

VISUAL FEEDBACK BASED OBJECT DETECTION AND MANIPULATION

SHRADDHA CHAUDHARY



**DEPARTMENT OF ELECTRICAL ENGINEERING
INDIAN INSTITUTE OF TECHNOLOGY DELHI**

OCTOBER 2018

VISUAL FEEDBACK BASED OBJECT DETECTION AND MANIPULATION

by

SHRADDHA CHAUDHARY

DEPARTMENT OF ELECTRICAL ENGINEERING

Submitted

in fulfillment of the requirements of the degree of Doctor of Philosophy

to the



INDIAN INSTITUTE OF TECHNOLOGY DELHI

OCTOBER 2018

Dedicated to
My Teachers and Family

Certificate

This is to certify that the thesis titled "**Visual Feedback based Object Detection and Manipulation**", submitted by **Ms. Shraddha Chaudhary** to the **Department of Electrical Engineering**, Indian Institute of Technology Delhi, for the award of **Doctor of Philosophy** is a record of bona-fide research work carried out by her under my guidance and supervision. In my opinion, the thesis has reached the standards fulfilling the requirements of the regulations relating to the degree.

The results contained in this thesis has not been submitted elsewhere, either in part or full, for the award of any other degree or diploma.

Prof. Sumantra Dutta Roy

Professor

Department of Electrical Engineering

Indian Institute of Technology Delhi

New Delhi - 110016, India.

Acknowledgements

Working on the Ph.D has been a wonderful and overwhelming experience. In any case, I am thankful to many people for making the time of working on my Ph.D. an forgettable experience. I would like to express my sincere gratitude to my advisor Prof. Sumantra Dutta Roy. This thesis would not have been possible without his guidance, critical review and encouragement. His intelligent ideas and comprehensive understanding helped me to establish the overall direction of the research, got me past many significant challenges and always motivated me to strive for excellence. I would like to express my sincere thanks to the Bhaba Atomic Research Centre (BARC) Mumbai for its financial aid during my Ph.D. This gave me the opportunity to work on the project Vision Guided Control of a Robotic Manipulator in the Programme for Autonomous Robotics Lab. Thanks are also due to the IRD section IIT Delhi, for processing the financial support to attend the conferences whenever required.

I thank the members of my thesis committee: Prof. S.D Joshi, Prof. Prem Kalra and Dr. Brejesh Lall for their insightful suggestions which encouraged me to widen my research from various perspectives. I am also grateful to all the professors at IIT Delhi from whom I have learned a lot during the courses.

I sincerely admire the contribution of my lab mates Mr.Riby Abraham Bobby, Mr. Ravi Prakash Joshi, Mr. Aamir Hayat, Mr. Anil Patidar and Ms. Tanya Raghuvanshi for extending their unstinting support and sympathetic attitude during the course of this study. I feel lucky to have so many incredible friends around. I thank each and every one of them for their trust and friendship.

I acknowledge the significant contribution of my parents in law. Their patience and support

will always be my inspiration. I would like to express my heartiest gratitude to my children, Master Pratyush and Master Prithvi for their love, and my indebtedness towards my husband, Piyush Arora, for his indefatigable assistance. Their infallible love and support has always been my strength. My Parents, Avinash Kumar Chaudhary and Shanti Chudhary, has always been a reassuring rock solid support for all my endeavours.

I thank the High Performance Computing facility at IIT Delhi for providing a multi-GPU enabled cloud computing platform which made training of deep models possible.

Shraddha Chaudhary

Abstract

Object detection along with pose estimation remains a quintessential computer vision problem. By pose estimation, we imply inferring the orientation by taking a 2D image and 3D point cloud of the object. This problem is motivated by the use of vision-based robotic manipulators. These are the robotic arms which rely at least in part, on cameras to obtain images of their surroundings. The thesis aims to provide an automated solution to the pick-and-place problem. Along with this, we have attempted subproblems of this broad scope objective, such as generalized object detection and pose estimation for different geometrical shapes using deep learning techniques.

