

A Robust Nonparametric Estimation Framework for Implicit Image Models

Himanshu Arora
University of Illinois
Urbana, IL61801, USA
harora1@uiuc.edu

Maneesh Singh
Siemens Corporate Research
Princeton, NJ08540, USA
msingh@scr.siemens.com

Narendra Ahuja
University of Illinois
Urbana, IL61801, USA
n-ahuja@uiuc.edu

Abstract

Robust model fitting is important for computer vision tasks due to the occurrence of multiple model instances, and, unknown nature of noise. The linear errors-in-variables (EIV) model is frequently used in computer vision for model fitting tasks. This paper presents a novel formalism to solve the problem of robust model fitting using the linear EIV framework. We use Parzen windows to estimate the noise density and use a maximum likelihood approach for robust estimation of model parameters. Robustness of the algorithm results from the fact that density estimation helps us admit an a priori unknown multimodal density function and parameter estimation reduces to estimation of the density modes. We also propose a provably convergent iterative algorithm for this task. The algorithm increases the likelihood function at each iteration by solving a generalized eigenproblem. The performance of the proposed algorithm is empirically compared with Least Trimmed Squares(LTS) — a state-of-the-art robust estimation technique, and Total Least Squares(TLS) — the optimal estimator for additive white Gaussian noise. Results for model fitting on real range data are also provided.

1. Introduction

Robust model fitting is central to many computer vision tasks. Examples include tracking or registration under Euclidian, affine, or projective transformations; surface normal and curvature estimation for 3D structure detection; and, fitting intensity models for object recognition and object registration. Robust estimation implies a framework which tolerates the presence of outliers — samples not obeying the relevant model. Consider the problem of segmenting a range image with planar patches: Here, each plane satisfies a linear parametric model. For estimating parameters of each plane, samples from all other planes should be considered as outliers, i.e. they should not contribute to the error in fit. Another scenario is when the noise model for the observed samples is not known. It is not possible to come up with a cost function which is optimal for every kind of (unknown) noise model. Robust estimation seeks to provide reliable

estimates in such cases — when data is contaminated with outliers in form of samples corrupted by unknown noise or when multiple structures are present in the data, some or all of which need to be detected.

Much work has been done in robust estimation in statistics, and more recently in vision. We refer the reader to [13] for a recent review of robust techniques used in computer vision. The two major classes of robust methods proposed in statistics, M-estimators, and least median of squares (LMedS), are regularly used by computer vision researchers to develop applications. M-estimators, a generalization of maximum likelihood estimators and least squares method, were first defined by Huber [6] and their asymptotic properties were studied by Yohai et al. [14], and Koenker et al. [8] in separate works. Least median of squares (LMedS) was proposed by Rousseeuw [10], wherein the sum of squared residuals in traditional least squares is replaced by median of squared residuals. The Hough Transform [7], [9], and RANSAC [4] were independently developed in computer vision community for robust estimation. For Hough Transform, entire parameter space is discretized and optimal parameters are estimated by a voting scheme due to each data sample. It can be viewed as a discrete version of M-estimation. RANSAC [4] uses the number of points with residual below a threshold as the objective function. It has similarities with both M-estimators, and LMedS. Recently, Chen et al. [2] showed that all robust techniques applied to computer vision, i.e. those imported from statistics, and those developed in computer vision literature, can be described as specific instances of the general class of M-estimators with auxiliary scale. In a separate work [1], Chen et al. explore the relationship between M-estimators and kernel density estimators, and propose a technique for robust estimation based on kernel density estimators.

Many parameter estimation problems in computer vision can be formulated with the linear errors-in-variable model (EIV) [15], where the observations are assumed to be corrupted by additive noise. Further, it is often desirable to use implicit functional form. For instance, consider the problem of range image segmentation mentioned earlier. We can extract the 3D world coordinates, (x_i, y_i, z_i) from the range data. If the range measurements, r_i are noisy, (x_i, y_i, z_i)

will be noisy. The linear EIV model can be used to fit a plane through these noisy observations. Further, we should not use any explicit scheme like $z = ax + by + c$, since it does not support the case when the original plane has the equation $ax + by + c = 0$, i.e. a plane perpendicular to z -axis. An implicit scheme is thus essential in this case. The linear EIV model has been used for analysis in some computer vision papers recently [2], [1].

