

DOI: 10.19650/j.cnki.cjsi.J2311797

基于视觉修正的激光雷达体积测量方法^{*}

郭隆强¹, 何贇泽^{1,2}, 杜旭¹, 郭昱良³, 付玉轩¹

(1. 湖南大学电气与信息工程学院 长沙 410082; 2. 湖南大学深圳研究院 深圳 518000;
3. 中南大学自动化学院 长沙 410083)

摘要:针对激光雷达点云数据稀疏、扰动、存在噪声和其他方法难以迁移,实时性差等难题,面向“L”型小尺寸目标研究了一种基于视觉修正的激光雷达体积测量方法。该方法首先通过联合标定和时间戳最近邻匹配实现相机与激光雷达数据的对齐;然后经过目标检测算法获取图像中目标的信息,与此同时对点云数据执行地面分割得到地面点云与非地面点云,利用视觉投影和点云聚类实现目标点云的分割,使用KDtree找到目标点云附近的地面点云;最后,设计了一种三维框的拟合算法初步完成点云目标三维框的粗拟合,并建立视觉修正模型对于目标三维框进行细修正,从而实现目标体积的计算。实验结果表明,对于武器箱道具、医疗箱和油桶等“L”型物体,提出的算法在一定范围内,体积测量的平均相对误差小于4.44%、最大误差小于6.12%、最大重复性小于5.61%,并且基于视觉的修正模型大幅提高了算法的精度和稳定性,在嵌入式平台的处理1帧用时55 ms,能够实现实时高精度的体积测量,具有良好的工程应用前景。

关键词: 体积测量;点云聚类;目标检测;视觉修正

中图分类号: TP29 TH702 **文献标识码:** A **国家标准学科分类代码:** 420.40

Volumetric measurement of Lidar based on visual correction

Guo Longqiang¹, He Yunze^{1,2}, Du Xu¹, Guo Yuliang³, Fu Yuxuan¹

(1. College of Electrical and Information Engineering, Hunan University, Changsha 410082, China;
2. Shenzhen Research Institute, Hunan University, Shenzhen 518000, China; 3. School of Automation,
Central South University, Changsha 410083, China)

Abstract: To address challenges faced by most volume measurement methods, such as difficulties in transferring, poor real-time performance, and sparse, noisy, and perturbed Lidar point cloud data, this article presents a vision-based correction method for Lidar volume measurement for small ‘L’-shaped objects. The proposed method first aligns camera and Lidar data through joint calibration and time-stamped nearest-neighbor matching. Subsequently, it leverages target detection algorithms to extract information from images while simultaneously performing ground segmentation on point cloud data to distinguish ground and non-ground points. By employing visual projection and point cloud clustering, the method segments target point clouds and utilizes KDtree to identify ground points in proximity to the target point cloud. Finally, a 3D box fitting algorithm is proposed to provide initial rough estimation of the point cloud target’s 3D box. A visual correction model is established to refine the target’s 3D box, and enable accurate volume calculation. Experimental results show that for ‘L’-shaped objects like weapon crates, medical boxes, and barrels, the proposed algorithm achieves promising results within a certain range. The average relative error in volume measurement is less than 4.44%, with a maximum error below 6.12% and a maximum repeatability error of 5.61%. In addition, the integration of the visual correction model significantly enhances the algorithm’s accuracy and stability. The processing of frame on an embedded platform takes 55 ms, demonstrating the capability to achieve real-time, high-precision volume measurement. This method holds great promise for practical engineering applications.

Keywords: volume measurement; point cloud clustering; target detection; visual correction

收稿日期: 2023-08-14 Received Date: 2023-08-14

^{*} 基金项目: 国家自然科学基金(62101184)、湖南省科技创新领军人才(2023RC1039)、湖南省自然科学基金重大项目(2021JC0004)、广东省基础与应用基础研究基金海上风电联合基金(2022A1515240050)、湖南省重点研发计划(2022GK2012)项目资助

0 引 言

近年来,由于激光雷达分辨率高,能获取目标的位置和轮廓信息等优点,使得其在自动驾驶、无人车、军事和农业等领域有着广泛的应用^[1-4]。视觉传感器可以获取现实世界丰富的语义信息,通过视觉传感器与激光雷达信息的融合^[5-6],可以实现对周围环境目标的可靠感知,对于各个领域的智能化极具帮助意义。

体积测量在许多领域都是至关重要的,提供了关键的信息,帮助人们更好地理解 and 处理目标与现象。例如在工程应用中准确测量物体的体积可以确定合适的材料用量、容量和空间布局;在贸易和商业活动中,体积测量是计算成本和定价的基础;在军事领域,物体的体积可作为威胁程度的判断依据之一。体积测量一般分为接触式和非接触式两大类。接触式步骤繁杂且效率低下,对于物体甚至可能造成损伤;相比于接触式,非接触式在保证精度的同时,效率更高和范围更广,对于测量物品没有损伤。

对于非接触式体积测量方式,主流的方法有基于机器视觉、基于激光雷达与视觉与激光结合等。Yunardi 等^[7]通过水平和垂直安装的两个相机对于分拣系统中的规则的包裹进行体积测量,其通过像素的实际大小计算每条边的长度从计算真实体积,测量对象要求在固定位置且需要正对,测量精度容易受相机安装位置的影响。尹文庆等^[8]基于结构光视觉对于谷粒体积流量进行测量,利用线激光获取非平整的谷粒流截面轮廓和滑轮测量流速实现体积测量,构建的测量系统特殊,只适用于特定的场景。吴琪波^[9]基于双目视觉对于长方体物体进行测量,固定双目相机通过物体上表面的面积和其与地面的深度之差计算体积,容易受到光照的影响且双目相机需要平行于长方体的上表面。Suzuki 等^[10]基于立体视觉实现对椭圆形药品体积的测量,通过计算表面积和其离碟形表面的乘积估计体积,但容易产生异常值,需要人工去除。张立斌等^[11]通过 3 个上左右固定位姿的单线激光雷达相互配合对固定区域内行驶的匀速车辆扫描构建其三维轮廓并使用曲线矫正模型从而进行尺寸测量。崔峰^[12]基于深度学习点云分割从点云数据中获取散料堆区域后对其进行三角网格化积分完成散料体积的计算,算法运行速度较慢且点云分割网络难以落地应用。郑国庆等^[4]基于激光雷达、视觉传感器和 IMU 的融合,通过环绕物体一周完成激光雷达闭环建图获取物体稠密的点云轮廓,结合视觉对物体进行定位,从而进行切片完成体积测量,精度较高但闭环建图对环境与场地要求较高且实时性差。通过上述调研可知,基于视觉的体积测量方式需要多个摄像头或者特定类型相机获取物体的深

度,虽然相对低成本但测量物体与传感器往往需要保持固定的位姿且受环境的影响大;基于激光雷达的体积测量方式往往需要多个激光雷达协同或者通过与其他传感器融合建图的方式获得稠密的物体轮廓,以此来解决点云稀疏性的问题,并且测量目标体积较大,点云的扰动性可忽略,去噪方法难以迁移到小尺寸目标体积测量中。

体积测量对象可以分为规则和非规则物体,但现实世界大多数场景下关注体积的对象基本都是类规则物体。以本文的测量对象武器箱道具、医疗箱和油桶为例,其形状是长方体,点云数据形状为“L”型。虽然武器箱的表面有各种凹凸部分和油桶的边缘是曲面,但是描述体积仍是长宽高之积。由于激光雷达的诸多优点和自动驾驶领域的火热,对于物体的三维检测也出现了 PointPillars^[13]、PointRCNN^[14]和 PartA2^[15]等深度学习方法,但因其大多针对的是道路场景、识别目标较大、需要制作数据集和落地部署困难等问题使得难以迁移到“L”型物体实时体积测量中。

