

基于光场图像序列的自适应权值块匹配深度估计算法^①



公冶佳楠, 李 轲

¹(中北大学 仪器与电子学院, 太原 030051)

²(中北大学 信息探测与处理山西省重点实验室, 太原 030051)

摘 要: 现有的深度估计算法中, 针对光场序列图像进行深度估计时, 在图像亮度变化较大和弱纹理区域, 其匹配效果较差, 鲁棒性较低. 针对这些问题, 本文提出了一种基于 CIELab 颜色空间的自适应权值块匹配算法. 由于彩色图像 RGB 颜色空间中颜色差异匹配影响因素较多, 本算法转换到 CIELab 空间进行颜色相似性匹配来计算权重值, 然后结合梯度和距离计算匹配图像和待匹配图像中匹配块得到综合权重值, 最后根据极平面图像 (EPI) 的线性特性对图像序列中匹配图像和待匹配图像块进行匹配计算, 求得深度图. 经过仿真验证, 本文算法能够较好的估计场景的深度信息, 精度上有较大的提升, 明显优于以往的深度估计算法, 可以广泛使用.

关键词: 光场图像序列; 极平面图像 (EPI); CIELab 颜色空间; 自适应权值块匹配; 深度估计

引用格式: 公冶佳楠, 李轲. 基于光场图像序列的自适应权值块匹配深度估计算法. 计算机系统应用, 2020, 29(4): 195–201. <http://www.c-s-a.org.cn/1003-3254/7387.html>

Adaptive Weight Block Matching Depth Estimation Algorithm Based on Light Field Image Sequence

GONGYE Jia-Nan, LI Ke

¹(School of Instrument and Electronics, North University of China, Taiyuan 030051, China)

²(Shanxi Provincial Key Laboratory of Information Detection and Processing, North University of China, Taiyuan 030051, China)

Abstract: In the existing depth estimation algorithm, when the depth estimation is performed on the image of the light field sequence, the matching effect is poor and the robustness is low when the image brightness changes and in the weak texture region. Aiming to solve these problems, this study proposes an adaptive weight block matching algorithm based on CIELab color space. Since the color difference matching in color image RGB color space has many influencing factors, the algorithm converts to CIELab space for color similarity matching to calculate the weight value, and then combines the gradient and distance to calculate the matching image and the matching block in the image to be matched to obtain the comprehensive weight value. Finally, according to the linear characteristics of the Epipolar Plane Image (EPI), the matching image and the image block to be matched in the image sequence are matched and calculated, and the depth map is obtained. After simulation, the proposed algorithm can estimate the depth information of the scene better, and the accuracy is greatly improved. It is obviously superior to the previous depth estimation algorithm and can be widely used.

Key words: light field image sequence; Epipolar Plane Image (EPI); CIELab color space; adaptive weight block matching; depth estimation

① 收稿时间: 2019-07-22; 修改时间: 2019-09-23, 2019-10-31; 采用时间: 2019-11-14; csa 在线出版时间: 2020-04-05

引言

光场是指空间中每一个点通过各个方向的光量,是包含了光的位置和方向信息的四维光辐射场的参数化表示^[1,2],通常用一组在空间中稠密采样的图像序列表示,称为光场图像。光场图像可以通过相机阵列^[3]、非结构化方法^[4]和孔径编码成像^[5]进行获取,相比于传统的2D成像方式,其多出了位置和方向信息2个自由度,这在计算成像中有很广的应用^[6]。基于光场图像信息的深度重建是指从获取的光场序列图像中提取有用的深度信息,且这些信息包含了丰富的位置和方向信息,所以在处理图像亮度变化较大和弱纹理区域可以得到很好的重建效果,尤其是在光场图像的深度估计中^[7]。