In this work, we assume that the image data may consist of a number of unknown structures, all of which obey the linear EIV model. We also assume that the observed samples are generated by additively corrupting unknown *true* samples with i.i.d. noise. However, the noise model is not available to us. We present a robust estimation algorithm that detects these structures irrespective of the number of structures or the noise model. The robustness is achieved as we use a nonparametric (kernel) estimator to estimate the noise density rather than assuming it to be known a priori. We also prove the convergence of the algorithm under mild conditions on the estimating kernels.

In Section 2, we prove that the parameter estimation problem for the linear EIV model amounts to solving a generalized eigenproblem. We then show in Section 3, that a robust estimation framework can be developed by modeling the pdf of the additive noise using nonparametric kernel density estimators. We then propose an iterative algorithm as a solution to the parameter estimation problem using the ML (maximum likelihood) framework and prove the convergence of this algorithm. In Section 4, we empirically compare the proposed approach with Least Trimmed Squares (LTS), a state-of-the-art robust estimation technique, and Total Least Squares (TLS), which is optimal for gaussian noise. We also present results of model fitting on real data extracted from range images.

2. Linear Errors-in-Variables

The linear errors-in-variables (EIV) approach assumes that the observed samples are generated from the *true* data samples by additively corrupting them by independent, identically distributed (i.i.d.) noise. The true samples obey some linear, functional constraints that capture the a-priori physical nature of the problem. Thus, we define,

Definition 2.1 (Linear EIV model). Let $S_x^o \doteq \{x_{io}\}_{i=1}^n$ be a data sample set of size n satisfying the constraints,

$$f(x_{io}) = x_{io}^T \theta - \alpha = 0 \quad i = 1, \dots, n \quad (1)$$

The observed data sample set $S_x \doteq \{x_i\}_{i=1}^n$ is related to S_x^o by i.i.d. samples from an unknown, additive noise process ϵ such that $x_i = x_{io} + \epsilon_i$. The ambiguity in parameters θ and α is resolved by imposing the constraint $\|\theta\| = 1$.

Consider the case when the noise samples are i.i.d. Gaussian i.e. $\epsilon_i \sim \mathcal{N}(0, \sigma^2 I_p)$. It is well known that the max-

imum likelihood estimate of the parameters and noise free samples is then given by

$$[\hat{\theta}, \hat{\alpha}, \hat{x}_{io}] = \operatorname{argmin}_{\theta, \alpha, x_{io}} \frac{1}{n} \sum_{i=1}^n \|x_i - x_{io}\|^2 \quad (2)$$

subject to the constraints on θ , α , and x_{io} as specified in Definition 2.1. Clearly, in minimization of (2), for fixed values of θ and α , the estimates for noise free samples x_{io} are given by the orthogonal projection of the observed samples x_i onto the hyperplane given by (1), with

$$\min_{x_{io}} \|x_i - x_{io}\| = \|x_i - \hat{x}_{io}\| = |x_i^T \theta - \alpha| \quad (3)$$

This indicates that the minimization can be reduced to just minimizing the sum of squared projections, $\sum_{i=1}^n |x_i^T \theta - \alpha|^2$, with respect to parameters θ and α . The theorem below shows that this is indeed true.

Theorem 2.2. Define $[\tilde{\theta}, \tilde{\alpha}]$ as the total least squares (TLS) solution as below,

$$[\tilde{\theta}, \tilde{\alpha}] \doteq \operatorname{argmin}_{\theta, \alpha} \frac{1}{n} \sum_{i=1}^n |x_i^T \theta - \alpha|^2, \quad \|\theta\| = 1 \quad (4)$$

Then, $[\tilde{\theta}, \tilde{\alpha}] = [\hat{\theta}, \hat{\alpha}]$ where $[\hat{\theta}, \hat{\alpha}]$ are as defined in (2) with the constraints as specified in Definition 2.1.