为解决激光雷达点云的数据稀疏、扰动和噪声对于小目标体积测量影响严重的问题。本文首次提出一种基于视觉修正的激光雷达体积测量方法实现小尺寸“L”型物体的体积测量。主要工作为先对相机与激光雷达进行数据对齐;之后对图像进行目标检测和对点云进行地面分割,以目标检测的结果作为点云目标分割约束条件并结合激光雷达线束分布特性确定点云聚类参数从而更好地分割出目标点云;最后设计了用于对“L”型目标点云进行三维框拟合的算法,并提出视觉修正模型对目标三维框进行修正得到目标体积。实验证明,本文方法可以实现对“L”型目标体积的实时、高精度、自动化测量,对于体积测量的研究有积极意义。

1 基于视觉修正的激光雷达目标体积测量算法框架设计

目前大多数目标体积测量方法都基于特殊的测试场景或者已知测量仪器的位姿条件,实际应用中受到的限制较大。为此,本文提出了一种基于视觉修正的激光雷达体积测量方法,只需要相机和激光雷达,传感器可以搭载在移动平台上,可以对于较小“L”型目标进行实时体积测量,算法框架如图 1 所示。

首先,利用 MATLAB 工具箱对采集的点云和图像进行空间上的标定,求得相机的内参与激光雷达到相机的外参,保证空间上对齐;对获取的图像和点云数据采用最近邻算法缩小时间戳上的误差,保证点云和图像数据同步,使得激光雷达与相机的数据在时空上对齐。然后,在图像上进行目标检测获得目标的种类与位置信息;对

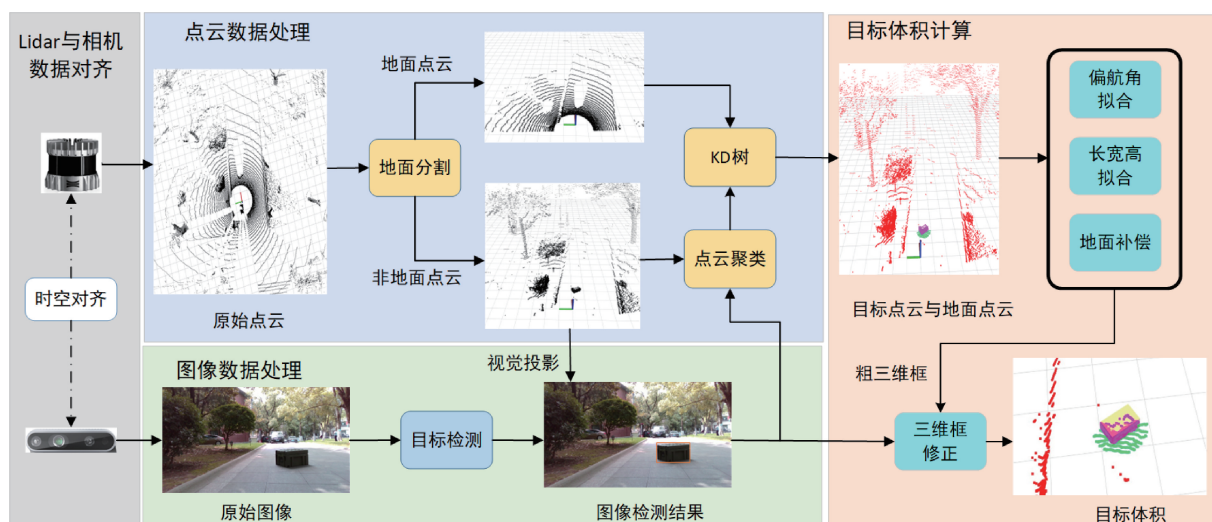


图1 基于视觉修正的激光雷达体积测量算法框架

Fig. 1 A framework for LiDAR volume measurement algorithm based on visual correction

点云数据执行地面分割得到地面点云与非地面点云,将非地面点云通过外参投影到图像中,根据目标的图像位置信息将对应点云进行聚类,聚类的输入参数根据激光雷达的线束分布和待聚类点云确定,从而分割出目标点云,利用KD树和目标点云的中心点寻找附近地面点云。最后,根据目标点云进行偏航角和长宽高的拟合,并且对目标附近地面点云进行平面拟合得到高度上的地面补偿;提出一种修正模型,利用图像上目标二维位置信息对于前面得到的粗三维框进行修正得到目标的体积,以此解决点云的噪声、扰动和稀疏等问题。

2 算法设计

2.1 目标检测算法

因为需要将图像中的目标信息作为约束和补充,所以本文采用YOLOv7-tiny^[16]算法进行目标的检测,其是一种单阶段目标检测算法,速度和精度达到目标检测领域的SOTA。框架主要分为特征提取模块,特征融合模块和预测输出模块。在网络中大量运用了扩展高效层聚合网络,通过控制不同长度梯度路径,运用分组卷积、扩张和混洗模块获得更加多样的特征,使得网络更有效地学习和收敛。为满足不同推理速度需求,提出了基于拼接的模型缩放,调整基础模块的放缩参数同时计算改变转移层的宽度缩放因子以此保持最优结构。在预测输出模块运用了重参数化技巧,通过增加网络分支增强网络的特征提取,推理时通过参数的合并来提升速度。在计算损失函数上,通过多标签分配策略,给模型添加额外的监督信息,提升学习能力。通过这些创新在精度和速度上取得了良好效果。

2.2 点云数据处理

1) 地面分割算法

在进行目标体积计算之前,需要辨别地面点和非地面点,避免地面点影响后面三维框拟合部分。点云中常用的地面分割方法有基于几何特征的高度法、坡度法和栅格法,基于聚类的区域生长法,基于统计和模型拟合的随机抽样一致(random sample consensus, RANSAC)、PMF和LineFit^[17],基于机器学习和深度学习的一些算法^[18-19]等。基于算法需要在端侧运行、测量对象相对较小、对周围地面点分割要求较为严格和点云具有扰动性等考虑,本文采用LineFit作为地面分割的算法。

LineFit是一种基于线的地面分割方法,有着高精准率并且运行速度非常快。其通过点云的有序化处理、点云的降维和直线的拟合完成点云的分割。如图2所示,在点云的有序化中先将 $x-o-y$ 平面转变成一个半径为无穷的圆面,将圆面按照 $\Delta\alpha$ 分成不同的segment扇面,然后通过设置距离参数将segment划分成不同的bin扇环。接下来在不同的segment将三维数据 $\{x, y, z\}$ 转化成二维数据 $\{d, z\}$,其中 d 是 $x-o-y$ 平面离圆心的距离,实现点云的降维。最后通过设置直线的斜率、直线在 z 轴的截距、直线的拟合误差和bins中开始点到直线的距离完成直线的拟合,设定点云到直线的 z 轴阈值完成地面点的筛选。

2) 点云聚类

得到非地面点云,需要提取目标的点云,本文通过视觉投影的方式结合目标检测的结果滤除大部分非目标点云,然后对剩余点云进行聚类从而分割出目标点云。视觉投影通过激光雷达到相机的外参和相机的内参计算完成,设任意一点 P 在激光雷达坐标系中坐标为 (P_x, P_y, P_z) ,在

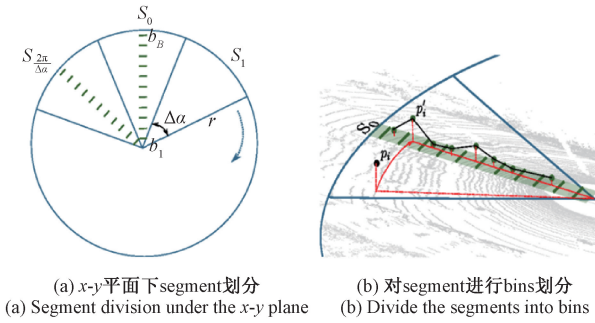


图 2 LineFit 点云有序化

Fig. 2 LineFit point cloud ordering

相机坐标系中坐标为 (X_w, Y_w, Z_w) , 在像素坐标系中坐标为 (x, y) , 点 P 从 Lidar 坐标系到相机坐标系则有如下转换关系:

$$\begin{bmatrix} X_w \\ Y_w \\ Z_w \\ 1 \end{bmatrix} = \begin{bmatrix} r_{x1} & r_{y1} & r_{z1} & t_x \\ r_{x2} & r_{y2} & r_{z2} & t_y \\ r_{x3} & r_{y3} & r_{z3} & t_z \\ 0 & 0 & 0 & 1 \end{bmatrix} \times \begin{bmatrix} P_x \\ P_y \\ P_z \\ 1 \end{bmatrix} \quad (1)$$