Regarding the methodology, the first contribution proposes an autonomous machine vision system which grasps a texture-less object from a single layer of cluttered featureless objects of varied sizes, rearranges it for proper placement and then places it at a proper position using computer vision. It contributes to a unique vision-based pose estimation algorithm, collision-free path planning and a dynamic changeover algorithm for final placement. This framework has a well-defined application for the assembly task in industries. The second contribution of this thesis is the use of a low-resolution 3-D scan to guide a localized image search, to optimally utilize information from multiple sensors. This information guides segmentation of the image for better object pose estimation, for a pick-and-drop. We apply this algorithm for emptying a bin where cylindrical texture-less objects are lying in multi-layered clutter. The third contribution models uncertainty of the autonomous system proposed in the previous contribution. Application of the pick-and-place task using a 5-DOF Kuka YouBot is demonstrated in the thesis. The fourth contribution widens the scope of the thesis by working on benchmarked

datasets and moving more towards a robust system, by taking into account different viewpoints and refining the search process. We formulate a novel object classification algorithm with the assistance of viewpoint using an SVM. The fifth part extends the previous task using a deep neural architecture. We show encouraging results for object recognition, object instance prediction and object pose estimation, with the help of an Inception-based convolution neural networks. This architecture requires only RGB images instead of RGB-D required by previous models. We show an improvement with respect to state-of-the-art object detection models: we achieve an accuracy of 89.8%. This is the first work to use an Inception architecture along with 1×1 filters, for object pose estimation.

The thesis presents architectural features, training procedures and shares futuristic recommendations. The proposed algorithms and their extensions, and techniques have been tested on publicly available *de facto* ‘standard’ datasets, such as CIFAR and SVHN. We show encouraging results of the same, indicating that the hybrid architecture can culminate in a very competitive and robust model for pose determination of generalized objects.

सार

पॉज़ अनुमान के साथ ऑब्जेक्ट डिटेक्शन एक उत्कृष्ट कंप्यूटर दृष्टि समस्या बनी हुई है। अनुमान लगाकर, हम ऑब्जेक्ट की 2-D छवि और 3-D पॉइंट क्लाउड ले कर ओरिएंटेशन का उल्लंघन करते हैं। यह समस्या दृष्टि-आधारित रोबोट मैनिपुलेटर्स के उपयोग से प्रेरित है। ये रोबोटिक हथियार हैं जो कैमरे पर अपने आसपास के चित्रों को प्राप्त करने के लिए कम से कम भाग में भरोसा करते हैं। थीसिस का लक्ष्य पिक-एंड-प्लेस समस्या के लिए एक स्वचालित समाधान प्रदान करना है। इसके साथ-साथ, हमने इस व्यापक दायरे के उद्देश्य की उप-समस्याओं का प्रयास किया है, जैसे कि सामान्यीकृत ऑब्जेक्ट डिटेक्शन और गहरी सीखने की तकनीक का उपयोग करके विभिन्न ज्यामितीय आकारों के लिए अनुमान लगाया गया है।

पद्धति के बारे में, पहला योगदान एक स्वायत्त मशीन दृष्टि प्रणाली का प्रस्ताव करता है जो विभिन्न आकारों के अव्यवस्थित फीचरलेस ऑब्जेक्ट्स की एक परत से एक बनावट-कम वस्तु को पकड़ता है, इसे उचित प्लेसमेंट के लिए पुनर्व्यवस्थित करता है और फिर इसे कंप्यूटर दृष्टि का उपयोग करके उचित स्थिति में रखता है। यह एक अद्वितीय दृष्टि-आधारित मुद्रा अनुमान एल्गोरिदम, टकराव मुक्त पथ योजना और अंतिम प्लेसमेंट के लिए गतिशील परिवर्तन एल्गोरिदम में योगदान देता है। इस ढांचे में उद्योगों में असेंबली कार्य के लिए एक अच्छी तरह से परिभाषित आवेदन है। इस थीसिस का दूसरा योगदान एकाधिक सेंसर से जानकारी का बेहतर उपयोग करने के लिए स्थानीयकृत छवि खोज को मार्गदर्शन करने के लिए निम्न-रिज़ॉल्यूशन 3-डी स्कैन का उपयोग है। यह जानकारी एक पिक-एंड-ड्रॉप के लिए बेहतर ऑब्जेक्ट पॉज़ अनुमान के लिए छवि का विभाजन करती है। हम एक बिन खाली करने के लिए इस एल्गोरिदम को लागू करते हैं जहां बहु-स्तरित अव्यवस्था में बेलनाकार बनावट-कम वस्तुएं झूठ बोल रही हैं। तीसरे योगदान मॉडल पिछले योगदान में प्रस्तावित स्वायत्त प्रणाली की अनिश्चितता।