Proof. For the optimization problem specified by (4), the solution $[\tilde{\theta}, \tilde{\alpha}]$ should satisfy the following equations obtained using the Lagrange multiplier method. These equations are obtained by setting the derivative with respect to θ and α equal to zero.

$$\theta : \quad \sum_{i=1}^n (x_i^T \tilde{\theta} - \tilde{\alpha}) x_i + \lambda \tilde{\theta} = 0 \quad (5)$$

$$\alpha : \quad \sum_{i=1}^n (x_i^T \tilde{\theta} - \tilde{\alpha}) = 0 \quad (6)$$

Similarly, the solution, $[\hat{\theta}, \hat{\alpha}]$, to the problem specified by (2) subject to the constraints in Definition 2.1, should satisfy the following equations:

$$x_{io} : \quad x_i - \hat{x}_{io} = \lambda_i \hat{\theta} \quad i = 1, \dots, n \quad (7)$$

$$\theta : \quad \sum_{i=1}^n \lambda_i \hat{x}_{io} + \gamma \hat{\theta} = 0 \quad (8)$$

$$\alpha : \quad \sum_{i=1}^n \lambda_i = 0 \quad (9)$$

subject to the constraints in Definition 2.1. Now, from (7), we get $\hat{x}_{io} = x_i - \lambda_i \hat{\theta}$. Also, taking the transpose of this equation and post-multiplying $\hat{\theta}$ gives us, $\lambda_i = x_i^T \hat{\theta} - \hat{\alpha}$. These two equations can be used in (8) to get,

$$\sum_{i=1}^n (x_i^T \hat{\theta} - \hat{\alpha}) x_i + (\gamma - \sum_{i=1}^n \lambda_i^2) \hat{\theta} = 0 \quad (10)$$

Substituting the values of λ_i in (9) also gives,

$$\sum_{i=1}^n (x_i^T \hat{\theta} - \hat{\alpha}) = 0 \quad (11)$$

Further, (5) and (6) can be used to show that $\lambda = \sum_{i=1}^n (x_i^T \hat{\theta} - \tilde{\alpha})^2$. Similarly, (7)-(9) imply $\gamma - \sum_{i=1}^n \lambda_i^2 = \sum_{i=1}^n (x_i^T \hat{\theta} - \hat{\alpha})^2$. Thus, solutions to (5)-(6) and (10)-(11) are the same. Hence both problems are equivalent. $\square$

Now, let us examine the problem in (4) more closely. Defining $\beta = (\theta, \alpha)^T$, $\mathcal{A} = \sum_{i=1}^n \begin{pmatrix} x_i x_i^T & x_i \\ x_i^T & 1 \end{pmatrix}$ and $\mathcal{B} = \begin{pmatrix} I_p & 0 \\ 0 & 0 \end{pmatrix}$, where I_p is the $p \times p$ identity matrix, we can rewrite (4) as

$$\hat{\beta} = \underset{\beta}{\operatorname{argmin}} \beta^T \mathcal{A} \beta; \quad \beta^T \mathcal{B} \beta = 1 \quad (12)$$

Solving for β leads to $\mathcal{A}\beta = \lambda \mathcal{B}\beta$, where λ is minimum eigenvalue for the generalized eigenproblem. Thus, the solution is the generalized (minimum) eigenvector of $\mathcal{A}$ with respect to $\mathcal{B}$. Consequently, The Maximum likelihood estimation of the linear EIV model parameters in case of Gaussian noise reduces to a generalized eigenproblem.

The assumption, made above, of Gaussian noise is not always desirable. The model of noise is often unknown, and the estimator proposed above may not be optimal in general. In particular, for heavy tailed distributions (e.g. log-normal distribution as will be discussed in section 4), the approach above might have a really bad performance. Also, the structure that we need to detect (in this case the model that we need to fit) might only be valid locally. For instance, while detecting multiple planer segments in a range image, the model parameters are valid only on (local) segments of data. Since the segmentation is not a priori available, robust estimation becomes important for the discovery of any local models. It tolerates the presence of data samples that do not obey the model that is to be estimated. In the next Section, we propose a principled approach to carry out robust estimation.

3. Robust EIV Estimation Using Parzen Windows

If we know the noise model for the linear EIV problem, then we can use the maximum likelihood approach to estimate the EIV model parameters. However, quite often, we do not have access to such a model. In that case, one can take recourse to estimating the noise density and then applying the maximum likelihood framework. In this section, we present such an approach. We first formulate the problem in terms of a noise density estimate using Parzen windows and subsequently, we propose a solution to the said problem.

