

复杂场景下融合深度信息的显著区域匹配

肖 锋,冯 飞,田鹏辉

(西安工业大学 计算机科学与工程学院,西安 710021)

摘 要: 针对复杂场景显著区域匹配过程中目标定位困难、冗余信息过多导致误匹配率高、匹配时间长的问题,提出一种融合深度信息的显著区域匹配算法。利用融合深度信息的视觉注意机制模型提取场景图像中的显著区域,得到场景目标的粗定位结果。使用基于局部特征点的匹配策略对有效目标在场景中的区域进行精确定位,并通过 FLANN 双向匹配实现对有效目标的匹配定位。实验结果表明,在光照强度和场景视角变化等复杂场景下,该方法能够有效减少匹配点,提高场景的匹配精度和匹配效率。

关键词: 复杂场景;深度信息;视觉注意机制;目标匹配;显著区域

中文引用格式: 肖 锋,冯 飞,田鹏辉. 复杂场景下融合深度信息的显著区域匹配[J]. 计算机工程,2017,43(5): 248-254,260.

英文引用格式: Xiao Feng, Feng Fei, Tian Penghui. Salient Region Matching Fused with Depth Information in Complex Scenes[J]. Computer Engineering, 2017, 43(5): 248-254, 260.

Salient Region Matching Fused with Depth Information in Complex Scenes

XIAO Feng, FENG Fei, TIAN Penghui

(School of Computer Science and Engineering, Xi'an Technological University, Xi'an 710021, China)

【Abstract】 In complex scenes, salient region matching has problems such as difficult targeting positioning and too much redundant information, resulting in high error matching rate and long matching time. Aiming at these problems, this paper proposes a salient region matching algorithm fused with depth information. First of all, adopting visual attention mechanism model fused with depth information, it extracts salient region of the scene image to get rough location of the scene target. Secondly, it pinpoints relevant target in the areas of the scene using matching strategy based on local feature points. At last, it uses FLANN bilateral matching method to match and locate relevant target. Experimental result shows that the proposed method can effectively reduce the number of match points while improving matching accuracy and matching efficiency, even in complex scenes with changed light intensity, changed view angle of the scene and other changes.

【Key words】 complex scenes; depth information; visual attention mechanism; target matching; salient region

DOI: 10.3969/j.issn.1000-3428.2017.05.040

0 概述

图像特征点匹配是计算机视觉研究中的热点,在视觉注意机制、机器人视觉导航以及遥感图像匹配等领域有着广泛的应用。特征点匹配主要是利用获取的场景中像素值计算得的特征点^[1-2]、特征线段^[3]与特征区域^[4]等实现复杂环境下的场景匹配^[5]。与直接利用像素值计算不同,特征点匹配的匹配点数相对较少,计算量较小,实时性相对较高,对匹配点位置变化敏感。其中,基于特征线段、区域

的匹配方法计算量较大且无法解决旋转、尺度不变性等问题,当场景目标发生比例改变、局部遮挡等变化时,就无法实现匹配,制约了方法的使用范围^[6]。

文献[7]提出的基于描述子 SIFT 的匹配方法能够很好地完成图像匹配的任务,而且还具有旋转、光照以及尺度不变性,但是存在匹配时间长、匹配正确率低的缺点。文献[8-9]在其基础上进行了改进,提高了匹配效率和匹配正确率。然而,它们都是在整幅图像中提取特征点,包括大量由图像背景产生的伪特征点,即冗余信息,会对图像的匹配产生干扰。

基金项目: 国家自然科学基金(61572392, 61671362)。

作者简介: 肖 锋(1976—),男,副教授、博士,主研方向为场景理解、智能信息处理;冯 飞,硕士研究生;田鹏辉,副教授、博士。

收稿日期: 2016-07-23 **修回日期:** 2016-09-27 **E-mail:** xffriends@163.com

实际上,真正有效的区域在图像中占很少一部分,这些有效的区域正是视觉注意的显著区域。因此,把2幅图像的匹配转化为2幅图像显著区域的匹配,可以很大程度上排除干扰信息,提高匹配效率和匹配正确率。文献[10-11]算法利用了这一思想,但在寻找有效目标时,它们找出的是目标物体所在的区域,不仅包含目标物体,还包含了一些冗余信息,因而会对匹配的精度和效率产生影响。同时,文献[10-11]中图像匹配时所用的SIFT算法是单方向的匹配,即只要在右视图找到能与左视图特征点匹配的,就认为它们是一对匹配点。这种匹配方法容易产生误配点,对匹配精度有一定的影响。

自底向上和自顶向下模型是2种基本视觉注意模型。其中,Itti模型是自底向上视觉注意模型中最具代表性的模型。但是,Itti模型在特征融合方面没有考虑与空间位置和空间关系相关的信息^[12]。

针对以上方法的局限性,同时借鉴心理学中有关视觉注意的研究成果,本文提出一种复杂场景下融合深度信息的显著区域匹配算法。首先通过融合深度信息的视觉注意模型提取场景显著区域,实现目标粗定位,得到含有冗余区域的显著图;然后采用基于特征点的匹配策略确定目标区域所在位置,即目标的精确定位,得到只含有目标区域的显著图;最后使用FLANN双向匹配方法对只含目标的显著图进行匹配。