5-डीओएफ(DOF) कुका यूबीओटी(YouBOT) का उपयोग करके पिक-एंड-प्लेस कार्य का उपयोग थीसिस में प्रदर्शित किया गया है। चौथा योगदान विभिन्न दृष्टिकोणों को ध्यान में रखते हुए और खोज प्रक्रिया को परिष्कृत करके, बेंचमार्क किए गए डेटासेट पर काम करके और एक मजबूत प्रणाली की ओर बढ़कर थीसिस के दायरे को बढ़ाता है। हम एक एसवीएम का उपयोग करके दृष्टिकोण की सहायता से एक उपन्यास ऑब्जेक्ट वर्गीकरण एल्गोरिदम तैयार करते हैं। पांचवां भाग एक गहरे तंत्रिका वास्तुकला का उपयोग कर पिछले कार्य को बढ़ाता है। हम एक प्रारंभ-आधारित संकल्प तंत्रिका नेटवर्क की सहायता से ऑब्जेक्ट मान्यता, ऑब्जेक्ट इंस्टेंस भविष्यवाणी और ऑब्जेक्ट पॉज़ अनुमान के लिए उत्साहजनक परिणाम दिखाते

हैं। इस आर्किटेक्चर के लिए पिछले मॉडल द्वारा आवश्यक आरजीबी-डी(RGB-D) की बजाय केवल आरजीबी(RGB) छवियों की आवश्यकता है। हम अत्याधुनिक ऑब्जेक्ट डिटेक्शन मॉडल के संबंध में सुधार दिखाते हैं: हम 89.8% की शुद्धता प्राप्त करते हैं। ऑब्जेक्ट पॉज़ अनुमान के लिए, 1X1 फ़िल्टर के साथ एक इंसेप्शन आर्किटेक्चर का उपयोग करने वाला यह पहला काम है।

थीसिस वास्तुशिल्प सुविधाओं, प्रशिक्षण प्रक्रियाओं और भविष्य की सिफारिशों के शेयर प्रस्तुत करता है। प्रस्तावित एल्गोरिदम और उनके एक्सटेंशन, और तकनीकों का परीक्षण सार्वजनिक रूप से उपलब्ध 'मानक' डेटासेट्स जैसे सीआईएफएआर(CIFAR) और एसवीएनएन(SVHN) पर किया गया है। हम इसके बारे में उत्साहजनक परिणाम दिखाते हैं, यह दर्शाते हुए कि हाइब्रिड आर्किटेक्चर सामान्यीकृत वस्तुओं के निर्धारण के लिए एक बहुत प्रतिस्पर्धी और मजबूत मॉडल में समाप्त हो सकता है।

Contents

Certificate	iii
Acknowledgements	v
Abstract	vii
List of Figures	xv
List of Tables	xxi
1 Introduction	1
1.1 Scope and Objectives	2
1.2 Major Contributions of the Thesis	7
1.2.1 Overview and layout	8
1.3 Summary	11
2 Bin Picking, Pose Estimation and its Applications: A Review	13
2.1 Introduction	13
2.2 Bin Picking Problem and its Various Solutions	13
2.2.1 Vision based solution for bin picking	15
2.2.2 RANSAC based algorithms for bin picking	16
2.2.3 3D sensors based solution for bin picking	16
2.3 Uncertainty Modeling	18

2.4	Machine Learning in Visual Servoing	19
2.5	Pose Estimation	20
2.5.1	Handcrafted features	20
2.5.2	Deep learning approaches	20
2.6	State of the Art	22
2.6.1	Gap addressed from current literature	23
3	Pose Estimation of the Cylindrical Pellets in Occluded Single Layer, and its Assembly	25
3.1	Introduction	25
3.1.1	Overview of the system	26
3.2	Pose Estimation	28
3.2.1	Position estimation of isolated pellets	28
3.2.2	Position estimation of occluded pellets	30
3.3	Path Planning and Rearrangement	33
3.3.1	Rearrangement/Changeover	34
3.4	Placement	37
3.4.1	Bin dropping:	37
3.4.2	Grid making:	38
3.4.3	Hole insertion:	38
3.5	Results	40
3.5.1	Multiple stages in pellet pick up	40
3.5.2	Reliability of the algorithm	41
3.5.3	Error in the estimated orientation	41
3.6	Summary	42
4	Pose Estimation of Pellets in Multilayer Configuration using Sensor Fusion	43
4.1	Introduction	43
4.1.1	Need for depth sensor fusion with Camera	44