Parzen windows or kernel density estimators are a popular non parametric density estimation technique in pattern recognition and computer vision [3]. The Parzen window estimate of the pdf from a given set of data samples can be defined as follows:

Definition 3.1 (Kernel Density Estimator). Let the observed samples $y_i \in \mathcal{R}^p$, $i = 1, \dots, n$ be generated independently from an underlying probability distribution function $f(x)$, $f : \mathcal{R}^p \rightarrow \mathcal{R}^+$. Then the kernel density estimate for f is defined as,

$$\hat{f}(y) \doteq \frac{1}{n \det(H)} \sum_{i=1}^n K(H^{-1}(y - y_i)) \quad (13)$$

where, H is a nonsingular bandwidth matrix, and $K : \mathcal{R}^p \rightarrow \mathcal{R}^+$ is the kernel function with zero mean, unit area, and identity covariance matrix.

The kernel function $K(\cdot)$ used above is often assumed to be rotationally symmetric. We find it convenient to define the profile of this rotationally symmetric kernel as a univariate kernel function $\kappa : \mathcal{R} \rightarrow \mathcal{R}^+$, where $K(y) = c_k \kappa(-\|y\|^2)$, c_k being a normalization constant.

The kernel density estimate of an arbitrary set of data samples can be computed as shown above. However, the above density estimate does not factor in any prior knowledge that one may have of the data. For example, the data might be generated using a parametric model. For such a case, we proposed a zero bias-in-mean kernel estimator in our earlier work [12]. We used this estimator for robust (parameter) estimation where the image is specified using an explicit parametric formulation. In this paper, we adapt the aforementioned approach to define a robust kernel maximum likelihood estimation framework for the EIV model. We draw the reader's attention to the fact that the EIV model is an implicit function formulation unlike our previous work.

Now we explain our approach in terms of noise density estimation: Let us assume that the noise free values x_{io} and the parameters $[\theta, \alpha]$ are known such that the constraints in Definition 2.1 are satisfied. Then, the noise can be estimated as $\epsilon_i = x_i - x_{io}$, with $x_{io}^T \theta = \alpha$, $i = 1, \dots, n$, and $\|\theta\| = 1$. The noise is nothing but the deviation of the observation from the model described by the parameters. The key question is to decide the metric that is to be chosen on these deviations to estimate the model parameters. If the noise density was known, one could easily formulate such a metric using the maximum likelihood framework. However, since the noise density is not known, we take the next best approach — we use Parzen windows to estimate the noise density using Definition 3.1.

For a set of observed data samples $\{x_i\}_{i=1}^n$, the kernel density estimate of noise given the noise free samples $\{x_{io}\}_{i=1}^n$ and parameters $[\theta, \alpha]$ can be written as

$$\hat{f}(\epsilon | \theta, \alpha, x_{io}) = \frac{1}{n} \sum_{i=1}^n K(H^{-1}(\epsilon - (x_i - x_{io}))) \quad (14)$$

under the constraint $x_{io}^T \theta = \alpha$, $\|\theta\| = 1$. Let us define space $\mathcal{S} \doteq \{\Theta \doteq [\theta, \alpha, \{x_{io}\}_{i=1}^n] \mid \|\theta\| = 1, x_{io}^T \theta = \alpha \forall i =$

$1, \dots, n\}$. Then the model parameters $\Theta \in \mathcal{S}$, and assuming the noise to be zero-mean, the maximum likelihood estimate of model parameters is given by

$$\Theta_{ML} = \underset{\Theta \in \mathcal{S}}{\operatorname{argmax}} \hat{f}(0|\Theta) \quad (15)$$

Note that the above definition is not restrictive, since any shift in the ϵ -space can be accounted for by a shift in x_{io} and α . In absence of a disambiguating prior, we assume a zero-mean noise process.

In general, there might be multiple structures in the data, all of which we might need to discover (akin to the Hough transform). Thus, the estimated density function $\hat{f}(0|\Theta)$ might be multimodal. In such a case, one seeks all local minima of the density function. Consequently, we define the parameter estimates as follows,

$$\Theta_{kml} = \underset{\Theta \in \mathcal{S}}{\operatorname{argLmax}} \hat{f}(\epsilon = 0|\Theta) \quad (16)$$

where $\operatorname{argLmax}$ denotes a local maximum. It can be shown that the estimator above is a redescending M-estimator [12].