1 复杂场景中目标图像定位及匹配

背景信息和目标信息是一幅图像必不可少的元素。目标信息一般是根据任务情况想要获取的区域,在图像中只占很少的部分;而背景信息则是用来混淆有效信息的获取,而一幅图像中,往往大部分都是背景信息。因此,在大部分的研究中,只需要研究图像的目标区域即可。图像中的目标区域具有很高的显著性,因此,将视觉注意模型应用到图像有效区域的搜索中,能够缩小在图像中的查询范围,增强实时性^[13]。

1.1 融合深度信息的视觉注意计算模型目标粗定位

随着视觉注意模型的发展,目前已形成多种有效的视觉注意模型,其中,最根本也最具代表性的是Itti模型^[14]。本文通过引入Itti模型来提取图像的有效目标。

1.1.1 Itti视觉注意模型

Itti视觉注意模型首先提取图像的初级特征,对输入图像 M 采样,获得其高斯金字塔 $M(\sigma)$,其中 σ 代表图像尺度。把图像每一层都划分为亮度、颜色和朝向3个通道,分别表示为 $I(\sigma)$, $C(\sigma, A)$ 和

$O(\sigma, \theta)$ 。其中, $C(\sigma, A)$ 是用高斯滤波器对 $M(\sigma)$ 滤波后得到的; $A = (R, G, B, Y)$,为4个颜色分量; $O(\sigma, \theta)$ 是用Gabor滤波器对 $M(\sigma)$ 进行滤波后获得的, $\theta = (0^\circ, 45^\circ, 90^\circ, 135^\circ)$,包含4个朝向。然后利用中央周边差算子对不同的通道计算特征图,中心尺度层用 $C = \{2, 3, 4\}$ 表示,周边尺度层用 $\delta \in \{3, 4\}$ 表示, $S = C + \delta, \delta \in \{3, 4\}$ ^[15]。对不同通道的各个金字塔层采用交叉尺度的减法 $\ominus$ 产生对应的特征图,如式(1)所示。

$$\begin{aligned} I(c, s) &= |I(c) \ominus I(s)| \\ O(c, s, \theta) &= |O(c, \theta) \ominus O(s, \theta)| \\ RG &= |(R(c) - G(c) \ominus R(s) - G(s))| \\ BY &= |(B(c) - Y(c) \ominus B(s) - Y(s))| \end{aligned} \quad (1)$$

特征图生成后,将获得的不同通道的特征图通过交叉尺度相加的方法分别合并成亮度 I 、颜色 C 和方向 O 3个特征显著图,如式(2)所示。

$$\begin{aligned} I &= \bigoplus_{c=2s=c+3}^4 \bigoplus_{c=2s=c+3}^{c+4} N(I(c, s)) \\ C &= \bigoplus_{c=2s=c+3}^4 \bigoplus_{c=2s=c+3}^{c+4} [N(RG(c, s)) + N(BY(c, s))] \\ O &= \sum_{\theta \in (0^\circ, 45^\circ, 90^\circ, 135^\circ)} \bigoplus_{c=2s=c+3}^4 \bigoplus_{c=2s=c+3}^{c+4} [N(O(c, s, \theta))] \end{aligned} \quad (2)$$

再对获得的3幅特征显著图进行归一化处理,如式(3)所示。

$$S = \frac{1}{3}(N(I) + N(C) + N(O)) \quad (3)$$

最后,采用“胜者为王”和“禁止返回”机制确定图像显著性目标的位置。图1是Itti模型生成显著图的过程。

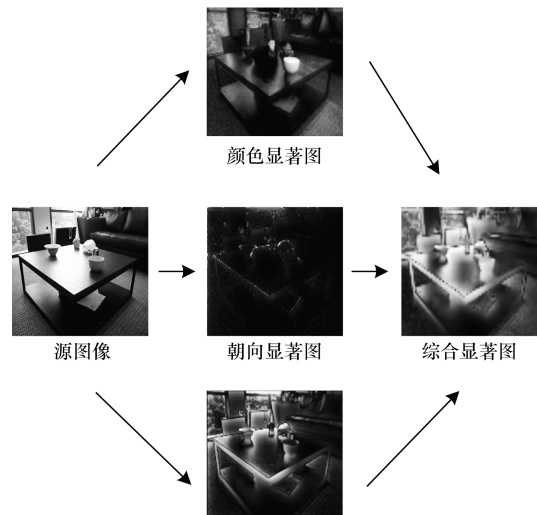


图1 Itti模型生成显著图的过程

1.1.2 深度信息提取

物体特征一般是由视觉刺激产生的,特征的提

取往往关系到显著区域的精确度。但是,若使用传统的 Itti 模型,没有考虑与空间位置和空间关系相关的信息,难以确保提取的目标区域的准确性。因此,本文在 Itti 模型的基础上通过融入深度信息来提高显著性检测的准确率^[16]。