式中: $\{r_{x1} \sim r_{x3}, r_{y1} \sim r_{y3}, r_{z1} \sim r_{z3}\}$ 代表了激光雷达坐标系到相机坐标系的旋转参数; t_x, t_y, t_z 分别表示在激光雷达坐标系 X, Y 和 Z 轴的平移量。根据相机的成像模型, 可得点 P 相机坐标系到像素坐标系的转换关系:

$$Z_w \begin{bmatrix} x \\ y \\ 1 \end{bmatrix} = \begin{bmatrix} f_x & 0 & u_0 \\ 0 & f_y & v_0 \\ 0 & 0 & 1 \end{bmatrix} \times \begin{bmatrix} X_w \\ Y_w \\ Z_w \end{bmatrix} \quad (2)$$

式中: f_x 和 f_y 代表了在相机坐标系 x 轴和 y 轴的焦距; u_0 和 v_0 代表了主点在 x 轴和 y 轴的实际位置。

将点云投影到图像中, 根据目标检测的结果 $(x_{\min}, y_{\min}, \text{width}, \text{height})$ 判断是否为待聚类的点, 其结果如图 3 所示, 在目标检测框内的实心点为待聚类的点, 在目标检测框之外的空心点为滤除点。

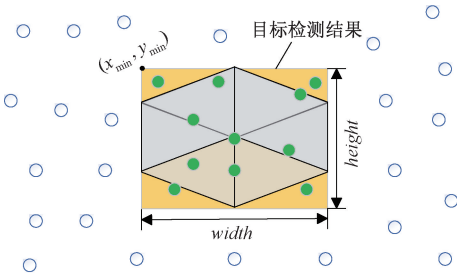


图 3 视觉投影结果

Fig. 3 Visual projection results

常用的点云聚类的方法多数基于距离和点云密度考虑。K-means^[20]将点云中的点划分为预先指定的 K 个簇, 通过迭代优化来使每个点分配到最近的簇中心;

欧氏聚类^[21]从点云中选择一个点作为种子点, 并找到与该种子点距离在阈值范围内的邻近点, 将它们划分为一个簇; DBSCAN^[22]将点云中的点分为核心点、边界点和噪声点, 基于点的密度来确定簇的边界, 并不需要预先指定簇的数量。由于本文针对的测量对象体积较小、距离越远点云越稀疏和激光雷达到相机的 x 轴偏移量较大导致出现背景噪声的原因, 采用 DBSCAN 点云聚类方法。

DBSCAN 根据数据点的密度来划分簇, 定义了两个参数, 邻域半径 ϵ 和邻域内最小数据点数量 MinPts。对于给定数据集中的一个数据点, 如果其邻域内至少包含 MinPts 个数据点 (包括该点本身), 则将该点标记为核心点。对于核心点 p 和 q , 如果 q 在以 p 为中心, 半径为 ϵ 的邻域内, 并且 p 的邻域包含 q , 则称 q 直接密度可达于 p , 如果存在一系列核心点 $p_1, p_2, \dots, p_n$, 使得 p_{i+1} 直接密度可达于 $p_i (1 \leq i \leq n-1)$, 并且 q 直接密度可达于 p_n , 则称 q 密度可达于 p 。其算法流程如下:

(1) 初始化, 将所有数据点标记为未访问状态;

(2) 随机选择一个未访问的数据点 p ;

(3) 如果 p 的邻域 ϵ 内包含至少 MinPts 个数据点, 则创建一个新簇, 并将 p 标记为该簇的一个核心点, 然后, 通过密度可达关系将 p 的邻域中的所有未访问点添加到该簇中;

(4) 重复步骤 (3), 直到所有核心点的邻域都被探索完毕;

(5) 将未访问的数据点标记为噪声点。

针对实际应用中参数难调的问题, 本文基于激光雷达的线束分布情况来自动确定输入的参数。实验中采用的激光雷达为 OS1-64, 水平分辨率为 0.35° , 垂直分辨率为 0.7° , 因此本文重点关注垂直方向的线束分布情况, 如图 4 所示。

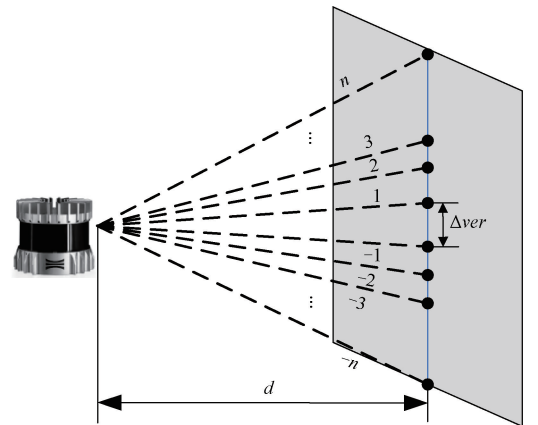


图 4 激光雷达垂直方向线束分布

Fig. 4 Vertical bundle distribution of Lidar

对于同一垂直方向,两点之间的间隔取决于离激光雷达的水平距离和垂直距离。设待聚类的点云重心在 z 轴上的值为 z_d ,到激光雷达的水平距离为 d ,则垂直间隔的计算公式如下:

$$\begin{cases} \Delta_{ver} = d \times \left\| \tan(0.7^\circ n - 0.35^\circ) - \tan(0.7^\circ n - 1.05^\circ) \right\| \\ n = \frac{\arctan(\|z_d\|/d)}{0.7^\circ} + 0.5 \\ n = 1, 2, 3, \dots, 32 \end{cases} \quad (3)$$

式中: Δ_{ver} 表示垂直方向的间隔; n 表示激光雷达的线束序号。根据式(3)设定 DBSCAN 的邻域半径 ε 为 Δ_{ver} 的 2 倍,最小数据点 MinPts 为 10 个。

当点云聚类成功后,如图 3 所示,由于背景点云的存在,一般会聚类出若干个点云簇,需要从中寻找目标所对应的点云簇。

(1) 计算每个点云簇的重心坐标 w_i 。

(2) 根据式(1)与(2)将点云簇的重心坐标 w_i 投影到像素坐标系中得到其像素坐标 (u_i, v_i) 。

(3) 计算每一个点云簇重心的像素坐标离图像目标二维检测框中心点 (c_x, c_y) 的距离 d_i , 计算公式为:

$$d_i = \sqrt{(u_i - c_x)^2 + (v_i - c_y)^2} \quad (4)$$

(4) 选取 d_i 最小的点云簇为目标对应的点云。

通过上述方法得到目标点云后,先对地面点云建立 KD 树,利用目标点云重心执行近邻查询算法,得到目标所对应地面点云,以便之后的地面高度补偿。

2.3 目标体积计算

1) 三维框拟合

由于点云基本集中于目标的前表面,直接对于目标点云进行三维框拟合会导致错误。对于三维框的拟合部分,本文的主要思想是先将目标点云投影到 $x-y$ 平面,拟合出偏航角(为简化算法的难度,不考虑俯仰角和翻滚角),之后再计算出长宽,结合其对应的地面点云计算目标的高度,其具有 7 个自由度 $(c_x, c_y, c_z, l, w, h, \alpha)$ 。