Θ_{kml} is a solution to a constrained nonlinear program. The local maximum of $\hat{f}(\cdot)$ can be sought in general by gradient ascent. We now propose an iterative algorithm to seek the modes of distribution $\hat{f}(\cdot)$ given a starting point. Under the constraint that the profile of the kernel, $\kappa(\cdot)$, is a convex bounded function, the algorithm is guaranteed to increase the objective function at each iteration and converges to a local maximum.

Since, Θ_{kml} is constrained to lie in the space $\mathcal{S}$, we can define the objective function $q : \mathcal{S} \rightarrow \mathcal{R}^+$ as

$$q(\Theta) = \frac{1}{n} \sum_{i=1}^n \kappa(-\|H^{-1}(x_i - x_{io})\|^2) \quad (17)$$

Now, let the derivative of the profile κ be $\kappa' = g$. Assuming that the initial estimate of Θ is $\Theta^{(0)}$, and using the convexity of the profile, we see that

$$q(\Theta) - q(\Theta^{(0)}) \geq \frac{1}{n} \sum_{i=1}^n g(-\|H^{-1}(x_i - x_{io}^{(0)})\|^2) (\|H^{-1}(x_i - x_{io}^{(0)})\|^2 - \|H^{-1}(x_i - x_{io})\|^2) \quad (18)$$

Defining the weights $w_i = g(-\|H^{-1}(x_i - x_{io}^{(0)})\|^2)$, we seek the next iterate $\Theta^{(1)}$ as the maximizer of right hand side of (18), i.e.,

$$\Theta^{(1)} = \underset{\Theta \in \mathcal{S}}{\operatorname{argmin}} \sum_{i=1}^n w_i (\|H^{-1}(x_i - x_{io})\|^2) \quad (19)$$

The problem in (19) is similar to the ML estimation for i.i.d Gaussian noise samples with identity covariance matrix, as discussed in Section 2. It can be reduced to the following

minimization on the space $\tilde{\mathcal{S}} = \{[\tilde{\theta}, \tilde{\alpha}, \{\tilde{x}_{io}\}_{i=1}^n] \mid \tilde{x}_{io}^T \tilde{\theta} = \tilde{\alpha}, \tilde{\theta}^T H^{-2} \tilde{\theta} = 1\}$

$$\tilde{\Theta}^{(1)} = \underset{\Theta \in \tilde{\mathcal{S}}}{\operatorname{argmin}} \sum_{i=1}^n w_i \|H^{-1}x_i - \tilde{x}_{io}\|^2 \quad (20)$$

$$\text{where } \theta = H^{-1} \tilde{\theta} \alpha = \tilde{\alpha} x_{io} = H \tilde{x}_{io} \quad (21)$$

where $\tilde{\Theta}^{(1)} = [\tilde{\theta}^{(1)}, \tilde{\alpha}^{(1)}, \{\tilde{x}_{io}^{(1)}\}_{i=1}^n]$. By Theorem 2.2, we can write

$$[\tilde{\theta}^{(1)}, \tilde{\alpha}^{(1)}] = \underset{\tilde{\theta}^T H^{-2} \tilde{\theta} = 1}{\operatorname{argmin}} \sum_{i=1}^n w_i ((H^{-1}x_i)^T \theta - \alpha)^2 \quad (22)$$

with $\tilde{x}_{io}$ being equal to the perpendicular projection of $H^{-1}x_i$ onto the plane defined by $[\tilde{\theta}, \tilde{\alpha}]$. Defining

$$\mathcal{A} = \sum_{i=1}^n w_i \begin{pmatrix} H^{-1}x_i x_i^T H^{-1} & -H^{-1}x_i \\ -x_i^T H^{-1} & 1 \end{pmatrix} \text{ and } \mathcal{B} = \begin{pmatrix} H^{-2} & 0 \\ 0 & 0 \end{pmatrix}, \text{ from the discussion in Section 2, we get}$$

$[\tilde{\theta}^T \tilde{\alpha}]$ as the generalized eigenvector of $\mathcal{A}$ with respect to $\mathcal{B}$ corresponding to the minimum eigenvalue. $\Theta^{(1)}$ can thus be estimated using (20).