本文引用文献[13]中基于图像分割的自适应立体匹配方法来获取图像深度信息,主要包括彩色图像分割、立体匹配和视差层 3 个步骤。

步骤 1 彩色图像分割。

本文采用均值平移算法来实现彩色图像分割,使密度模型不同的像素能够划分到相应的模型下。

步骤 2 立体匹配。

立体匹配大致分为对像素周围小区域进行约束的局部匹配和对扫描线甚至整幅图进行约束的全局匹配。将大于一定面积的分割区域用视差图描述,定义的视差图如式(4)所示。

$$d = c_1 x + c_2 y + c_3 \quad (4)$$

其中, c_1, c_2, c_3 为模板参数; x, y 为图像像素坐标; d 为像素 (x, y) 处的视差值。为更好地衡量视差图的精度,选用对应像素的平方和作为衡量依据,为此引入绝对强度误差和(Sum of Absolute Intensity Differences, SAD)和梯度(Gradient, GRAD)指标,如式(5)所示。

$$\begin{aligned} C_{\text{SAD}}(x, y, d) &= \sum_{(i, j) \in N(x, y)} |I_1(x, y) - I_2(i + d, j)| \\ C_{\text{GRAD}}(x, y, d) &= \sum_{(i, j) \in N_x(x, y)} |\nabla_x I_1(i, j) - \nabla_x I_2(i + d, j)| \\ &\quad + \sum_{(i, j) \in N_y(x, y)} |\nabla_y I_1(i, j) - \nabla_y I_2(i + d, j)| \end{aligned} \quad (5)$$

其中, $N(x, y)$ 为 (x, y) 处的小窗口,大小为 3×3 ,窗口 $N_x(x, y)$ 四周没有最右列,而窗口 $N_y(x, y)$ 四围没有最低行, ∇_x 表示 x 的梯度, ∇_y 表示 y 的梯度。设 ω 为优化权值,通过式(6)计算视差为 d 时点 (x, y) 的相似性匹配代价。

$$\begin{aligned} C(x, y, d) &= (1 - \omega) \times C_{\text{SAD}}(x, y, d) \\ &\quad + \omega \times C_{\text{GRAD}}(x, y, d) \end{aligned} \quad (6)$$

步骤 3 视差层构建。

视差层的目的是找到一个视差函数,使得全局能量达到最小。构造全局能量函数表示模板在 (x, y) 的模板视差,如式(7)所示。

$$E(f) = E_{\text{data}}(f) + E_{\text{smooth}}(f) \quad (7)$$

其中:

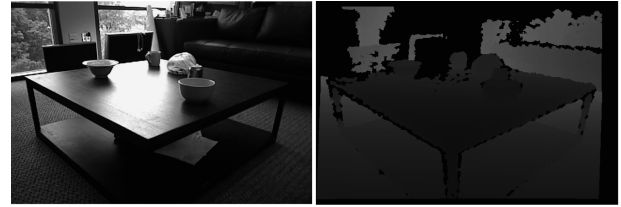
$$E_{\text{smooth}}(f) = \sum_{(S, S')} u_{s, s'} \times \delta(f(S \neq f(S'))))$$

$$E_{\text{data}}(f) = \sum_{S \in \mathbb{R}} C(S, f(S))$$

$E_{\text{data}}(f)$ 定义为 2 个分割之间的差异; $E_{\text{smooth}}(f)$

为视差平滑函数; $u_{s, s'}$ 为 2 个分割区域 S 和 S' 的边界长度。

最后通过计算最小切割面得到最优视差图,即深度信息图,如图 2 所示,其中,图 2(a)是源图像,图 2(b)是源图像对应的深度图像。



(a)源图像 (b)深度图像

图 2 利用本文算法生成的深度图像

1.1.3 融合深度信息的视觉注意模型

由于传统的 Itti 视觉注意模型没有考虑图像深度信息对人类视觉的影响,因此在对图像特征的选择上可能与人类的视觉选择机制存在很大的不同。为更好地模拟人类视觉注意机制,得到更加准确的显著区域,本文在 Itti 视觉注意模型的基础上,将深度信息与之融合,形成一种新的视觉注意模型,其结构如图 3 所示。

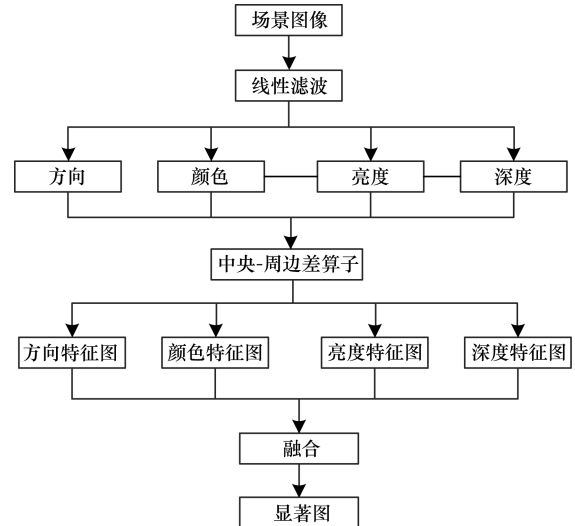


图 3 融合深度信息的视觉注意模型结构

该模型在原有 Itti 模型的基础上,在提取图像的颜色、亮度、朝向等特征时,同时提取出图像的深度图,然后通过中央-周边差算子分别得到各自的特征显著图,通过融合所有的特征显著图,得到图像显著图,该显著图就是含有冗余区域的目标显著图。图 4 所示是用不同模型得到的图像显著区域,其中,图 4(a)是用传统 Itti 模型得到的显著区域,图 4(b)是传统 Itti 模型融合图像深度信息后得到的显著区域。