深度图像估计的准确性是深度重建的基础,在光场中,较多的光场视图数量对于获得准确和稠密深度图仍是一个很大的挑战。为了得到准确和可靠的深度图,许多具有代表性的方法已被采用。传统的立体匹配是一个很重要的研究方向^[8],通过两幅图或者是多幅图之间像素点的对应关系,获得场景的深度信息,但这些方法在处理弱纹理区和有遮挡区,效果很差,鲁棒性低^[9]。Bolles等人^[10]提出了极平面图像(ExtremePlane Image, EPI)的概念,根据EPI的线性特点,拟合直线后计算斜率达到估计深度的目的,得到了较好的深度图。Criminisi等^[11]采用分层思想,利用迭代优化方法将EPI分成不同深度的EPI-tube,通过对3D光场的反射特性进行分类,可以去除EPI-tube中由镜面反射造成的影响,得到较好的效果,但是其算法的复杂度较高,比较费时。Kim等^[12]提出由精到粗的深度扩散方法:首先根据边缘置信度,计算场景的边缘轮廓信息;再依次通过降采样来计算轮廓内部的深度,但是对于较小视差的估计,精确度比较低。丁伟利^[13]等提出了一种改进的Kim算法进行视差的估计:采用交叉检测模型检测边缘进行视差的估计,引入了权值的计算,但其算法没有解决较小视差估计不精确的问题。Wanner等^[14-16]提出了结构张量的方法获取EPI中的斜率进而求得视差,采用全局优化的方法将结果整合到深度图中,结果虽然较为平滑,但是边缘信息丢失严重。Tao等人^[17]提出了一种同时结合散焦和对应深度融合的方法估计深度,根据在不同焦线处的模糊程度,来得到其对应的深度,然而,对于离主透镜焦平面比较远的区域,深度估计误差较大。

针对上述问题,本文提出了一种改进的基于CIELab

颜色空间、梯度和距离的自适应权值块匹配深度估计方法。本文算法有3个贡献点:1)将块匹配应用到EPI域中,根据EPI的线性特点进行光场图像序列的块匹配。2)在匹配中,先将图像由RGB颜色空间转到CIELab颜色空间,计算EPI集中匹配块和待匹配块中的每个像素的对于中心像素的相似性。3)采用梯度和距离作为平滑项,计算块中每个像素对于中心像素的权值,结合CIELab颜色空间的权重值进行综合权重计算,之后计算参考图像和各视角图像中块的匹配成本值,比较不同斜率下的成本值大小,确定最小成本和对应的最优斜率,得到最佳的深度图。

1 光场获取及EPI深度估计

1.1 光场获取

本文采用相机沿直线运动的方式得到一系列的光场图像,如图1(a)。由于相机直线运动,因此这些图像对应的相机光心在同一直线上,假设为 L_s ,获取的光场图像的平面我们假设为 Π_{us} ,因此我们可以将获得的3D光场,如下表述:

$$\Pi_{us} \times L_s \rightarrow R^3, (u, v, s) \rightarrow E(u, v, s) \quad (1)$$

式中, $E(u, v, s)$ 表示空间中一点通过位置 (u, v, s) 的光线的亮度。这里用CIELab颜色空间中 L, a, b 三通道的值表示。

图1(a)是获取的光场图像序列,我们展示了一组光场图像中其中的4张图像。将得到的光场序列图像按顺序叠加起来,组成一个三维立体的合集,如图1(b),称为EPI集。其中中间的横向切片表示具有线性特点像素的集合,即为单个EPI。

图1(b)中,随着时间的推移,相机沿着箭头的方向移动,使目标物点 p 会在光场序列图像的不同位置出现,因此光场图像记录了目标物点 p 在不同视角下的信息。在形成的EPI中,目标物点 p 在不同视角下的成像点分别为 a, b, c 三点,其在EPI中为一条斜线 abc 。

图1(c)表示由多个目标物点形成的具有线性特点的EPI图。

1.2 光场深度估计

对于三维空间中的一点 $p(x, y, z)$,如图2(a)所示,假设它在图像序列中同一行 v^* 的投影分别为 $\{p_1, p_2, \dots, p_n\}$,相机对应的光心可以表示为 $\{c_1, c_2, \dots, c_n\}$ 其中相机光心之间的距离均是 Δs ,那么点 p 在相邻两幅图像之间均有相等的视差 $\Delta u = p_{i+1} - p_i$, p 点的深度 z 可通过三角测量原理求得,公式如下:

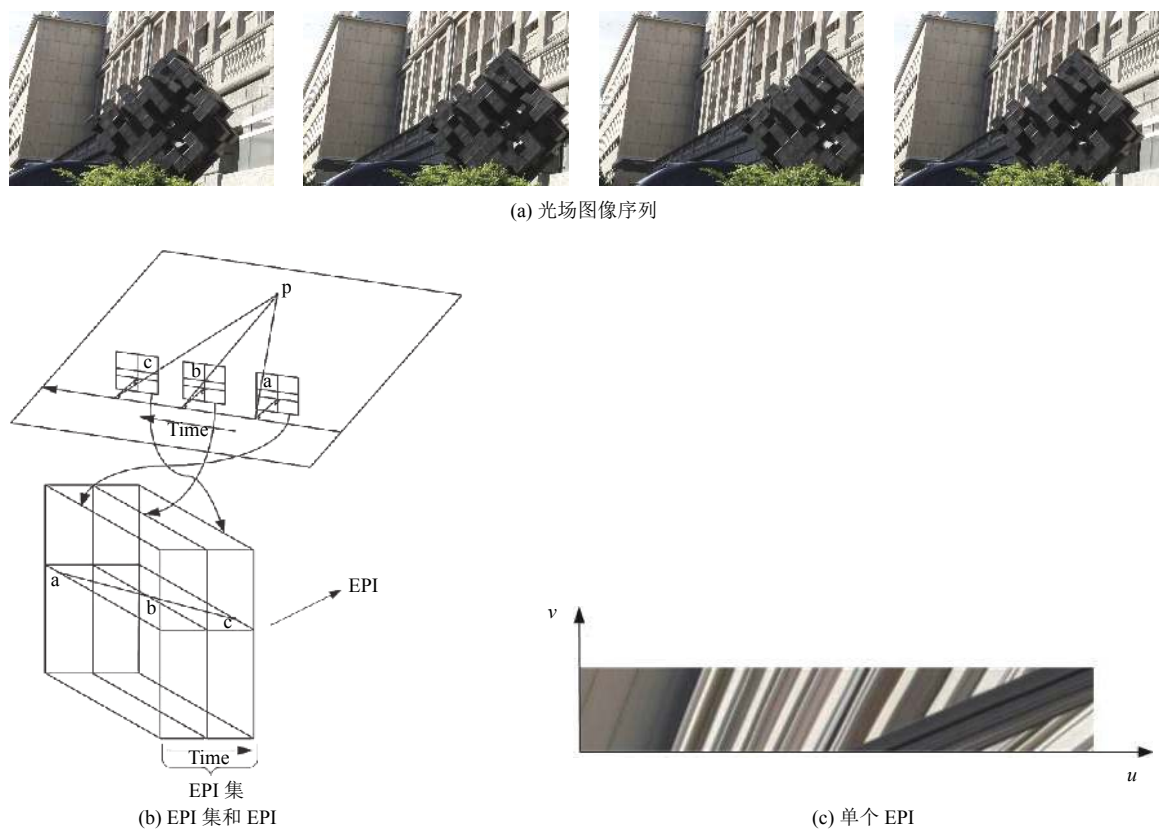


图1 EPI 形成

$$z = \frac{f}{\Delta u} \cdot \Delta s \tag{2}$$

式中, f 为相机的焦距. 由 p 点在各个相机中的投影 $\{p_1, p_2, \dots, p_n\}$ 组成了 EPI 上斜率为 $\frac{\Delta s}{\Delta u}$ 的直线 l_p , 如图 2(b), 直线的斜率与直线上 p 点的深度成正比关系. 因此, 我们计算场景点 p 的深度可以转换成求 EPI 中直线的斜率.

2 自适应权值算法

在斜率求解过程中, 受光照和噪声等因素的影响, 单个像素点之间匹配求斜率时误差较大, 本文采用基于自适应权值块匹配的方法来计算斜率.

匹配块是以各视角图像中待匹配像素点为中心像素的 5×5 像素块, 因此基于像素点求斜率可以转换成基于像素块匹配求斜率. 计算不同斜率下参考图像中的匹配块和各视角下目标图像中的待匹配块之间的成本值, 成本值最小时对应的斜率值最优.

2.1 CIE Lab 颜色空间

CIE Lab 颜色空间是由 CIE(国际照明委员) 于 1976 年制定的一种色彩模式, 它由亮度 L (Luminance)

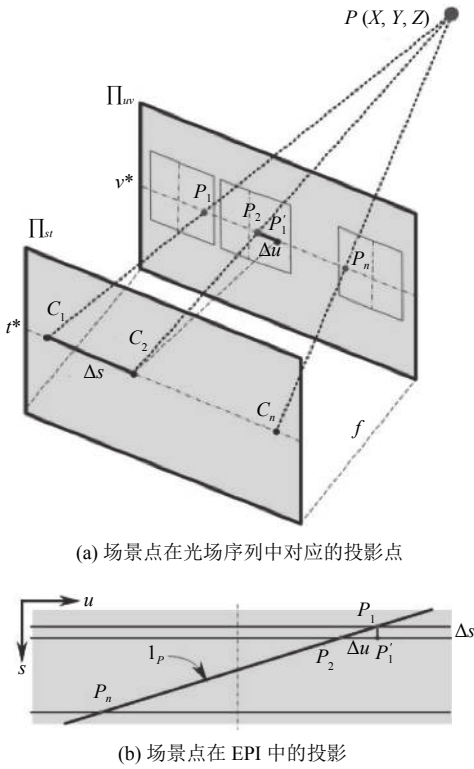


图2 三维空间中一点 p 在 EPI 中投影点的斜率

和色度信息 (a, b) 组成. CIELab 颜色空间采用坐标 Lab, 其中 a 的正向代表红色, 负向代表绿色, b 的正向代表黄色, 负向代表蓝色. CIELab 颜色空间对于色彩有较强的感知力, 其中 L 分量可以密切匹配亮度感知, 因此可以通过修改 a 和 b 的分量来感知颜色的相似程度, 由 RGB 颜色空间转换到 CIELab 颜色空间的公式如下:

$$X = 0.49 \times R + 0.31 \times G + 0.2 \times B \quad (3)$$

$$Y = 0.177 \times R + 0.812 \times G + 0.011 \times B \quad (4)$$

$$Z = 0.01 \times G + 0.99 \times B \quad (5)$$

由 XYZ 颜色空间转换到 Lab 颜色空间, 表示如下:

$$\begin{cases} L = 116f(Y) - 16 \\ a = 500 \left(f\left(\frac{X}{0.982}\right) - f(Y) \right) \\ b = 200 \left(f(Y) - f\left(\frac{Z}{1.183}\right) \right) \end{cases} \quad (6)$$

其中, $f(X) = 7.787X + 0.138$, $X \leq 0.008 856$, $f(X) = X^{\frac{1}{3}}$, $X > 0.008 856$. 式中由于 CIELab 颜色空间不包含人类感知的所有颜色, 而 XYZ 颜色空间几乎包含所有感知的颜色, 所有先将 RGB 颜色空间转换到 XYZ 颜色空间, 如式 (3) 至式 (5) 所示, 再转换到 CIELab 颜色空间, 如式 (6) 所示.

2.2 基于 CIELab 颜色空间自适应权值算法

在块匹配中, 匹配窗中像素的权值由 3 个因素决定: 梯度、CIELab 颜色空间中颜色差异和距离. 匹配窗中的某一像素离中心像素越近, 颜色差异越小, 梯度差异越小, 则该像素的权值越大. 设参考图像中匹配窗的中心像素为 P_0 , 窗内任一像素为 P , 如图 3 所示, 则像素 P 的权值可以表示为:

$$w(\vec{p}_0, \vec{p}) = f(\Delta c_{\vec{p}_0 \vec{p}}, \Delta g_{\vec{p}_0 \vec{p}}, \Delta gra_{\vec{p}_0 \vec{p}}) \quad (7)$$

式中, $\Delta c_{\vec{p}_0 \vec{p}}$ 表示像素 $\vec{p}_0$ 和像素 $\vec{p}$ 在颜色空间 CIELab 中颜色的差异, $\Delta g_{\vec{p}_0 \vec{p}}$ 表示像素 $\vec{p}_0$ 和像素 $\vec{p}$ 之间的距离, $\Delta gra_{\vec{p}_0 \vec{p}}$ 表示 $\vec{p}_0$ 和像素 $\vec{p}$ 之间的梯度差异. 由于这 3 个因素之间是相互独立的, 则匹配窗中任一个像素 P ,

相对于中心像素 p_0 的权值可表示为:

$$w(\vec{p}_0, \vec{p}) = f_s(\Delta c_{\vec{p}_0 \vec{p}}) \times f_p(\Delta g_{\vec{p}_0 \vec{p}}) \times f_g(\Delta gra_{\vec{p}_0 \vec{p}}) \quad (8)$$

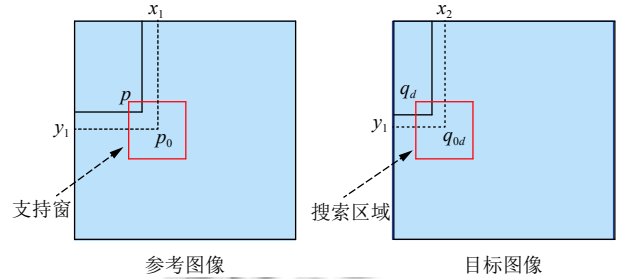


图3 自适应权值算法示意图

式中, $f_s(\Delta c_{\vec{p}_0 \vec{p}})$ 为颜色差异函数, $f_p(\Delta g_{\vec{p}_0 \vec{p}})$ 为距离函数, $f_g(\Delta gra_{\vec{p}_0 \vec{p}})$ 为梯度差异函数, 可以看出, 函数对于计算 $w(\vec{p}_0, \vec{p})$ 有着很重要的影响, 同时对匹配的精度影响较大, 函数的具体表示如下:

$$f_s(\Delta c_{\vec{p}_0 \vec{p}}) = \exp\left(-\frac{\Delta c_{\vec{p}_0 \vec{p}}}{r_s}\right) \quad (9)$$

$$f_p(\Delta g_{\vec{p}_0 \vec{p}}) = \exp\left(-\frac{\Delta g_{\vec{p}_0 \vec{p}}}{r_p}\right) \quad (10)$$

$$f_g(\Delta gra_{\vec{p}_0 \vec{p}}) = \exp\left(-\frac{\Delta gra_{\vec{p}_0 \vec{p}}}{\tau_{gra}}\right) \quad (11)$$