The above mentioned process is repeated iteratively with new estimates to yield a sequence $\{\Theta^{(n)}\}_{i=1}^\infty$ of parameter estimates, and a sequence $\{q(\Theta^{(n)})\}_{i=1}^\infty$ of function values. Clearly, for $\Theta = \Theta^{(1)}$, the right hand side of (18) is positive, implying that $q(\Theta^{(1)}) \geq q(\Theta^{(0)})$. The sequence $\{q(\Theta^{(n)})\}_{i=1}^\infty$ is thus increasing and bounded above (since κ is bounded), implying that it is convergent.

4. Experiments and Results

First, we empirically compare the performance of the algorithm proposed in Section 3 with (a) the total least squares (TLS) solution, and (b) Least Trimmed Squares (LTS) [11]. TLS, as discussed in Section 2, is the optimal estimator for additive, white Gaussian noise (AWGN) and comparison with TLS shows the comparable performance of our algorithm to the optimal solution in case of often-used AWGN model. To test the robustness, we compare our algorithm with Least Trimmed Squares (LTS) which is a state-of-the-art method for robust regression.

We use line fitting in 2D space as the testbed for our experiment. The true samples (x_{io}, y_{io}) satisfy $ay_{io} + bx_{io} + c = 0$ where $b = c = 1, a = -1$. The true values are set as $x_{io} = \frac{i}{50} - 1$, and $y_{io} = x_{io} + 1, i = 0, \dots, 100$. The data samples (x_i, y_i) are generated by adding uncorrelated noise samples to (x_{io}, y_{io}) . The noise samples are generated from the Gaussian distribution (a standard noise model) and two-sided log-normal distribution (to simulate outliers) with several different variance values.

Figure 1 shows two sample realizations. At each value of variance, we generated 1000 realizations and computed the

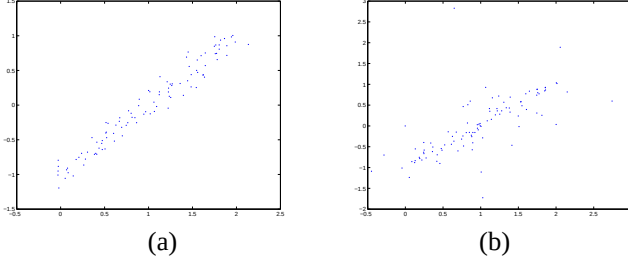


Figure 1: Points generated according to the model $(x_i, y_i) = (x_{io}, y_{io}) + \epsilon_i$ and $\{\epsilon_i\} \in \mathcal{R}^2$ (a) ϵ_i is gaussian with mean 0 and variance 0.09, (b) ϵ_i is log-normal with $M = -4$ and $S = 1.5$.

Table 1: TLS, LTS, and KML estimates of (b, c) for $(x_i, y_i) = (x_{io}, y_{io}) + \epsilon_i$ and $\{\epsilon_i\} \in \mathcal{R}^2$ are i.i.d. Gaussian with mean 0 and variance $\sigma^2 \mathcal{I}_2$. Ground-truth values are $(b, c) = (1, 1)$. Mean and deviation of the estimated values for 1000 experiments are presented in the top and bottom four rows, respectively.

Mean						
σ	TLS		LTS		KML	
	b	c	b	c	b	c
0.03	1.000	1.000	0.999	0.999	1.000	1.000
0.06	1.002	1.000	0.991	1.000	1.003	1.000
0.09	1.001	0.998	0.979	1.000	1.001	0.998
0.12	1.001	1.000	0.963	0.999	1.003	1.001
Standard Deviation						
0.03	0.007	0.004	0.008	0.004	0.007	0.004
0.06	0.015	0.007	0.015	0.009	0.015	0.007
0.09	0.020	0.013	0.023	0.013	0.020	0.013
0.12	0.027	0.016	0.032	0.020	0.029	0.016

means and variances of the estimated parameters. Table 1 shows the results for Gaussian noise. The estimated parameters here are normalized with respect to a since the LTS algorithm is implemented only for explicit function model. The upper and lower halves of the table shows means and variances of the estimated parameters respectively. TLS, LTS, and KML denote Total Least Squares, Least Trimmed Squares, and Kernel Maximum likelihood (proposed algorithm). As we can see, the TLS is the best for this case, i.e. has means closest to 1 and lowest variances, but the performance of our algorithm is comparable. LTS has a bias in estimation and the variances are higher as well. This shows that the proposed algorithm is comparable to LTS, which is the optimal estimator for this case. Table 2 shows the performance for log-normal noise. The table exposes the non robustness of TLS. Its variance blows up as the noise variance increases. Both KML and LTS perform well, with KML being better for higher noise variances. We also note that our algorithm is simpler than LTS and is faster by almost one order of magnitude.