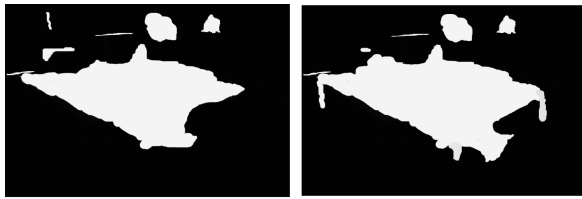


图4 不同视觉注意模型获得的显著区域

1.2 局部特征点匹配的目标精确定位与匹配

上文中,目标图像已经被检测出来。伴随大量冗余区域,若此时对目标图像进行匹配,冗余信息会对匹配精度和匹配效率产生干扰。为了提高匹配精度和效率,本文提出基于显著区域引导的特征点匹配策略,以实现目标区域的精确定位、消除冗余区域。

基于显著区域引导的特征点匹配策略首先对上文得到的显著区域进行特征点检测,然后利用参考图与获取的显著区域进行特征点匹配,最后根据匹配的特征点数确定目标所在区域。如图5所示,其中,图5(a)是参考图,是源图像中的最显著部分;图5(b)是场景图像已获取的显著区域。



图5 参考图与场景图像中的显著区域

该策略具体步骤如下:

步骤1 检测显著区域特征点。

经过显著区域提取出的目标图像过滤了一部分冗余区域,针对已获取出的场景图像的显著区域,用SURF算法提取特征点,以获取场景图像显著区域所存在的匹配点。

步骤2 对参考图与显著区域进行匹配。

在完成步骤1后,为了获取更稳定、鲁棒的关键点,利用最邻近距离比例的方法进一步对匹配点进行筛选。具体过程为:对场景的显著区域定义其需要匹配的关键点 P ,在显著区域中与 P 对应的最近及次近的SURF关键点为 P' 和 P'' ,如果满足式(8),则判断关键点对 (P, P') 得到了匹配。

$$D(P, P')/D(P, P'') \leq TRation \times Sa(x, y) \times \alpha \quad (8)$$

其中, $D(\cdot)$ 表示2个带匹配关键点SURF特征向量间的欧氏距离; $TRation$ 表示距离的比例,此比例决定了最邻近距离比例方法的效果,依据经验将此设置为0.7; $Sa(x, y)$ 表示有效目标中 (x, y) 位置的值,该位置由 $P'(x, y)$ 决定; α 为控制常数,将其设置为4,即可以满足大部分图像处理要求^[17]。

对参考图与场景图像的显著区域图,使用SURF算法进行匹配,匹配情况如图6所示。

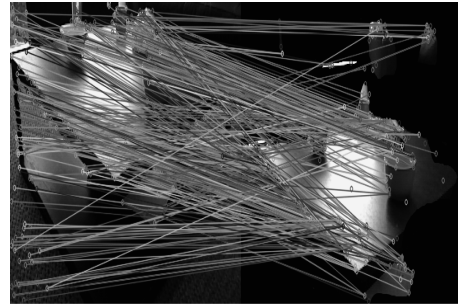


图6 参考图与显著区域的匹配情况

通过图6可以看出,虽然匹配图像会产生误匹配点,但是在有效目标中的匹配点个数仍然占多数,根据此特征,可以快速定位有效目标的区域。

步骤3 确定目标所在区域。

分别计算图6中不同区域的匹配点个数;然后遍历所有区域,直到第 n 个区域中的匹配点个数 R_n 满足式(9)和式(10)。

$$R_n \geq R_{all} \times T_N \quad (9)$$

$$R_n = \max_{R_n \in \mathbb{Z}} (R_i), i = 1, 2, \dots, N \quad (10)$$

其中, R_{all} 表示图中所有匹配点的个数; T_N 表示范围在 $[0, 1]$ 之间的比例常数; R_i 表示第 i 个区域内匹配点的个数; N 表示显著区域个数,在图5(b)中, $N=5$ 。最后把符合条件的区域提取出来,这部分区域就是有效目标所在的区域,通过上述步骤获得场景图像中有效目标所在区域,如图7所示。



图7 有效目标所在区域

步骤 4 对目标所在区域进行匹配。

传统匹配算法只是在单方向上寻找可以匹配的
点,势必会产生较多的误匹配点,影响匹配精度。文献
[18]通过实验对比,证明 FLANN 相比传统匹配算法具
有更好的匹配精度和实时性。因此,本文将提取出的有
效目标区域用 FLANN 双向匹配方法进行匹配。