式中, r_s , r_p 和 τ_{gra} 为常数; 则 $\Delta c_{\vec{p}_0 \vec{p}}$, $\Delta g_{\vec{p}_0 \vec{p}}$, $\Delta gra_{\vec{p}_0 \vec{p}}$ 分别为像素 $\vec{p}$ 和中心像素 $\vec{p}_0$ 之间的 CIELab 颜色, 距离和梯度差异, 可如下表示:

$$\Delta c_{\vec{p}_0 \vec{p}} = \|\vec{p}_0 - \vec{p}\|_2 = \sqrt{(\vec{p}_{0L} - \vec{p}_L)^2 + (\vec{p}_{0a} - \vec{p}_a)^2 + (\vec{p}_{0b} - \vec{p}_b)^2} \quad (12)$$

$$\Delta g_{\vec{p}_0 \vec{p}} = \|\vec{g}_{\vec{p}_0} - \vec{g}_{\vec{p}}\|_2 = \sqrt{(x_{\vec{p}_0} - x_{\vec{p}})^2 + (y_{\vec{p}_0} - y_{\vec{p}})^2} \quad (13)$$

$$\Delta gra_{\vec{p}_0 \vec{p}} = \Delta grax_{\vec{p}_0 \vec{p}} + \Delta gray_{\vec{p}_0 \vec{p}} \quad (14)$$

其中, $\Delta grax_{\vec{p}_0 \vec{p}}$ 为匹配窗中像素和像 $\vec{p}_0$ 素 $\vec{p}$ 的水平梯度差异, $\Delta gray_{\vec{p}_0 \vec{p}}$ 为匹配窗中像素 $\vec{p}_0$ 和像素 $\vec{p}$ 的垂直梯度差异, 如下所示:

$$\Delta grax_{\vec{p}_0 \vec{p}} = \|\text{grax}_{\vec{p}_0} - \text{grax}_{\vec{p}}\|_2 = \sqrt{(\text{grax}_{\vec{p}_0 r} - \text{grax}_{\vec{p} r})^2 + (\text{grax}_{\vec{p}_0 g} - \text{grax}_{\vec{p} g})^2 + (\text{grax}_{\vec{p}_0 b} - \text{grax}_{\vec{p} b})^2} \quad (15)$$

$$\Delta gray_{\vec{p}_0 \vec{p}} = \|\text{gray}_{\vec{p}_0} - \text{gray}_{\vec{p}}\|_2 = \sqrt{(\text{gray}_{\vec{p}_0 r} - \text{gray}_{\vec{p} r})^2 + (\text{gray}_{\vec{p}_0 g} - \text{gray}_{\vec{p} g})^2 + (\text{gray}_{\vec{p}_0 b} - \text{gray}_{\vec{p} b})^2} \quad (16)$$

综上所述, $w(\vec{p}_0, \vec{p})$ 可如下表示:

$$w(\vec{p}_0, \vec{p}) = \exp\left(-\left(\frac{\Delta c_{\vec{p}_0 \vec{p}}}{r_s} + \frac{\Delta g_{\vec{p}_0 \vec{p}}}{r_p} + \frac{\Delta gra_{\vec{p}_0 \vec{p}}}{\tau_{gra}}\right)\right) \quad (17)$$

设像素 q_{od} 为目标图像上与参考图像中像素 p_0 可能的匹配像素, 像素 q_d 为参考图像中 q_{od} 对应的匹配窗中的像素, 像素 p 为参考图像中 p_0 对应的匹配窗的像素, 则 p_0 和 q_{od} 之间的差异表示如下:

$$E(p_0, p_{0d}) = \frac{\sum_{p \in N_{p_0} q_d \in N_{q_{od}}} w(p_0, p) w(q_{0d}, q_d) e(p, q_d)}{\sum_{p \in N_{p_0} q_d \in N_{q_{od}}} w(p_0, p) w(q_{0d}, q_d)} \quad (18)$$