在二维平面中,偏航角与矩形的长宽常合在一起计算,相当于根据点集拟合旋转矩形框,常用的评价拟合性能的经典准则是最小二乘法。由于角度所产生的分区问题导致直接求解最小二乘法的优化过于困难,并且武器箱道具、医疗箱等目标的点云形状类似于“L”,用最小面积的准则拟合效果不理想。因此本文采用遍历搜索的方式寻找损失最小的偏航角,其中采用的损失函数 Varcriterion 是点到边的平方误差^[23],进而计算长宽和中心点。先遍历每一个角度,将点云旋转得到两个垂直方向上的距离 C_1 与 C_2 ,通过点到边的最小平方误差准则找到最优的角度作为旋转矩形框的偏航角 α ; 计算偏航角对应的 C_1 与 C_2 ,其代表了点云在长度和宽度方向上

的值,将其所在区间等分成若干份,找到点云数量最多的区间,由于距离激光雷达最近的矩形框角点附近云密度最高,所以对于 C_1 而言,点云数量峰值所在的区间离其最小值更近;对于 C_2 而言,点云数量峰值所在区间离其最大值更近;以 C_1 与 C_2 的峰值区间至离其最近的值(最大值或者最小值),寻找其峰值的 m 倍所在位置作为角点,即可以滤除两边的点云毛刺(比如说未分割完全的部分地面点云)又可以减小点云扰动性的影响,此时可以确定矩形框的长 l 与宽 w 和中心点 (c_x, c_y) ,其算法如算法 1 所示。

算法 1 基于搜索的旋转矩形框拟合

输入: 目标点云 $P \in R^{n \times 3}$

输出: 矩形框的属性 $\{c_x, c_y, length, width, \alpha\}$

```

1:  将  $P$  投影至  $x-y$  平面, 得到  $X \in R^{n \times 2}$ 
2:   $Q \leftarrow 1\ 000, \alpha \leftarrow 0$ 
3:  for  $\theta = 0$  to  $\pi/2 - \delta$  step  $\delta$  do
4:       $e_1 \leftarrow (\cos \theta, \sin \theta), e_2 \leftarrow (-\sin \theta, \cos \theta)$ 
5:       $C_1 \leftarrow X \cdot e_1^T, C_2 \leftarrow X \cdot e_2^T$ 
6:       $q \leftarrow \text{Varcriterion}(C_1, C_2)$ 
7:      if  $q < Q$  then
8:           $\alpha \leftarrow \theta, Q \leftarrow q$ 
9:      end if
10: end for
11:  $e_1 \leftarrow (\cos \alpha, \sin \alpha), e_2 \leftarrow (-\sin \alpha, \cos \alpha)$ 
12:  $C_1 \leftarrow X \cdot e_1^T, C_2 \leftarrow X \cdot e_2^T$ 
13:  $k \leftarrow 0.02, cp \leftarrow (0, 0)$ 
14: for  $C \in \{C_1, C_2\}$  do
15:     delta  $\leftarrow k \times (\max(C) - \min(C)), \text{Count} \leftarrow \emptyset$ 
16:     for  $i = 0$  to  $1/k - 1$  step 1 do
17:         对  $[i, i+1] \cdot \text{delta}$  区间进行点云数量计数
18:         将  $\{\text{序号 } i, \text{数量}\}$  插入到 Count 中
19:     end for
20:     对 Count 按照数量进行排序
21:     找出 Count 中数量最大值  $\{\max Id, \max Val\}$ 
22:      $m \leftarrow 0.3$ 
23:     if  $\max Id \geq (1/k - 1 - \max Id)$  then
24:         滤除 Count 中序号  $i < \max Id$  元素得  $M$ 
25:          $M$  中数量离  $m \times \max Val$  最近的序号  $Id$ 
26:          $width \leftarrow Id \times \text{delta}$ 
27:          $cp(1) \leftarrow \min(C) + width/2$ 
28:     else
29:         滤除 Count 中序号  $i > \max Id$  元素得  $M$ 
30:          $M$  中数量离  $m \times \max Val$  最近的序号  $Id$ 
31:          $length \leftarrow Id \times \text{delta}$ 
32:          $cp(0) \leftarrow \min(C) + length/2$ 
33:     end if
34: end for
35:  $(c_x, c_y)^T \leftarrow (e_1; e_2)^{-1} \cdot cp^T$ 

```

当不考虑目标是否放在地面上时,可直接根据目标点云簇进行高度和中心值的计算。设给定目标点云簇 P 为 $\{(x_i, y_i, z_i) \mid i=1, 2, 3, \dots, k\}$, k 为点云数量, h 为目标三维框高度, c_z 为目标三维框在 z 轴方向的中心值, 则有计算公式

$$\begin{cases} h = \max\{z_1, z_2, \dots, z_k\} - \min\{z_1, z_2, \dots, z_k\} \\ c_z = \frac{\max\{z_1, z_2, \dots, z_k\} + \min\{z_1, z_2, \dots, z_k\}}{2} \end{cases} \quad (5)$$

如果目标放置在地面,为提高其高度上的可靠性,可先通过地面点云拟合一个平面,由目标中心 (c_x, c_y) 计算出地面所对应的 z 轴方向的值,结合目标点云簇的 z 轴最大值推出目标的 h 和 c_z 。设给定目标附近地面点集 Q 为 $\{(x_i, y_i, z_i) \mid i=1, 2, 3, \dots, n\}$, n 为点云数量,地面平面方程可以表示为:

$$z = ax + by + c \quad (6)$$

式中: a 、 b 、 c 是平面的系数。根据最小二乘法拟合原理,平面拟合的残差之和为:

$$E = \sum_{i=1}^n (ax_i + by_i + c - z_i)^2 \quad (7)$$

当 E 取最小值时, $\frac{\partial E}{\partial a} = \frac{\partial E}{\partial b} = \frac{\partial E}{\partial c} = 0$, 即有:

$$\begin{bmatrix} \sum x_i^2 & \sum x_i y_i & \sum x_i \\ \sum x_i y_i & \sum y_i^2 & \sum y_i \\ \sum x_i & \sum y_i & n \end{bmatrix} \times \begin{bmatrix} a \\ b \\ c \end{bmatrix} = \begin{bmatrix} \sum x_i z_i \\ \sum y_i z_i \\ \sum z_i \end{bmatrix} \quad (8)$$

求解式(8)可得到平面方程系数 a^* 、 b^* 、 c^* 。则目标 h 和 c_z 计算公式为:

$$\begin{cases} z_g = a^* c_x + b^* c_y + c^* \\ h = \max\{z_i \mid (x_i, y_i, z_i) \in P\} - z_g \\ c_z = 0.5h + z_g \end{cases} \quad (9)$$

至此,得到三维框在激光雷达坐标系中的中心点坐标、长宽高和偏航角 7 个自由度上的值 $(c_x, c_y, c_z, l, w, h, \alpha)$ 。

2) 视觉修正模型

不同于稀疏的激光雷达数据,图像数据含有丰富的语义信息,对于目标的二维边界判断更为精准,将激光雷达坐标系中的拟合出来的三维框投影至像素坐标系中,结合目标检测结果对于三维框进行修正,可以在很大程度上解决由于点云稀疏性和噪声点带来的边界不准问题。本文基于此提出基于目标检测结果的旋转三维框修正模型,如图 5 所示,其中包含多个坐标系的转换和相机成像模型,其中 $O_l-X_l Y_l Z_l$ 是激光雷达坐标系, o_b-xyz 是激光雷达坐标系下物体自身坐标系,沿 x 轴是长度所在方向,沿 y 轴是宽度所在方向,沿 z 轴方向是高度方向, $O_c-X_c Y_c Z_c$ 是相机坐标系, o_o-uv 是像素坐标系。