4.1.2	Set up and work flow	45
4.2	Calibration and Sensor Fusion	47
4.2.1	Transformation between tool and laser coordinate systems	48
4.2.2	Transformation between camera and laser scanner frames	50
4.3	Segmentation and Pose Estimation	52
4.3.1	3-D Segmentation	52
4.3.2	Mathematical formulation for pose estimation	53
4.3.3	Image segmentation	56
4.4	Graspability and Approachability	58
4.4.1	Collision avoidance	59
4.4.2	Disturbance detection	61
4.4.3	Oriented placing	62
4.4.4	Shake criterion	62
4.5	Experiments and Results	62
4.5.1	Segmentation and pose estimation	63
4.5.2	System repeatability	64
4.6	Summary	65
5	Uncertainty Modeling of Deterministic Robotic Picking System using Sensor Fusion	67
5.1	Introduction	67
5.1.1	Uncertainty propagation	69
5.1.2	Uncertainty ellipsoid	71
5.2	Picking System Using Sensor Fusion	71
5.2.1	Disturbance detection using sensor fusion	72
5.3	Uncertainty During Sensor Fusion Calibration	75
5.3.1	Uncertainty in robot	76
5.3.2	Uncertainty in camera calibration	78

5.3.3	Uncertainty during laser calibration	82
5.3.4	Uncertainty propagation to end effector coordinate system	83
5.4	Experimental Results	85
5.4.1	Effect of uncertainty in camera calibration matrix	86
5.4.2	Effect of uncertainty in laser frame	89
5.4.3	Effect of uncertainty propagated to end effector/tool frame	92
5.5	Summary	94
6	Application: Pick and Place task with KUKA YouBot	95
6.1	Navigation and Object Recognition	96
6.2	Methodology	96
6.2.1	3D Map Generation and Robot Localization	97
6.2.2	Object Recognition	97
6.2.3	Motion Planning	98
6.2.4	Robust Object Recognition Algorithm	99
6.3	Pose Estimation using Monocular camera mounted on the YouBot arm	101
6.3.1	Stereo Camera Calibration	101
6.3.2	Depth estimation	101
6.3.3	Grasp Detection	102
6.4	Object Pick and place using Image based Visual Servoing	104
6.5	Experimental results	105
6.5.1	Grasp Detection	105
6.5.2	Machine Learning Approach to Object Picking	107
6.6	Summary	110
7	A Hybrid Model For Pose Estimation	113
7.1	Introduction	113
7.1.1	Dataset	114
7.1.2	Problem Analysis	115

7.2	An Overview of an Approach Without Using Deep Learning	115
7.2.1	Feature extraction	116
7.2.2	Best view estimation	117
7.2.3	View point assisted object classification	118
7.2.4	Object classification and viewpoint estimation accuracy using proposed approach	119
7.2.5	Need for deep learning based approach	121
7.3	The Proposed Model	121
7.3.1	Extracting Features From An Inception Model	121
7.3.2	Three Types Of Specialized Neural Networks For Prediction	124
7.4	Evaluation Criteria	128
7.5	Results	128
7.5.1	Object And Instance Recognition	131
7.5.2	Pose Prediction	131
7.6	Summary	134
8	Supervised Generative Modeling	135
8.1	Generative Models	135
8.2	Generative Adversarial Networks	136
8.3	A Novel Supervised Generative Network Using GANs As Backbone	137
8.3.1	Discriminator	137
8.3.2	Generator	137
8.3.3	Architecture	137
8.3.4	Loss Functions	140
8.3.5	Training algorithm	140
8.3.6	Prediction algorithm	140
8.4	Experimental Results	141
8.5	Summary	144

9 Conclusion	145
9.1 Summary of the contributions	145
9.2 Future Scope	147
Bibliography	149
Publications	163
Biography	165