We next demonstrate the ability of the algorithm to detect multiple structures in data: The algorithm was used to esti-

Table 2: TLS, LTS, and KML estimates of (b, c) for $(x_i, y_i) = (x_{io}, y_{io}) + \epsilon_i$ and $\{\epsilon_i\} \in \mathcal{R}^2$ are i.i.d. log-normal with parameter $\mu = -4$ and $S^2 = \sigma^2 \mathcal{I}_2$. Ground-truth values are $(b, c) = (1, 1)$. Mean and deviation of the estimated values for 1000 experiments are presented in the top and bottom four rows, respectively.

Mean						
σ	TLS		LTS		KML	
	b	c	b	c	b	c
0.5	1.000	1.000	0.990	1.000	1.000	1.000
1.0	1.002	1.000	0.976	1.000	1.003	1.000
1.5	1.149	0.996	0.953	1.000	1.001	0.998
2.0	5.314	1.497	0.919	0.996	1.003	1.001
Standard Deviation						
0.5	0.016	0.009	0.019	0.011	0.016	0.009
1.0	0.038	0.019	0.025	0.014	0.025	0.015
1.5	2.087	0.084	0.038	0.019	0.038	0.020
2.0	186.2	24.79	0.065	0.032	0.044	0.024

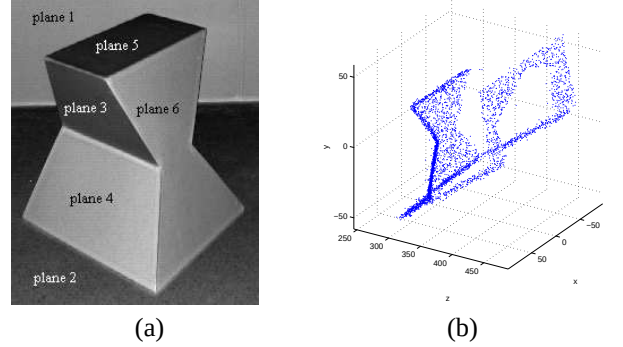


Figure 2: (a) Intensity image from perceptron Ladar USF Range Database (b) Cartesian coordinates extracted from the range data corresponding to (a).

mate plane parameters from 3D data extracted from range images with planer patches. We used the perceptron ladar range images from the USF Range Database [5]. Cartesian coordinates (x_i, y_i, z_i) corresponding to points r_i in the range image are first extracted. Estimation of (all) plane parameters is formulated as a robust *EIV* model parameter estimation problem. The algorithm has following steps: (1) Estimate TLS estimate Θ_p for the parameters for each data point due to points within δ neighborhood of p to provide an initial guess. (2) Arrange the points and corresponding parameter values in decreasing order of likelihood $q(\Theta_p)$ and put on a stack S . (3) Choose the value of parameters from top of the stack, and apply the iterations according to (19) till convergence. Append this value in the estimated parameter list. (4) Remove all points from stack S which are within a perpendicular distance τ from the estimated plane. (5) Repeat steps (3),(4) till all points are exhausted.

The above steps were applied to the data depicted in Fig-

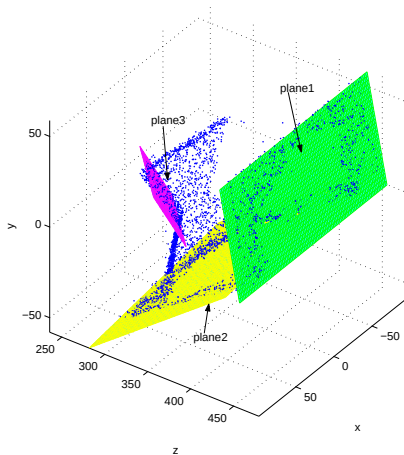


Figure 3: Three estimated planes 1,2 and 3 overlaid on the data samples. The planes are shown in different colors. (rotated version of 2(b) for ease of illustration)