对于激光雷达坐标系下任意一点 P ,由式(1)和(2)可以得到在相机坐标系和像素坐标系下的所应的点。

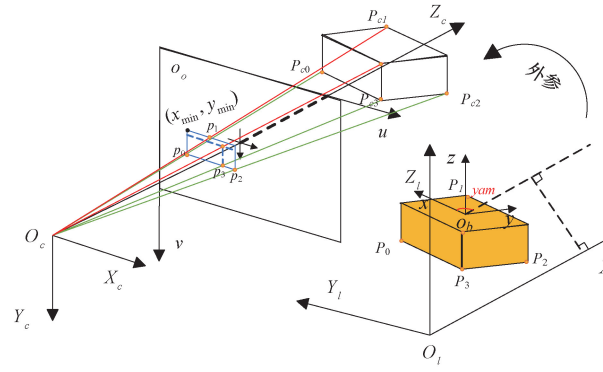


图 5 视觉修正模型

Fig. 5 Visual correction model

图 5 中, P_0 、 P_1 、 P_2 、 P_3 为激光雷达下物体的关键点, P_{c0} 、 P_{c1} 、 P_{c2} 、 P_{c3} 为相机坐标系下物体的关键点, p_0 、 p_1 、 p_2 、 p_3 为像素坐标系下物体的关键点,其中 y_{am} 为物体坐标系与激光雷达坐标的 x 轴的夹角,设其值为 α ,即物体的偏航角,其值在 $0 \sim \pi/2$ 。

在 $X_l-O_l-Y_l$ 平面中,物体的中心点为 (c_x, c_y) ,拟合的长宽为 l 和 w ,则 P_3 的坐标为:

$$\begin{cases} (x|P_3) = \frac{-l \cos \alpha - w \sin \alpha}{2} + c_x \\ (y|P_3) = \frac{-l \sin \alpha + w \cos \alpha}{2} + c_y \end{cases} \quad (10)$$

将目标投影至 $X_l-O_l-Y_l$ 平面,由于 P_3 离激光雷达最近,附近点云密度最高,因此位置的可信度相比于其余 3 个关键点更高,本文以此重新推算 P_0 与 P_2 的坐标。设目标的修正后的长为 l' ,宽为 w' ,则 P_0 和 P_2 的坐标为:

$$\begin{cases} (x|P_2) = (x|P_3) + w' \sin \alpha \\ (y|P_2) = (y|P_3) - w' \cos \alpha \\ (x|P_0) = (x|P_3) + l' \cos \alpha \\ (y|P_0) = (y|P_3) + l' \sin \alpha \end{cases} \quad (11)$$

根据式(1)可得到 P_0 与 P_2 在相机坐标系中的坐标 P_{c0} 与 P_{c2} 在 $X_c-O_c-Z_c$ 平面的值为:

$$\begin{cases} (x|P_{c0}) = r_{x1}(x|P_0) + r_{y1}(y|P_0) + t_x \\ (z|P_{c0}) = r_{x3}(x|P_0) + r_{y3}(y|P_0) + t_z \\ (x|P_{c2}) = r_{x1}(x|P_2) + r_{y1}(y|P_2) + t_x \\ (z|P_{c2}) = r_{x3}(x|P_2) + r_{y3}(y|P_2) + t_z \end{cases} \quad (12)$$

根据相机的成像模型式(2)可以得到 P_0 与 P_2 在像素坐标系所对应的 p_0 与 p_2 的横坐标值为:

$$\begin{cases} (u|p_0) = f_x \frac{(x|P_{c0})}{(z|P_{c0})} + u_0 \\ (u|p_2) = f_x \frac{(x|P_{c0})}{(z|P_{c0})} + u_0 \end{cases} \quad (13)$$

设目标检测的二维框左上角坐标为 $(x_{\min}, y_{\min})$, 右下角坐标为 $(x_{\max}, y_{\max})$, 令 $(u|p_0) = x_{\min}, (u|p_2) = x_{\max}$, 长与宽的修正公式则有:

$$\begin{cases} w' = \frac{C_1 T_2 - T_1}{(r_{x1} - C_1 r_{x3}) \sin \alpha - (r_{y1} - C_1 r_{y3}) \cos \alpha} \\ l' = \frac{C_2 T_2 - T_1}{(r_{x1} - C_2 r_{x3}) \cos \alpha + (r_{y1} - C_2 r_{y3}) \sin \alpha} \end{cases} \quad (14)$$

式中: $C_1 = \frac{x_{\max} - u_0}{f_x}, C_2 = \frac{y_{\min} - v_0}{f_y}$; T_1 与 T_2 为 P_3 在相机坐标系下的 X_c 轴和 Z_c 轴上的值。

$$\begin{cases} T_1 = r_{x1}(x|P_3) + r_{y1}(y|P_3) + t_x \\ T_2 = r_{x3}(y|P_3) + r_{y3}(y|P_3) + t_z \end{cases} \quad (15)$$

由 P_3 和修正后的 l' 与 w' , 结合式 (10) 可以得到物体在 $X_l-O_l-Y_l$ 平面上的新中心点 (c'_x, c'_y) 为:

$$\begin{cases} c'_x = (x|P_3) + \frac{l' \cos \alpha + w' \sin \alpha}{2} \\ c'_y = (y|P_3) + \frac{l' \sin \alpha - w' \cos \alpha}{2} \end{cases} \quad (16)$$

因为激光雷达的放置位置往往高于物体, 点 P_1 投影至像素坐标系中会落在目标二维框的顶边, 所以可利用点 P_3 与 P_1 对于物体的 h 和 c_z 进行修正。利用已修正的信息, 可得点 P_1 在 $X_l-O_l-Y_l$ 平面的坐标为:

$$\begin{cases} (x|P_1) = \frac{l' \cos \alpha + w' \sin \alpha}{2} + c'_x \\ (y|P_1) = \frac{l' \sin \alpha - w' \cos \alpha}{2} + c'_y \end{cases} \quad (17)$$

与前面类似原理, 可得修正后的 h' 与 c'_z 的联立方程组:

$$\begin{cases} c'_z - \frac{h'}{2} = \frac{C_3(z|P_{c3}) - (y|P_{c3})}{r_{z2} - C_3 r_{z3}} \\ c'_z + \frac{h'}{2} = \frac{C_4(z|P_{c1}) - (y|P_{c1})}{r_{z2} - C_4 r_{z3}} \end{cases} \quad (18)$$

式中: $C_3 = \frac{y_{\max} - v_0}{f_y}, C_4 = \frac{y_{\min} - v_0}{f_y}$ 。求解式 (18) 可得三维框的高 h' 和在 Z_l 轴上的中心点 c'_z 。

3 实验验证与分析

3.1 实验平台搭建

为了验证本文所提算法的性能, 以较为规则、容易计算真实体积的武器箱道具、医疗箱和油桶作为实验对象进行测试, 其中各物体的尺寸如表 1 所示。

实验平台使用小车作为移动装置, 移动电源为 OS1-64 线激光雷达、Inter RealSense D435i 相机、Nvidia Jetson AGX Xavier 嵌入式平台和显示屏供电, 实验平台如图 6 所示。其中 OS1-64 线激光雷达水平视场角为 360° , 垂直

表 1 物体尺寸和体积

Table 1 Object size and volume

种类	长 / 10^{-1} m	宽 / $(\times 10^{-1} \text{ m})$	高 / $(\times 10^{-1} \text{ m})$	体积 / $(\times 10^{-3} \text{ m}^3)$
武器箱道具	7.05	4.50	3.75	118.97
医疗箱	3.56	2.31	2.31	19.00
油桶	3.40	1.60	4.50	24.48

视场角为 $\pm 22.5^\circ$, 水平角分辨率为 0.35° , 垂直角分辨率为 0.7° , 在 $1 \sim 20 \text{ m}$ 测距精度为 $\pm 1 \text{ cm}$, 扫描频率为 10 Hz ; Inter RealSense D435i 相机有两个近红外摄像头和一个可见光摄像头, 实验只采用可见光摄像头, 其视场角为 $69^\circ \times 42^\circ$, 分辨率配置为 1280×720 , 帧率为 30 fps 。