式中, $e(p, q_d) = \min\left(\sum_{c \in \{r, g, b\}} |I_c(p) - I_c(q_d)|, T\right)$, I_c 为 R 、 G 、 B 的颜色值, T 代表截断阈值。

综上所述, 最终的视差值可以通过 WTA 算法得到:

$$d_p = \arg \min_{d \in D} E(p_0, q_{0d}) \quad (19)$$

其中, $D = \{d_{\min}, \dots, d_{\max}\}$ 为所有可能视差的集合, 即不同斜率的集合。

3 实验结果与分析

本文算法基于文献[12]提供的数据集和斯坦福数

据集进行测试. 为了验证本文算法的性能, 本文结果同时与文献[12,16]提供的流行深度估计算法进行比较, 从定性和定量两个方面来分析算法的估计结果. 在 Windows 7 操作系统下, Intel Core(TM)i7-2600 2.6 GHz CPU 以及 Matlab R2015b 的仿真软件下进行验证.

在定量分析中, 采用均方根误差 (RMSE) 和相对深度误差 (B) 作为量化指标评价算法性能, 其中 RMSE 和 B 值越小, 表示深度估计结果越好.

$$RMSE = \sqrt{\frac{1}{MN} \sum_{(x,y)} (d(p_{(x,y)}) - d_{GT}(p_{(x,y)}))^2} \quad (20)$$

$$B = \frac{1}{MN} \sum_{(x,y)} (|d(p_{(x,y)}) - d_{GT}(p_{(x,y)})| > \delta_d) \quad (21)$$

其中, M, N 表示图像的宽和高, $d(p_{(x,y)})$ 表示实验获取的深度估计值, $d_{GT}(p_{(x,y)})$ 表示深度图真值. δ_d 表示相对深度允许的误差, 本文实验中取值为 0.3.

3.1 定性分析

图 4 展示了文献[12]、文献[16]与本文方法深度估计结果, 可以发现, 本文算法在很好的保留图像边缘信息的同时对于平滑区域也保留了更多的细节信息, 如图 4(d) 中黑色边框所示, 较好的展示了娃娃脖子区域的细节, 而文献[12]的方法虽然娃娃边缘获得了较好的效果, 但是对于娃娃脖子区域的细节没有展现. 文献[16]方法, 在娃娃边缘上表现的结果较差.

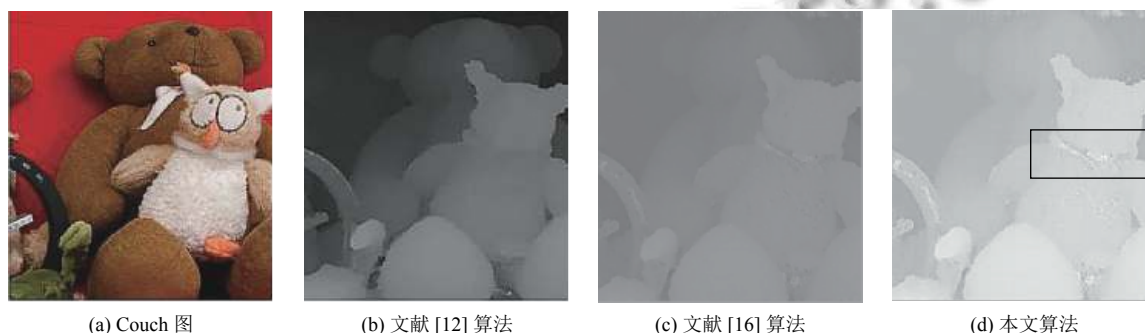


图 4 深度估计实验结果图

图 5 是各种深度估计方法求得的深度图, 通过对比可以发现, 本文算法和文献[12]算法一样都可以得到较好的深度图, 但是对于一些微小细节, 如图 5(d) 中白色边框中的教堂塔尖, 本文算法能够很好的展现. 文本算法在房子和前排植物等平滑区域展现了较好的灰度平缓变化, 即深度的平缓变化过程, 如图 5(d) 中黑色边

框所示, 而文献[16]算法表现的结果较差.

图 6 中, 本文算法与文献[12,16]算法均实现了较好的边缘深度估计, 但在平滑区域本文算法较文献[12,16]算法能够显示更多的细节信息, 如图 6(d) 中白色边框所示, 通过灰度值的变化显示物体不同深度的变化, 而文献[12,16]算法表现的结果图较差.