2 实验结果与分析

为证明本文算法能够很好地实现复杂场景下的
目标匹配,首先通过改变场景图像的光照、视角来验
证本文算法的鲁棒性;然后通过与基于视觉注意机
制的其他算法进行比较,突出本文算法的可行性。
实验使用文献[10]和文献[11]提出的基于视觉显
著区域的 SIFT 匹配算法。

2.1 实验结果

通过场景图像在不同环境下的匹配实验验证本
文算法的可行性和鲁棒性。实验匹配效果如图 8 所
示,其中,图 8(a)是在场景图像在普通环境下的匹
配结果;图 8(b)是场景图像视角发生改变时的匹配
结果;图 8(c)是场景图像在光照发生变化时的匹配
结果,具体数据如表 1 所示。

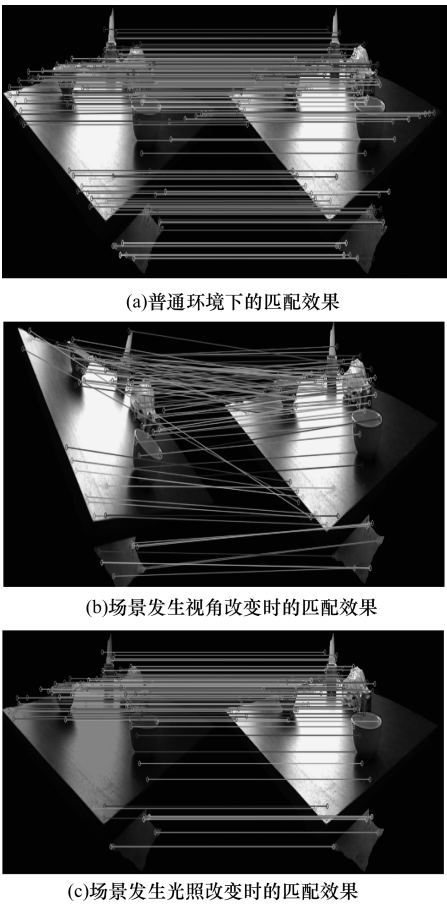


图 8 不同场景下本文算法的匹配效果

表 1 不同场景下本文算法的匹配数据

实验场景	匹配对数	匹配精度/%
普通场景	172	93
视角发生改变	121	82
光照发生改变	106	88

结合图 8 和表 1 可以看出,本文算法在普通场
景条件下具有较高的匹配对数(172),当场景依次发
生视角和光照改变时,图像的匹配对数减少(分别为
121 和 106);匹配精度在普通场景下为 93%,发生视
角和光照变化时,分别为 82% 和 88%。从总体分析
来看,本文算法能够很好地满足场景图像在不同环
境下的匹配要求,匹配精度平均在 88% 左右,具有良
好的可行性和鲁棒性。

2.2 算法分析与比较

上述实验通过在不同环境下的匹配验证了本文
算法的可行性,本实验主要通过与其他基于视觉显
著性的匹配算法进行对比,来验证本文算法的有效
性。用本文算法分别与文献[10]和文献[11]中的
方法进行对比,并通过改变光照、视角等场景图像所
处环境进行改变来体现本文算法的有效性。图 9 是
场景图像在普通环境下 3 种算法的匹配效果对比,
具体数据如表 2 所示。

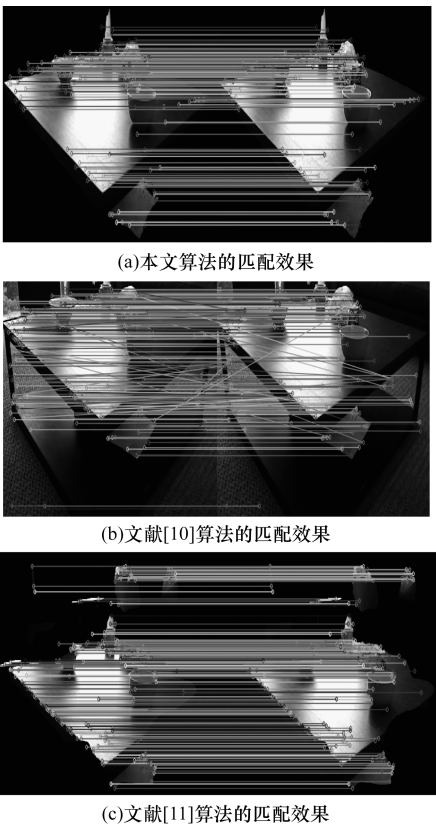


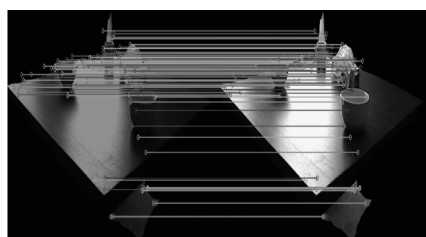
图 9 普通环境下 3 种算法的匹配效果

表 2 普通环境下 3 种算法的匹配数据

算法	匹配对数	匹配精度/%	匹配时间/s
本文算法	172	93	1.617
文献[10]算法	294	78	2.732
文献[11]算法	243	84	2.541