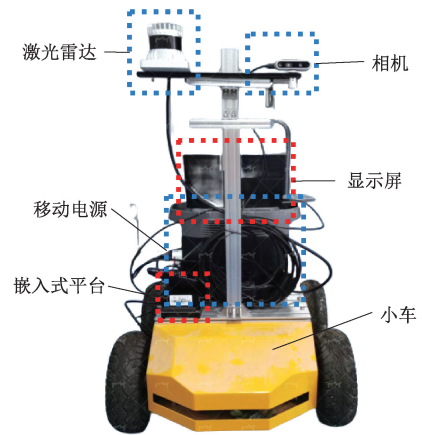


图 6 实验平台示意图

Fig. 6 Diagram of experimental platform

算法在 Nvidia Jetson AGX Xavier 嵌入式平台中基于 C++ 开发, 在 ROS 框架下订阅激光雷达与相机驱动发布的话题, 对两个传感器的数据进行实时处理与计算, 其中软件环境如表 2 所示

表 2 软件环境需求

Table 2 Software environment requirements

软件名称	版本
CUDA	10.2
PCL	1.8.1
ROS	Melodic
OpenCV	4.1.1
TensorRT	7.1.3

3.2 实验指标

实验过程激光雷达离地面的高度 $H_{\text{Lidar}} \approx 0.75 \text{ m}$, 由激光雷达的垂直和水平角分辨率, 结合式 (3) 易得地平面上的水平和垂直方向上的分辨率, 其随距离的变化如

图 7 所示。由激光雷达的垂直方向的视场角,可得最小测量距离为:

$$d_{\min} \approx 2.4H_{\text{Lidar}} + d_0$$

(19)

为保证地面分割算法的准确性,取 $d_0=0.2\text{ m}$ 。

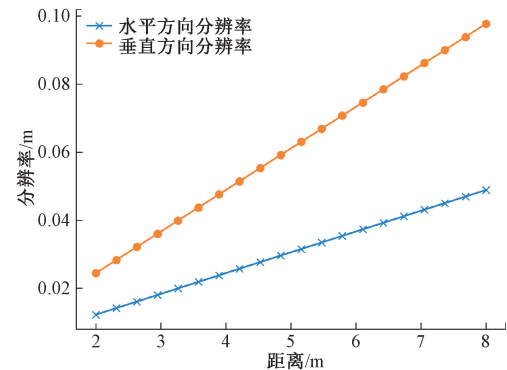


图 7 分辨率随距离变化

Fig. 7 Resolution varies with distance

测试过程中,点云的点会在 2~3 cm 波动,如图 8 所示,将相邻 4 帧点云用不同的颜色和形状区分,每个小方格的边长为 1 cm,可以证实点云具有波动性,其对于小尺寸物体的测量不容小觑。基于以上分析,本文对于小尺寸物体的定义为在测量点,点云的垂直分辨率加上点云波动性产生的点云最大间隔大于物体最小边长的 1/10。由表 1 和图 7 可知,本文测量对象在最小测量距离处均满足小尺寸物体的定义。

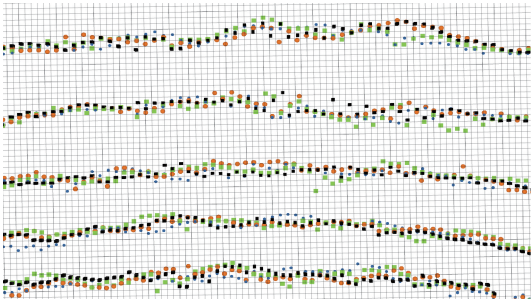


图 8 点云扰动性示意图

Fig. 8 Diagram of point cloud disturbance

虽然通过多帧累计滤波下采样对于大尺寸物体测量是一个解决点云扰动性的有效办法,但对于小尺寸物体滤波下采样会使得原本稀疏的物体点云变得更加稀疏,且多帧累计对于地面点噪声而言也是累计的作用。

基于上述分析,本文采用平均相对误差 e 和重复性 δ 的指标来评价算法的性能,其计算公式为:

$$e = \frac{1}{N} \sum_{i=1}^N \frac{\|v_i - v^*\|}{v^*} \times 100\%$$

(20)

$$\delta = \frac{\sqrt{\sum_{i=1}^N (v_i - \bar{v})^2}}{\sqrt{(N-1)v^2}} \times 100\%$$

(21)

式中: N 为在相同测试条件下点云和图像的帧数; v^* 为物体真实体积; $\bar{v}$ 为物体测量的平均体积。

3.3 实验结果与分析

本文通过录制每种物体在不同距离下 3 s 的数据包进行测试。受相机实际传输速率不稳定的影响,点云和图片的对数有一定浮动,对每个数据包随机取其中 15 帧点云和图片作为本文的测试数据来源。在不同距离下油桶体积、武器箱道具和医疗箱测量结果如表 3~5 所示,在一定的范围内,各类目标体积测量的精度在一个值附近浮动,但超出一定距离后,随着作用距离的变远,体积测量的精度开始下降并且浮动变大。如对油桶而言,宽度较小、边缘是弧角且点云对于其反射率不佳,导致距离一远,点云分辨率降低,宽度方向点云缺失严重,精度下降严重。基于这类考虑,本文将油桶、武器箱和医疗箱的测量范围分别限定为 2.0~3.7、2.0~5.9 与 2.0~5.5 m。在 2~3.7 m 范围内,油桶的最大平均误差小于 4.92%,所有距离平均误差为 4.44%,重复性小于 5.61%;在 2~5.9 m 范围内,武器箱道具的最大平均误差不超过 5.47%,重复性小于 3.42%,所有距离的平均误差为 4.41%;在 2~5.5 m 范围内,医疗箱的最大平均误差不超过 6.12%,重复性小于 3.26%,所有距离的平均误差为 3.74%。

表 3 油桶体积测量结果

Table 3 Measurement results of oilbarrel volume

真实体积 /($\times 10^{-3}\text{ m}^3$)	距离/m	平均测量体积 /($\times 10^{-3}\text{ m}^3$)	平均误差 /%	重复性 /%
24.48	2.0	24.01	3.47	3.94
	2.5	23.86	4.77	5.60
	3.0	24.88	4.62	4.98
	3.7	23.62	4.92	4.62
	4.5	13.62	44.37	34.26
	5.5	11.27	54.14	23.34

表 4 武器箱道具体积测量结果

Table 4 Measurement results of weapon box prop volume

真实体积 /($\times 10^{-3}\text{ m}^3$)	距离/m	平均测量体积 /($\times 10^{-3}\text{ m}^3$)	平均误差 /%	重复性 /%
118.97	2.0	112.44	5.47	1.19
	2.5	114.28	4.05	1.79
	3.0	120.61	4.17	2.41
	3.7	124.09	4.31	1.80
	4.8	115.84	3.70	3.42
	5.9	124.62	4.77	2.54
	6.5	129.35	8.71	3.78
	7.5	135.27	13.71	7.56

表 5 医疗箱体积测量结果

Table 5 Medical box volume measurement results

真实体积 /($\times 10^{-3} \text{ m}^3$)	距离/m	平均测量体积 /($\times 10^{-3} \text{ m}^3$)	平均误差 /%	重复性 /%
19.00	2.0	18.30	3.79	2.26
	2.5	19.01	2.05	2.68
	3.0	19.24	3.26	2.41
	3.7	19.72	3.85	1.96
	4.3	20.16	6.12	1.25
	5.5	19.67	3.89	2.45
	6.5	20.16	8.63	7.53
	7.5	20.02	9.92	13.34

