及 AUC 上均有一定的提高,算法

在 iSUN 数据集上的平均 AUC 为 0.893 9,相似度系数为 0.702 5,相较于主流算法有一定优势,说明了空间卷积神经网络的重要性,以及融合局部与全局特征进行显著性检测的有效性。本文提出的空间卷积神经网络模型能够快速计算图像的显著值,实现了轻量级架构。但是模型参数有限,所以下一步工作是加深该模型的网络深度,进一步挖掘图像中的深层次信息。

### 参考文献

- [1] ZHANG Xilin, LI Zhaoping, ZHOU Tianguang, et al. Neural activities in v1 create a bottom-up saliency map [J]. *Neuron*, 2012, 73(1): 183-192.
- [2] ITTI L, KOCH C, NIEBUR E. A model of saliency-based visual attention for rapid scene analysis [J]. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 1998, 20(11): 1254-1259.
- [3] HAREL J, KOCH C, PERONAP. Graph-based visual saliency [C]//*Proceedings of the 20th Annual Conference on Neural Information Processing Systems*. Cambridge, USA: MIT Press, 2007: 545-552.
- [4] ACHANTA R, ESTRADA F, WILS P, et al. Saliency region detection and segmentation [C]//*Proceedings of the 6th International Conference on Computer Vision Systems*. Berlin, Germany: Springer, 2008: 66-75.
- [5] RAHTU E, KANNALA J, SALO M, et al. Segmenting salient objects from images and videos [C]//*Proceedings of European Conference on Computer Vision*. New York, USA: ACM Press, 2010: 366-379.
- [6] WEN Shifeng, HAN Junwei, ZHANG Dingwen, et al. Saliency detection based on feature learning using deep Boltzmann machines [C]//*Proceedings of IEEE International Conference on Multimedia and Expo*. Washington D. C., USA: IEEE Press, 2014: 1-6.
- [7] 易生, 梁华刚, 茹锋. 基于多列深度 3D 卷积神经网络的手势识别 [J]. *计算机工程*, 2017, 43(8): 243-248.
- [8] 王奎奎, 玉振明. 基于 DCT 零系数与局部结构张量的局部模糊检测 [J]. *计算机工程*, 2017, 43(6): 207-211.
- [9] 江帆, 刘辉, 王彬, 等. 基于 CNN-GRNN 模型的图像识别 [J]. *计算机工程*, 2017, 43(4): 257-262.
- [10] LI Guangbin, YU Yizhou. Visual saliency detection based on multiscale deep CNN features [J]. *IEEE Transactions on Image Processing*, 2016, 25(11): 5012-5024.
- [11] LI Hongyang, CHEN Jiang, LU Huchuan, et al. CNN for saliency detection with low-level feature integration [J]. *Neurocomputing*, 2017, 226(3): 212-220.
- [12] QU Liangqiong, HE Shengfeng, ZHANG Jiawei, et al. RGBD salient object detection via deep fusion [J]. *IEEE Transactions on Image Processing*, 2017, 26(5): 2274-2285.
- [13] 李岳云, 许悦雷, 马时平, 等. 深度卷积神经网络的显著性检测 [J]. *中国图象图形学报*, 2016, 21(1): 53-59.
- [14] KRUTHIVENTI S S, GUDISA V, DHOLAKIYA J H, et al. Saliency unified: a deep architecture for simultaneous eye fixation prediction and salient object segmentation [C]//*Proceedings of IEEE Conference on Computer Vision and Pattern Recognition*. Washington D. C., USA: IEEE Press, 2016: 5781-5790.
- [15] He Shengfeng, LAU R W H, LIU Wenxi, et al. SuperCNN: a superpixelwise convolutional neural network for salient object detection [J]. *International Journal of Computer Vision*, 2015, 115(3): 330-344.
- [16] JADERBERG M, SIMONYAN K, ZISSERMAN A, et al. Spatial transformer networks [C]//*Proceedings of the 28th International Conference on Neural Information Processing Systems*. Cambridge, USA: MIT Press, 2015: 2017-2025.
- [17] JIANG Ming, HUANG Shengsheng, DUAN Juanyong, et al. SALICON: saliency in context [C]//*Proceedings of IEEE Conference on Computer Vision and Pattern Recognition*. Washington D. C., USA: IEEE Press, 2015: 1072-1080.
- [18] XIAO Jianxiong, HAYS J, EHINGER K A, et al. SUN database: large-scale scene recognition from abbey to zoo [C]//*Proceedings of IEEE Conference on Computer Vision and Pattern Recognition*. Washington D. C., USA: IEEE Press, 2010: 3485-3492.
- [19] IBARZ J, SZEGEDY C, VANHOUCKE V, et al. Large-scale scene understanding (LSUN) [EB/OL]. [2017-04-20]. <http://lsun.cs.princeton.edu/leaderboard/#saliencyisun>.
- [20] RICHEL N, MANCAS M, DUVINAGE M, et al. Rare2012: A multi-scale rarity-based saliency detection with its comparative statistical analysis [J]. *Signal Processing: Image Communication*, 2013, 28(6): 642-658.
- [21] XIA Chen, QI Fei, SHI Guangming. Bottom-up visual saliency estimation with deep autoencoder-based sparse reconstruction [J]. *IEEE Transactions on Neural Networks and Learning Systems*, 2016, 27(6): 1227-1240.
- [22] ZHANG Jianming, SCLAROFF S. Saliency detection: a boolean map approach [C]//*Proceedings of IEEE International Conference on Computer Vision*. Washington D. C., USA: IEEE Press, 2013: 153-160.

编辑 顾逸斐