List of Figures

1.1	The motivation for pose estimation is augmenting robotic vision to better interact with the environment which could lead to more robust manufacturing.[1] . . .	2
1.2	Experimental Setup with information fused setup	4
1.3	Work flow of the Thesis. (Rectangles in orange color in the diagram highlights work done in the thesis.)	5
2.1	Gabor like filters typical of initial layers of CNNs.[2]	21
3.1	Standing pellet	27
3.2	Sleeping pellet	27
3.3	Slant pellet	27
3.4	Zhang Calibration Equation	28
3.5	SVD orientation	29
3.6	Standing Pellet Detection using Hough Circles	29
3.7	Principal Components(a,b,c,d) 1st, 2nd, 3rd, 4th	30
3.8	Actual-Error Analysis	30
3.9	Actual-Prediction Analysis	30
3.10	Results for sleeping pellets	32
3.11	Wrong results	32
3.12	Results after hypothesis fix	32
3.13	(a,b,c) Result Using Heuristic Approach	32
3.14	Final image with positions	33

3.15	Segmentation of the same image	33
3.16	Nomenclature for sections around pellet	34
3.17	Placing for Standing Pellet	36
3.18	Pick-up for sleeping pellet	36
3.19	Placing for Sleeping Pellet	36
3.20	Pick-up for standing pellet	36
3.21	Bin dropping	37
3.22	Grid making with 2nd pellet	38
3.23	Grid Making with 3rd pellet	38
3.24	Holes detected using hough circles	39
3.25	Grippers approach for pellet hole insertion	39
3.26	Pellet Insertion in Hole by Manipulator	39
3.27	Multiple Stages in pellet pickup	40
3.28	Success rate for hole insertion	41
3.29	Error Plot for Orientation	42
4.1	The pellets are kept in a pile, lacking any specific order.	44
4.2	Flowchart for the pose estimation in Multilayered configuration	46
4.3	Set up with sensor fusion	46
4.4	Laser scanner frame and robot Tool frame: (a) Schematic, (b) Experimental Set-up.	48
4.5	Laser scan, (a) Profile of scan, (b) Schematic of 3 points	48
4.6	(a) Projection on calibration grid, (b) Back projection on objects in bin.	52
4.7	(a) Data Acquired using Laser scanner, (b) Filtered data.	54
4.8	3D segmentation a) Region growing in Voxel, b) Normals estimation, c) Axis calculated, d) Cylindrical cluster projected onto the plane axis.	55
4.9	a)Image for Aligned Cylinders, b)Segmentation Results	58

4.10	After extracting cluster from 3D a) shows ROI and b) show segmentation result in standing pellet case.	58
4.11	(a)XY Plane (Top View),(b) Cylinder-axis and Z-axis plane (Side View)	60
4.12	(a)Pixel Difference Observed without Disturbance,(b) Disturbance Detection. .	62
4.13	Oriented Placement: Standing pellets are placed in the queue	63
4.14	Oriented Placement: Sleeping pellets are placed on the rail guide	63
4.15	Green color represents of surfaces as classified by 3D analysis. Red colour is associated when results from 3D segmentation are non-conclusive.	64
5.1	(a) 3D scan points of bin	72
5.2	(b) projection of 3D points on image	72
5.3	Flowchart for Disturbance Detection using sensor fusion for bin picking system	73
5.4	Disturbance detection. a) Initial Configuration and b) Image after disturbance .	73
5.5	Subtracted image, difference shown in grey.	74
5.6	Effect of various uncertainties on the final pose estimates	74
5.7	Transformation from Grid coordinate system to Robot end-effector coordinate system	77
5.8	Flowchart for Disturbance Detection using sensor fusion for bin picking system	79
5.9	Experimental procedure for uncertainty in back projected world coordinates	80
5.10	3D checker board pattern used for calibration	81
5.11	3D checker board pattern used for calibration	86
5.12	Probability Density Function of Laser Uncertainty	89
5.13	Triangulation Principle	90
5.14	Effect of Distance from Laser Sensor	91
5.15	Error Ellipsoid for the uncertainty in the Tool/End-effector frame	92
6.1	Experimental Setup for YouBot Navigation	97
6.2	Kinect RGB Image	97
6.3	Kinect Depth Image	97