从表 2 中的数据可知,在普通场景下,本文算法的匹配对数为 172,相比文献[10]算法的 294 减少了 41%,相比文献[11]算法的 243 减少了 29%;在匹配精度上,本文算法为 93%,相比文献[10]算法的 78% 提升了 15%,相比文献[11]算法的 84% 提升了 9%;在匹配效率上,本文算法用时 1.617 s,相比文献[10]算法的 2.732 s 提升了 40%,相比文献[11]算法的 2.541 s 提升了 36%。

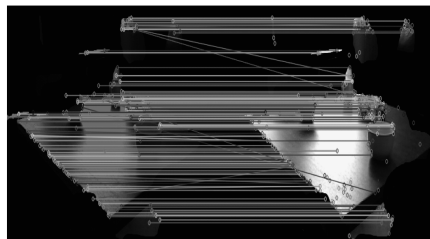
图 10 是场景图像在光照环境发生变化时 3 种算法的匹配效果对比,具体数据如表 3 所示。



(a)本文算法的匹配效果



(b)文献[10]算法的匹配效果



(c)文献[11]算法的匹配效果

图 10 光照改变环境下 3 种算法的匹配效果

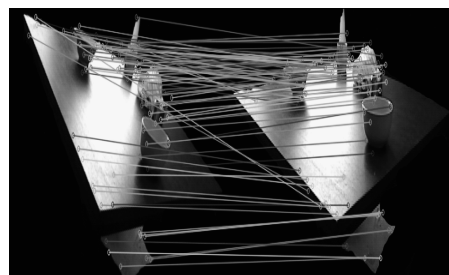
表 3 光照改变环境下 3 种算法的匹配数据

算法	匹配对数	匹配精度/%	匹配时间/s
本文算法	106	88	1.482
文献[10]算法	334	72	2.803
文献[11]算法	301	82	2.784

从表 3 中的数据可得,在场景光照发生改变的情况下,本文算法的匹配对数为 106,相比文献[10]算法的 334 减少了 68%,相比文献[11]算

法的 301 减少了 65%;在匹配精度上,本文算法为 88%,相比文献[10]算法的 72% 提升了 16%,相比文献[11]算法的 82% 提升了 6%;在匹配效率上,本文算法用时 1.482 s,相比文献[10]算法的 2.803 s 提升了 47%,相比文献[11]算法的 2.784 s 提升了 47%。

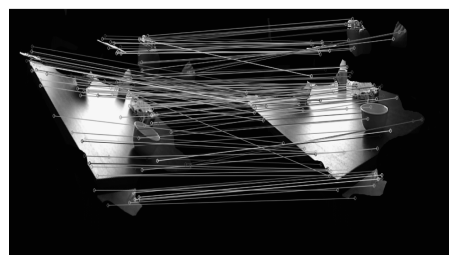
图 11 是场景图像在视角发生变化时 3 种算法的匹配效果对比,具体数据如表 4 所示。



(a)本文算法的匹配效果



(b)文献[10]算法的匹配效果



(c)文献[11]算法的匹配效果

图 11 视角发生变化时 3 种算法的匹配效果

表 4 视角发生变化时 3 种算法的匹配数据

算法	匹配对数	匹配精度/%	匹配时间/s
本文算法	121	82	1.424
文献[10]算法	237	64	2.657
文献[11]算法	251	73	2.762

从表 4 中的数据可知,在普通场景下,本文算法的匹配对数为 121,相比文献[10]算法的 237 减少了 49%,相比文献[11]算法的 251 减少了 52%;在匹配精度上,本文算法为 82%,相比文献[10]算法的 64% 提升了 18%,相比文献[11]算法的 73% 提升了 9%;在匹配效率上,本文算法用时 1.424 s,相比文

献[10]算法的 2.657s 提升了 46%, 相比文献[11]算法的 2.762 s 提升了 48%。

为更直观地对比 3 种算法的匹配性能, 对 3 种算法在不同环境下的匹配情况进行对比, 结果如图 12 所示。可以看出, 在相同的环境条件下, 本文算法具有更高的匹配精度和匹配效率, 同时, 其匹配精度持续保持在 87% 以上, 虽然匹配时间随着环境的改变而增加, 但是相比其他算法, 本文算法仍然具有较好的实时性。