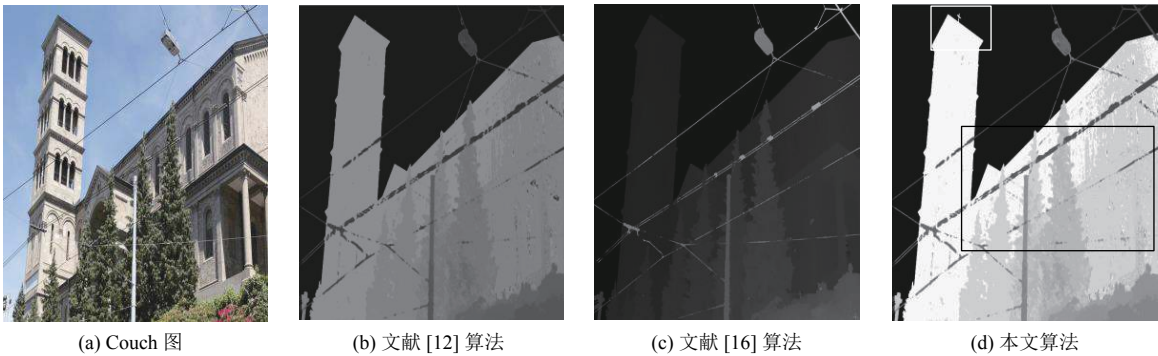


图5 深度估计实验结果图

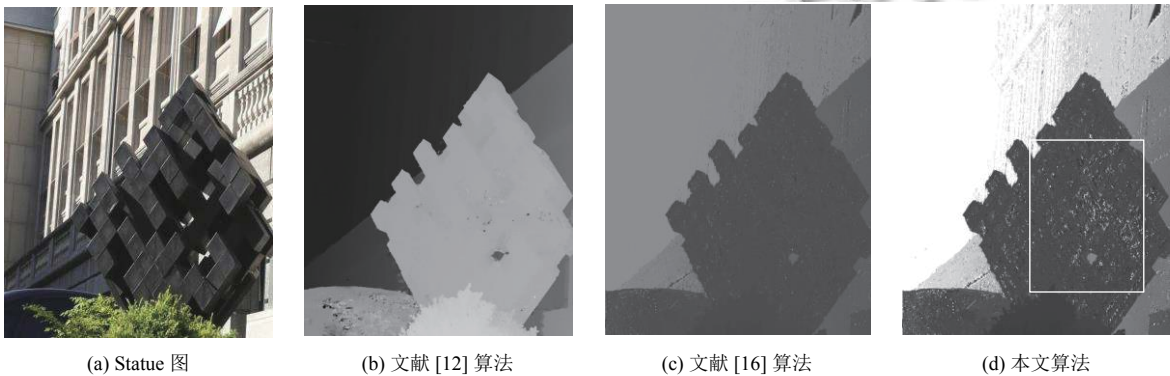


图6 深度估计实验结果图

3.2 定量分析

根据给定的数据集深度图真值, 可以进行定量分析. 表1给出了文献[12]、文献[16]以及本文算法的深度估计结果评价指标. 可以看出, 本文方法和文献[12]方法在边缘保存效果明显优于文献[16], 所以均方根误差与相对深度误差值较小. 本文方法在一些微小表现细节方面比文献[12]方法更具有一定的优势, 因此实验数值更小, 与本文中定性评价结果一致.

表1 各方法定量指标

图像	文献[12]方法		文献[16]方法		本文方法	
	RMSE	B	RMSE	B	RMSE	B
Couch 图	2.344	0.198	3.169	0.335	2.174	0.179
Courch 图	3.861	0.265	4.365	0.331	3.352	0.208
Statue 图	5.215	0.305	7.195	0.462	3.329	0.271

另外, 本文算法对于图像平滑区域的细节信息和边缘信息有较好的保持, 获得了较好的深度图. 同时由于本算法支持基于图像中多个匹配块的并行化计算, 算法运行时间大大缩短, 对比文献[12,16]的方法, 如表2所示, 本文算法的时间复杂度明显降低, 更适合于快速深度估计.

表2 算法运行时间 (单位: s)

算法图像	Couch 图	Courch 图	Statue 图
文献[12]	1951	2303	1712
文献[16]	2371	2716	2026
本文算法	793	913	676

4 结论

相比传统的深度估计算法, 本文提出了一种基于 CIELab 颜色空间的自适应权值块匹配算法. 本算法是在 EPI 上利用其线性特点进行匹配, 通过线性匹配求得最优斜率并确定最佳深度. 在线性匹配的过程中, 应用基于 CIELab 颜色空间的自适应权值算法, 求得匹配窗和待匹配窗的权值, 进而通过 WTA 算法, 确定最优斜率, 并求得深度. 通过对比深度图可以发现, 本文算法不仅能够有效的保留边缘信息, 同时对于内部的平滑区域, 也很好的展现了细节信息. 但对于图中的深度平滑性, 本文算法仍还有所不足. 因此, 在之后的工作中, 我们应该更加关注平滑性, 将深度表现的更加平滑, 得到更完善的深度信息.