6.4	COIL-100 Dataset Images	100
6.5	Real Dataset images	100
6.6	2D Map generation in Gazebo simulator	100
6.7	Object Recognition results	101
6.8	Stereo camera calibration using MATLAB app	102
6.9	Objects in Cornell grasping dataset	103
6.10	Result of Robotic Grasp Detection Algorithm	103
6.11	(a) Spectacle View 1, (b) Spectacle View 2.	106
6.12	(a) Remote View 1, (b) Remote View 2.	106
6.13	(a) Standing Pellet View, (b) Sleeping Pellet View.	106
6.14	Mean Square Error and R value for the NN model	107
6.15	Real Images Grasp Point	108
6.16	Regression plots for Neural Network model	108
6.17	Results of the Support Vector Regression Model	109
7.1	Sample of images from the dataset. The dataset consists of a wide variety of objects which makes the problem even more challenging.	115
7.2	Histogram encoded DAISY features for sub images of size 48 x 48	117
7.3	Configuration 1 : Assignment of view points (All coloured azimuth angles ex- cept black are assigned corresponding view points)	118
7.4	Orientation of object category "car" for the corresponding view point assign- ment	118
7.5	Joint classification confusion matrix corresponding to Fig 7.3	120
7.6	Configuration 2 : Assignment of view points (All colored azimuth angles except black are assigned corresponding view points)	120
7.7	Schematic view of an Inception module[3].	122

7.8	Visualizations of activations of an Inception network[3] as an image of a ‘plier’ is passed through it. The parallel branches of the Inception modules are called ‘towers’.	123
7.9	Schematic diagram of the network used for object recognition. Number of hidden neurons per layer are denoted in [] brackets. Fully connected layers with ReLU activations are used throughout.	124
7.10	Schematic diagram of the network used for object instance prediction. [Number of hidden neurons per layer] in brackets. Fully connected layers with ReLU activations are used throughout.	125
7.11	Schematic diagram of the network used for pose prediction. [Number of hidden neurons per layer] in brackets. Fully connected layers with ReLU activations are shown in gray color.	126
7.12	Schematic diagram of the entire hybrid model.	127
7.13	Plot of the t-SNE reduced ‘deep’ features. We represent different objects with different colors here.	129
7.14	Plot of the t-SNE reduced ‘shallow’ and ‘deep’ features for a fixed object but with varying poses. We clearly observe a good amount of variation with pose in both type of features.	129
7.15	Visual demonstration of how 1×1 filters take raw ‘shallow’ features and convert them into more meaningful ones while reducing feature space dimensionality.	130
7.16	A consolidated comparison of pose prediction performance for different types of input features for all classes of objects.	133
8.1	A way to visualize deconvolution/transposed convolution process. Blue is the input while green is the output. Shown here is a mechanism to apply zero padded input to achieve an output with greater spatial extent([4]).	139
8.2	Images from the actual MNIST dataset.	141
8.3	Images generated by SGN for MNIST. These are not from the dataset.	142

8.4	Images from the actual SVHN dataset.	142
8.5	Images generated by SGN for SVHN. These are not from the dataset.	143
8.6	Images from the actual CIFAR-10 dataset.	143
8.7	Images generated by SGN for CIFAR-10. These are not from the dataset. . . .	144

List of Tables

3.1	Algorithm Accuracy	40
4.1	Timing and Efficiency as compared to state of the art algorithms	65
5.1	Uncertainty ellipses for a single corner point in world frame with different noise standard deviation	87
5.2	Uncertainty ellipses for different corner points in world frame	88
5.3	standard deviation of X and Z for for single corner point with different noise deviations	89
5.4	Effect of Distance from Laser Sensor	91
5.5	Effect of Speed of Robotic Manipulator on Laser Sensor	91
5.6	Uncertainty ellipses for a single corner point Tool/End-Effector frame with different noise standard deviation	93
7.1	We compare our model with some state of the art models for object recognition. Our model uses only RGB data but still beats models using RGB-D data.	132
7.2	Comparison with some state of the art models for instance prediction showing that our model using only RGB data can outperform models using RGB-D data.	132

7.3	We follow the evaluation criteria in [5] for consistency, MedPose error is the error in pose expressed in median of the test set, a penalty of 180° is imposed if object is wrongly classified in terms of either the class or the instance. MedPose(C) only includes the error when the object class is correctly predicted, here also a penalty of 180° is imposed for every incorrect prediction of instance. MedPose(I) only includes the cases where both class and instance are correctly identified.	133
-----	---	-----