如图 9 所示,点云中细框线表示粗三维框,粗框线表示修正后的三维框,图像二维框内实心点表示投影的目标点云,二维框表示图像目标检测的结果,左右边界分别带有实心点的表示粗三维框的投影,左右边界无实心点表示修正后的三维框投影,其中粗三维框拟合算法对于物体的偏航角拟合度较好,加上其对于物体坐标系长度和宽度负方向有抗干扰性,使得离激光雷达最近的角点接近物体边缘,而基于视觉的三维框修正模型物体坐标系长度和宽度正方向有着很好地去噪声点的作用和使得物体贴合图像的检测结果,使得武器箱体体积修正相对误差由 26.2% 下降到-0.3%、医疗箱体体积修正相对误差由 10.7% 下降到-4.8% 和油桶的体积修正相对误差由 15.9% 下降到-5.7%。

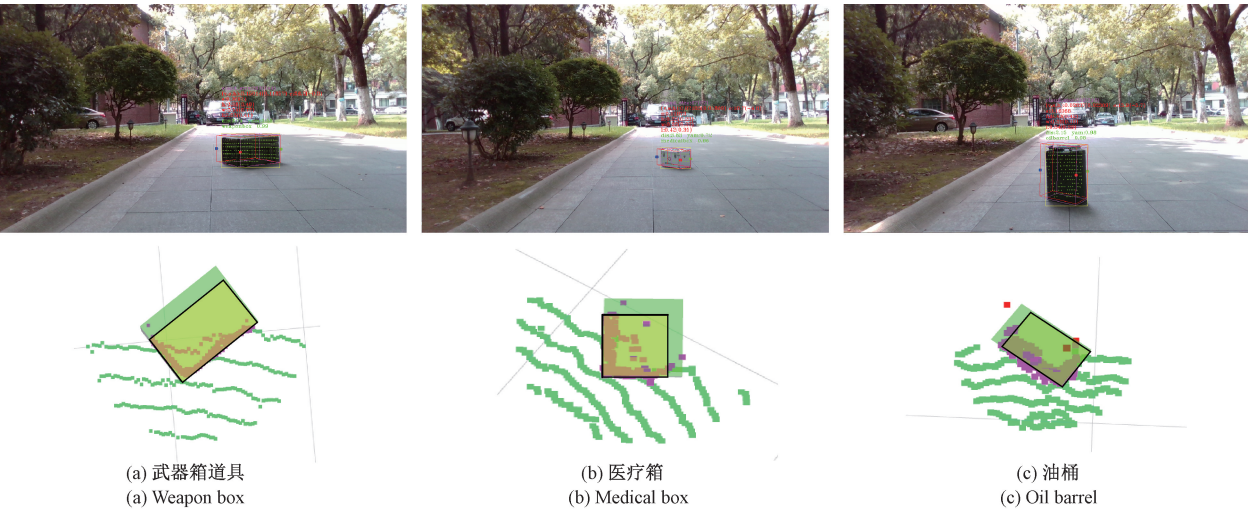


图 9 不同物体视觉修正模型结果示意图

Fig.9 Diagram of visual correction model results

如表 6 所示,武器箱道具体积经过修正后所有测量点的平均误差下降了 8.20%,最大误差下降了 13.7%,最大重复性由 18.09% 下降到 3.42%;医疗箱体体积经过修正后所有测量点的的平均误差下降了 19.90%,最大误差下降了 32.95%,最大重复性由 22.38% 下降到 3.26%;油桶体积经过修正后所有测量点的的平均误差下降了 21.49%,最大误差下降了 43.55%,最大重复性由 22.38% 下降到 3.26%。可知视觉的修正模型大幅提高了体积测量的精度和稳定性。如表 7 所示,以医疗箱在 2.0 m 处 4 个时刻的视觉修正模型前后体积测量精度为例,可以明显地看出在本文算法中视觉修正模型对于体积测量精度和稳定性的重要性。

当多个目标同时存在并且目标检测二维框不重叠时,本文所提算法仍然适用,如图 10 所示,多目标同时存在时,医疗箱、油桶和武器箱的体积测量误差分别为

表 6 体积修正前后精度变化

Table 6 Accuracy changes before and after volume correction

种类	%					
	平均误差		最大误差		最大重复性	
	修正前	修正后	修正前	修正后	修正前	修正后
武器箱道具	12.61	4.41	19.17	5.47	18.07	3.42
医疗箱	23.64	3.74	39.07	6.12	22.38	3.26
油桶	25.93	4.44	48.47	4.92	23.96	5.61

5.6%、2.5% 和-1.6%。并且值得一提的是本文算法平均每帧数据的处理时间为 55 ms,其中目标检测算法时间约为 14 ms,点云的处理与体积的计算时间约为 41 ms,能够达到实时处理的要求。由于本文使用的是机械式激光雷达,其通过电机旋转激光发射模块和接收模块完成周围环境的三维感知,能在白天和黑夜的环境下工作,但

表 7 医疗箱在 2.0 m 不同时刻测量视觉修正前后结果

Table 7 Measurement results before and after visual correction at different times at 2.0 m of the medical box

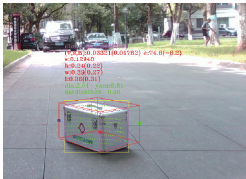
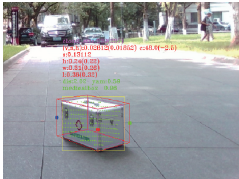
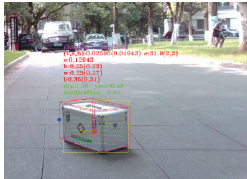
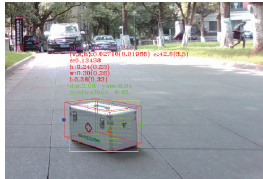
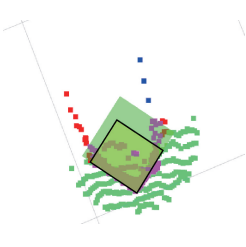
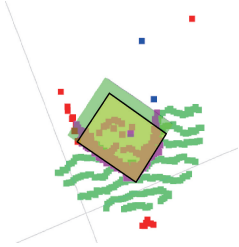
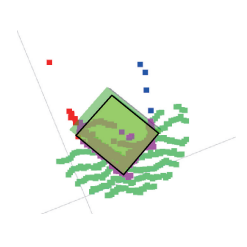
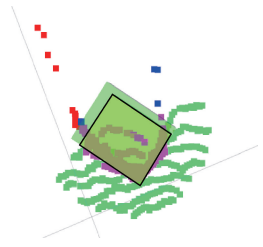
时刻	t_0	t_1	t_2	t_3
				
示意图				
$e/\%$	74.8→-6.2	48.0→-2.5	31.9→2.3	42.6→3.5



图 10 多目标体积测量结果

Fig. 10 Multi target volume measurement results

在中大雨天、浓雾浓烟等环境中,激光的传播易受影响,点云的质量急速下降,可见光传感器易受环境光的影响,本文算法依赖于图像目标检测精度和点云的质量,所以本文算法适用于光线充足、空气纯净的天气。实验证明在这样的环境下,本文能够实现精度高、速度快的体积测量,但由于实验条件和实验设备的限制,未能在其他复杂环境下做进一步的研究。

4 结 论

针对大部分体积测量方法迁移困难、实时性差与激光雷达点云的数据稀疏、扰动、存在噪声等问题,本文首次提出了一种针对“L”型状点云小目标的实时视觉与激光雷达融合体积测量的方法。算法的主要创新点为结合激光雷达数据特点,设计了一种对小目标点云三维框的

粗拟合算法,并提出视觉修正模型对体积进行细修正。实验结果表明,本文提出的算法在一定范围内可以实现“L”型状点云小目标的高精度体积实时测量,可应用于物流与仓储中物品的体积测量、道路上车辆的检测等领域。

但本文算法也存在一定的局限性,例如点云的聚类 and 视觉修正模型都需要依靠于准确的激光雷达到相机的外参,视觉修正模型依靠于“L”型状点云的假设,因此在本文算法中油桶的测量范围比另外两类窄。本文也将算法应用于车这类大目标的体积计算中,遗憾地是尽管三维框能够拟合地很好,但由于单帧数据的视角窄,导致计算的体积偏小,因此多视角数据融合是本文接下来的工作之一。

参考文献

[1] 张又文. 激光雷达电子行业助力自动驾驶汽车产业升级的路径研究[J]. 激光杂志, 2023, 44(5): 215-218.
ZHANG Y W. Research on the path of Lidar electronics industry helping autonomous vehicle industry upgrade[J]. Laser Journal, 2023, 44(5): 215-218.
[2] 樊博璇, 陈桂明, 常亮, 等. 激光雷达技术在军事领域应用现状及发展趋势[J]. 航天制造技术, 2021(3): 66-72.
FAN B X, CHEN G M, CHANG L, et al. The current application status and development trend of LiDAR technology in the military field [J]. Aerospace Manufacturing Technology, 2021 (3): 66-72.