参考文献

- 1 Gortler SJ, Grzeszczuk R, Szeliski R, *et al.* The lumigraph. Proceedings of the 23rd Annual Conference on Computer Graphics and Interactive Techniques. New York, NY, USA. 1996. 43–54.
- 2 Levoy M, Hanrahan P. Light field rendering. Proceedings of the 23rd Annual Conference on Computer Graphics and Interactive Techniques. New York, NY, USA. 1996. 31–42.
- 3 Wilburn B, Joshi N, Vaish V, *et al.* High performance imaging using large camera arrays. ACM Transactions on Graphics (TOG), 2005, 24(3): 765–776. [doi: [10.1145/1073204.1073259](https://doi.org/10.1145/1073204.1073259)]
- 4 Davis A, Levoy M, Durand F. Unstructured light fields. Computer Graphics Forum, 2012, 31(2pt1): 305–314. [doi: [10.1111/j.1467-8659.2012.03009.x](https://doi.org/10.1111/j.1467-8659.2012.03009.x)]
- 5 Veeraraghavan A, Raskar R, Agrawal A, *et al.* Dappled photography: Mask enhanced cameras for heterodyned light fields and coded aperture refocusing. ACM Transactions on Graphics, 2007, 26(3): 69. [doi: [10.1145/1276377.1276463](https://doi.org/10.1145/1276377.1276463)]
- 6 Zhang YB, Lv HJ, Liu YB, *et al.* Light-field depth estimation via epipolar plane image analysis and locally linear embedding. IEEE Transactions on Circuits and Systems for Video Technology, 2017, 27(4): 739–747. [doi: [10.1109/TCSVT.2016.2555778](https://doi.org/10.1109/TCSVT.2016.2555778)]
- 7 聂云峰, 相里斌, 周志良. 光场成像技术进展. 中国科学院研究生院学报, 2011, 28(5): 563–572.
- 8 丁伟利, 陈瑜, 马鹏程, 等. 基于阵列图像的自适应光场三维重建算法研究. 仪器仪表学报, 2016, 37(9): 2156–2165. [doi: [10.3969/j.issn.0254-3087.2016.09.030](https://doi.org/10.3969/j.issn.0254-3087.2016.09.030)]
- 9 郑淼, 赵红颖, 杨鹏, 等. 基于光场数据的无纹理景物视差估计方法. 北京大学学报 (自然科学版), 2018, 54(2): 336–340.
- 10 Bolles RC, Baker HH, Marimont DH. Epipolar-plane image analysis: An approach to determining structure from motion. International Journal of Computer Vision, 1987, 1(1): 7–55. [doi: [10.1007/BF00128525](https://doi.org/10.1007/BF00128525)]
- 11 Criminisi A, Kang SB, Swaminathan R, *et al.* Extracting layers and analyzing their specular properties using epipolar-plane-image analysis. Computer Vision and Image Understanding, 2005, 97(1): 51–85. [doi: [10.1016/j.cviu.2004.06.001](https://doi.org/10.1016/j.cviu.2004.06.001)]
- 12 Kim C, Zimmer H, Pritch Y, *et al.* Scene reconstruction from high spatio-angular resolution light fields. ACM Transactions on Graphics, 2013, 32(4): 73.
- 13 丁伟利, 马鹏程, 陆鸣, 等. 基于先验似然的高分辨光场图像深度重建算法研究. 光学学报, 2015, 35(7): 0715002.
- 14 Wanner S, Goldluecke B. Variational light field analysis for disparity estimation and super-resolution. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2014, 36(3): 606–619. [doi: [10.1109/TPAMI.2013.147](https://doi.org/10.1109/TPAMI.2013.147)]
- 15 Wanner S, Goldluecke B. Globally consistent depth labeling of 4D light fields. Proceedings of 2012 Computer Vision and Pattern Recognition. Providence, RI, USA. 2012. 41–48.
- 16 Wanner S, Goldluecke B. Spatial and angular variational super-resolution of 4D light fields. Proceedings of the 12th European Conference on Computer Vision-ECCV 2012. Florence, Italy. 2012. 608–621.
- 17 Tao MW, Hadap S, Malik J, *et al.* Depth from combining defocus and correspondence using light-field cameras. Proceedings of 2013 IEEE International Conference on Computer Vision (ICCV). Sydney, NSW, Australia. 2013. 673–680.