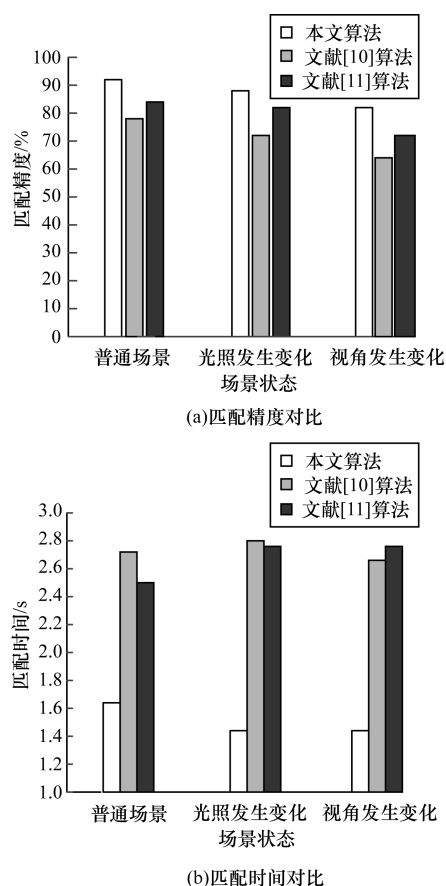


图 12 不同环境下 3 种算法的匹配性能对比

结合以上分析可知, 本文算法能够很好地完成场景匹配的任务, 当场景图像所在环境发生视角、光照等变化时, 仍然能够满足匹配的要求, 具有很强的鲁棒性; 同时, 本文算法相比文献[10]和文献[11]中的方法, 不仅匹配精度上得到了提高, 同时还提升了匹配效率。

通过引入时间复杂度和空间复杂度来研究本文算法的复杂度。把时间复杂度记作: $T(n) = O(f(n))$, $f(n)$ 是 $T(n)$ 的同数量级函数。本文算法是复杂场景下图像中特征点的匹配, 设图 5 参考图中特征点的个数为 n , 场景图像中的特征点个数为 m , 首先参

考图中每个特征点分别在场景图像中寻找对应的特征点, 即相当于参考图中每个特征点在场景图像中要遍历 m 次, 那么 n 个特征点需要遍历 $n \times m$ 次; 同理, 再用场景图像中的 m 个特征点依次去遍历其在参考图中所对应的特征点, 同样需要 $n \times m$ 次, 因此, 本文匹配算法的时间复杂度为 $T(n) = n \times m + m \times n$, 即时间复杂度为 $O(mn)$ 。本文算法最后得到 2 幅图像匹配点的交集, 个数小于任何一幅图的匹配点, 所占的内存空间也最少。

3 结束语

针对复杂场景下图像定位困难、冗余信息较多导致匹配效率和匹配精度低的问题, 本文提出一种融合深度信息的显著区域匹配算法。首先通过获取场景图像的深度信息, 融合 Itti 模型形成一种更有效的视觉注意模型, 提取出更准确、完整的目标区域, 进而实现对目标区域的粗定位; 然后根据有效目标区域匹配点较多的特性, 通过特征点匹配实现目标区域的精确定位; 最后利用 FLANN 双向匹配算法对目标区域分别在普通环境、改变光照强度和改变场景视角的环境下进行匹配。实验结果表明, 本文算法能够快速有效地完成场景在复杂环境下的匹配, 对图像的匹配精度平均保持在 88% 左右。在时间复杂度不变的情况下, 该算法能够减少空间复杂度, 提升匹配效率。

参考文献

- [1] Wu Yue, Ma Wenping, Gong Maoguo, et al. A Novel Point-matching Algorithm Based on Fast Sample Consensus for Image Registration[J]. IEEE Geoscience and Remote Sensing Letters, 2015, 12(1): 43-47.
- [2] Sarathi M P, Ansari M A. A Novel Method for Image Retrieval Based on Visually Significant Feature Point Maps[J]. International Journal of Computational Intelligence Systems, 2015, 8(2): 198-207.
- [3] Minster G, Moreau G, Saito H. Geolocation for Printed Maps Using Line Segment-based SIFT-like Feature Matching[C]//Proceedings of 2015 IEEE International Symposium on Mixed and Augmented Reality Workshops. Washington D. C., USA: IEEE Press, 2015: 88-93.
- [4] Pan Xunyu, Lü Siwei. Region Duplication Detection Using Image Feature Matching[J]. IEEE Transactions on Information Forensics and Security, 2010, 5(4): 857-867.
- [5] 李 博, 杨 丹, 张小洪. 基于 Harris 多尺度角点检测的图像匹配新算法[J]. 计算机工程与应用, 2006, 42(35): 37-40.

(下转第 260 页)