- [3] 牛润新, 张向阳, 王杰, 等. 基于激光雷达的农业机器人果园树干检测算法[J]. 农业机械学报, 2020, 51(11): 21-27.
NIU R X, ZHANG X Y, WANG J, et al. Agricultural robot orchard trunk detection algorithm based on LiDAR[J]. Journal of Agricultural Machinery, 2020, 51 (11): 21-27.
- [4] 郑国庆, 吕品, 赖际舟, 等. 无人车基于激光雷达/视觉的目标体积自动测量方法[J]. 导航定位与授时, 2023, 10(2): 65-73.
ZHENG G Q, LYU P, LAI J ZH, et al. Automatic measurement method for target volume of unmanned vehicles based on LiDAR/vision [J]. Navigation, Positioning and Timing, 2023, 10 (2): 65-73.
- [5] 冯明驰, 高小倩, 汪静姝, 等. 基于立体视觉与激光雷达的车辆目标外形位置融合算法研究[J]. 仪器仪表学报, 2021, 42(10): 210-220.
FENG M CH, GAO X Q, WANG J SH, et al. Research on vehicle target shape and position fusion algorithm based on stereo vision and LiDAR[J]. Chinese Journal of Scientific Instrument, 2021, 42 (10): 210-220.
- [6] 何珍, 楼佩煌, 钱晓明, 等. 多目视觉与激光组合导航 AGV 精确定位技术研究[J]. 仪器仪表学报, 2017, 38(11): 2830-2838.
HE ZH, LOU P H, QIAN X M, et al. Research on precise positioning technology for AGV using multi vision and laser integrated navigation[J]. Chinese Journal of Scientific Instrument, 2017, 38 (11): 2830-2838.
- [7] YUNARDI R T. Contour-based object detection in automatic sorting System for a parcel boxes [C]. International Conference on Advanced Mechatronics, Intelligent Manufacture, and Industrial Automation (ICAMIMIA), IEEE, 2015: 38-41.
- [8] 尹文庆, 浦浩, 胡飞, 等. 基于结构光视觉的联合收获机谷粒体积流量测量方法[J]. 农业机械学报, 2020, 51(9): 101-107.
YIN W Q, PU H, HU F, et al. Measurement method for grain volume flow rate of combine harvesters based on structured light vision [J]. Journal of Agricultural Machinery, 2020, 51 (9): 101-107.
- [9] 吴琪波. 基于双目视觉的特定物体体积测量方案设计[D]. 厦门: 厦门大学, 2021.
WU Q B. Design of a specific object volume measurement scheme based on binocular vision[D]. Xiamen: Xiamen University, 2021.
- [10] SUZUKI T, FUTATSUSHI K, KOBAYASHI K. Food volume estimation using 3D shape approximation for medication management support [C]. 3rd Asia-Pacific Conference on Intelligent Robot Systems (ACIRS), IEEE, 2018: 107-111.
- [11] 张立斌, 吴岛, 单洪颖, 等. 基于激光点云的车辆外廓尺寸动态测量方法[J]. 华南理工大学学报(自然科学版), 2019, 47(3): 61-69.
ZHANG L B, WU D, SHAN H Y, et al. A dynamic measurement method for vehicle outline dimensions based on laser point cloud [J]. Journal of South China University of Technology (Natural Science Edition), 2019, 47 (3): 61-69.
- [12] 崔峥. 基于深度学习点云分割的散料堆体测量方法的研究[D]. 济南: 山东大学, 2023.
CUI ZH. Research on the volume measurement method of bulk material piles based on deep learning point cloud segmentation[D]. Jinan : Shandong University, 2023.
- [13] LANG A H, VORA S, CAESAR H, et al. Pointpillars: Fast encoders for object detection from point clouds[C]. Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, 2019: 12697-12705.
- [14] ZHOU Q, YU C. Point RCNN: An angle-free framework for rotated object detection[J]. Remote Sensing, 2022, 14(11): 2605-2627.
- [15] SHI S, WANG Z, SHI J, et al. From points to parts: 3D object detection from point cloud with part-aware and part-aggregation network [J]. IEEE Transactions on Pattern Analysis and Machine Intelligence, 2020, 43(8): 2647-2664.
- [16] WANG C Y, BOCHKOVSKIY A, LIAO H Y M. YOLOv7: Trainable bag-of-freebies sets new state-of-the-art for real-time object detectors[C]. Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition, 2023: 7464-7475.
- [17] HIMMELSBACH M, HUNDELSHAUSEN F V, WUENSCH H J. Fast segmentation of 3D point clouds for ground vehicles[C]. 2010 IEEE Intelligent Vehicles Symposium, IEEE, 2010: 560-565.
- [18] XU J, ZHANG R, DOU J, et al. Rpnnet: A deep and efficient range-point-voxel fusion network for Lidar point

cloud segmentation [C]. Proceedings of the IEEE/CVF International Conference on Computer Vision, 2021: 16024-16033.

- [19] YE D, CHEN W, ZHOU Z, et al. Lidarmultinet: Unifying lidar semantic segmentation, 3D object detection, and panoptic segmentation in a single multi-task network [J]. Computer Science, 2022, DOI: 10.48550/arXiv.2206.11428.
- [20] AHMED M, SERAJ R, ISLAM S M S. The k-means algorithm: A comprehensive survey and performance evaluation[J]. Electronics, 2020, 9(8): 1295-1306.
- [21] LIU H, SONG R, ZHANG X, et al. Point cloud segmentation based on Euclidean clustering and multi-plane extraction in rugged field [J]. Measurement Science and Technology, 2021, 32(9): 095106.
- [22] KHAN K, REHMAN S U, AZIZ K, et al. DBSCAN: Past, present and future [C]. The Fifth International Conference on the Applications of Digital Information and Web Technologies (ICADIWT 2014), IEEE, 2014: 232-238.
- [23] ZHANG X, XU W, DONG C, et al. Efficient L-shape fitting for vehicle detection using laser scanners[C]. 2017 IEEE Intelligent Vehicles Symposium (IV), IEEE, 2017: 54-59.

作者简介



郭隆强, 2022 年于湖南大学获得学士学位, 现为湖南大学硕士研究生, 主要研究方向为点云处理、尺寸测量。

E-mail: 1653768150@qq.com

Guo Longqiang received his B. Sc. degree from Hunan University in 2022. He is currently a postgraduate at Hunan University. His main research interests include point cloud processing and size measurement.



何贇泽 (通信作者), 2012 年于国防科学技术大学获得博士学位, 现为湖南大学教授, 主要研究方向为嵌入式人工智能与边缘计算、红外热成像与机器视觉。

E-mail: hejicker@163.com

He Yunze (Corresponding author) received his Ph. D. degree from the University of Defense Science and Technology in 2012. He is currently a professor at Hunan University. He is main research interests include embedded artificial intelligence and edge computing, infrared thermal imaging and machine vision.



杜旭, 2021 年于贵州大学获得学士学位, 现为湖南大学硕士研究生, 主要研究方向为图像处理。

E-mail: 2676871308@qq.com

Du Xu received his B. Sc. degree from Guizhou University in 2021. He is currently a master student at Hunan University. His research interest is image processing.