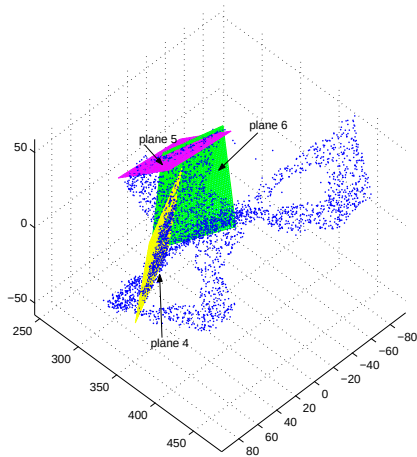


Figure 4: Estimated planes 4,5 and 6 overlaid on the data samples. The planes are shown in different colors. (rotated version of 2(b) for ease of illustration)

ure 2(b). This data was extracted from the range image shown in Figure 2(a). There are six planes to be detected in this image. Figures 3 and 4 show the estimated planes overlaid onto the data points. The algorithm was, thus, able to detect all the instances, along with their parameters, in this multiple structure detection problem .

5. Conclusions

In this paper, we presented a novel formalism to solve the problem of robust model fitting using the linear EIV framework. We used nonparametric density estimation to estimate the unknown noise density and use a maximum likelihood approach for robust estimation of model parameters. The implication of such an approach was that the data could consist of multiple model instances and unknown noise. We also proposed a provably convergent iterative al-

gorithm to solve the resultant optimization problem. The algorithm uses iterative quadratic approximation to the likelihood function based on a variation formulation using the convexity of density kernel functions. The performance of the proposed algorithm favorably compares to two popular algorithms — the Least Trimmed Squares (LTS), and the Total Least Squares (TLS). Results for model fitting on real range data are also provided.

References

- [1] H. Chen and P. Meer. Robust Computer Vision through kernel density estimation. In *Proc. 7th ECCV*, volume 1, pages 236 – 250, 2002.
- [2] H. Chen, P. Meer, and D. E. Tyler. Robust Regression for Data with Multiple Structures. In *CVPR, IEEE Conf.*, volume 1, pages 1069 – 1075, 2001.
- [3] R. O. Duda, P. E. Hart, and D. G. Stork. *Pattern Classification*. Wiley, second edition, 2000.
- [4] M. A. Fischler and R. C. Bolles. Random Sample Consensus: A Paradigm for Model Fitting with Applications to Image Analysis and Automated Cartography. *CACM*, 24(6):381 – 395, June 1981.
- [5] A. Hoover, X. Liang, P. Flynn, H. Bunke, D. Goldgof, K. Bowyer, D. Eggert, A. Fitzgibbon, and R. Fisher. An experimental comparison of range image segmentation algorithms. *PAMI, IEEE Trans.*, 18(7):673 – 689, 1996.
- [6] P. J. Huber. Robust Regression: Asymptotics, Conjectures, and Monte Carlo. *The Annals of Statistics*, 1(5):799 – 821, September 1973.
- [7] J. Illingworth and J. Kittler. A survey of the Hough Transform. *Computer Vision, Graphics, and Image Processing*, 44:87 – 116, 1988.
- [8] R. Koenker and S. Protnoy. M-estimation of multivariate regression. *Journal of Amer. Stat. Assoc.*, 85(412):1060 – 1068, December 1990.
- [9] V. Leavers. Survey: Which Hough transform? *CVGIP*, 58:250 – 264, 1993.
- [10] P. J. Rousseeuw. Least Median of Squares regression. *Journal of Amer. Stat. Assoc.*, 79:871 – 880, 1984.
- [11] P. J. Rousseeuw and K. V. Driessen. Computing LTS regression for large data sets. Technical report, University of Antwerp, December 1990.
- [12] M. Singh, H. Arora, and N. Ahuja. A Robust Probabilistic Estimation Framework for Parametric Image Models. In *8th ECCV, Proceedings, Part I*, pages 508 – 522.
- [13] C. V. Stewart. Robust Parameter Estimation in Computer Vision. *SIAM Reviews*, 41(3):513 – 537, 1999.
- [14] V. J. Yohai and R. A. Maronna. Asymptotic behaviour of M-estimators for the linear model. *The Annals of Statistics*, 7(2):258 – 268, March 1979.
- [15] R. H. Zamar. Robust estimation in the errors in variables model. *Biometrika*, 76:149 – 160, 1989.