- [3] Wang Haijun, Gao Xinbo. Single-image Super-resolution Using Active-sampling Gaussian Process Regression [J]. IEEE Transactions on Image Processing, 2016, 25 (2): 935-948.
- [4] Zhang Lei, Wu Xiaolin. An Edge-guided Image Interpolation Algorithm via Directional Filtering and Data Fusion [J]. IEEE Transactions on Image Processing, 2006, 15 (8): 2226-2238.
- [5] Sun Jian, Xu Zongben, Shum H. Image Super-resolution Using Gradient Profile Prior [C]//Proceedings of 2008 IEEE Conference on Computer Vision and Pattern Recognition. Washington D. C., USA: IEEE Press, 2008: 11-18.
- [6] Dai Shengyang, Han Mei, Xu Wei, et al. Soft Edge Smoothness Prior for Alpha Channel Superresolution [C]//Proceedings of 2007 IEEE Conference on Computer Vision and Pattern Recognition. Washington D. C., USA: IEEE Press, 2007: 1-8.
- [7] Sun Jian, Zheng Nanning, Tao Hai, et al. Image Hallucination with Primal Sketch Prior [C]//Proceedings of 2010 IEEE Conference on Computer Vision and Pattern Recognition. Washington D. C., USA: IEEE Press, 2010: 729-736.
- [8] Takeda H, Farisu S, Milanfar P. Kernel Regression for Image Processing and Reconstruction [J]. IEEE Transactions on Image Processing, 2007, 16 (2): 349-366.
- [9] Chang Hong, Yeung D, Xiong Yimin. Super-resolution Through Neighbor Embedding [C]//Proceedings of 2004 IEEE Conference on Computer Vision and Pattern Recognition. Washington D. C., USA: IEEE Press, 2004: 275-282.
- [10] 徐忠强, 朱秀昌. 压缩视频超分辨率重建技术 [J]. 电子与信息学报, 2007, 29 (2): 499-505.
- [11] 韩玉兵, 陈小菁, 吴乐南. 一种视频序列的超分辨率重建算法 [J]. 电子学报, 2005, 33 (1): 126-130.
- [12] 杨浩, 安国成, 陈向东, 等. 一种基于实例的文本图像超分辨率重建算法 [J]. 东南大学学报 (自然科学版), 2008, 38 (2): 191-194.
- [13] Brooks A, Zhao Xiaonan, Pappas T. Structural Similarity Quality Metrics in a Coding Context: Exploring the Space of Realistic Distortions [J]. IEEE Transactions on Image Processing, 2008, 17 (8): 1261-1273.
- [14] 金郭赞. 基于 POCS 的压缩视频超分辨率重建 [D]. 南京: 南京邮电大学, 2008.
- [15] Park S, Park M, Kang M. Super Resolution Image Reconstruction: A Technical Overview [J]. IEEE Single Processing Magazine, 2003, 20 (3): 21-36.
- [16] Dong Weisheng, Zhang Lei, Shi Guangming. Centralized Sparse Representation for Image Restoration [C]//Proceedings of 2011 IEEE International Conference on Computer Vision. Washington D. C., USA: IEEE Press, 2011: 1259-1266.
- [17] Hardie RC, Barnard K J, Armstrong E E. Joint MAP Registration and High-resolution Image Estimation Using a Sequence of Undersampled Images [J]. IEEE Transactions on Image Processing, 1997, 6 (12): 1621-1633.

编辑 金胡考

(上接第 254 页)

- [6] Zitova B, Flusser J. Image Registration Methods: A Survey [J]. Image and Vision Computing, 2003, 21 (11): 977-1000.
- [7] 傅卫平, 秦川, 刘佳, 等. 基于 SIFT 算法的图像目标匹配与定位 [J]. 仪器仪表学报, 2011, 32 (1): 163-169.
- [8] 曾 雷, 王元钦, 谭久彬. 改进的 SIFT 特征提取和匹配算法 [J]. 光学精密工程, 2011, 19 (6): 1391-1397.
- [9] 罗 宇, 陈 勃, 李山山, 等. 基于空间约束 SIFT 的光学与 SAR 图像配准 [J]. 计算机工程, 2015, 41 (12): 182-187.
- [10] 李 颖, 郑 芳, 陆雪松, 等. 基于视觉注意模型的 SIFT 图像配准算法研究 [J]. 现代科学仪器, 2012 (2): 13-18.
- [11] 尹春霞, 徐 德, 李成荣, 等. 基于显著图的 SIFT 特征检测与匹配 [J]. 计算机工程, 2012, 38 (16): 189-191, 195.
- [12] 曾志宏, 李建洋, 郑汉垣. 融合深度信息的视觉注意计算模型 [J]. 计算机工程, 2010, 36 (20): 200-202.
- [13] 靳 薇, 张建奇, 张 翔. 基于视觉注意力模型的红外目标检测 [J]. 红外技术, 2007, 29 (12): 720-723.
- [14] Bellina C R, Bottigli U, Guzzardi R, et al. Interactive Transit Time Imaging (ITTI) for Improved Cardiac and Vascular Studies in Dynamic Scintigraphy [J]. The Journal of Nuclear Medicine and Allied Sciences, 1980, 24 (3/4): 201-207.
- [15] 冯俊丽. 基于 Itti 模型的计算机视觉注意模型研究 [J]. 科技风, 2011 (20): 118.
- [16] 刘 坤, 张晓烽, 陈宁纪, 等. 融合深度信息的雾霾情况下显著性目标提取 [J]. 河北工业大学学报, 2015, 44 (2): 10-15.
- [17] 郝文欣, 高立宁. 基于视觉注意机制与局部描述子的物体检测 [J]. 计算机工程与设计, 2012, 33 (5): 1918-1922.
- [18] 孔 军, 汤心溢, 蒋 敏. 多尺度特征提取的双目视觉匹配 [J]. 计算机工程与应用, 2010, 46 (33): 1-5.

编辑 金胡考