



Call for Papers

Joint Workshops of VCL'09 and ViSU'09



June 25, 2009, Fontainebleau Resort, Miami Beach, Florida

VCL'09: Workshop on Visual and Contextual Learning from Annotated Images and Videos

www.umiacs.umd.edu/conferences/vclworkshop09/

Abhinav Gupta (U Maryland)
Jianbo Shi (U Penn)
David Forsyth (UIUC)

There has been a significant interest in the computer vision community on utilizing visual and contextual models for high level semantic reasoning. Large datasets of weakly annotated images and videos, along with other rich sources of information such as dictionaries, can be used to learn visual and contextual models for recognition.

The goal of this workshop is to investigate how linguistic information available in the form of captions and other sources can be used to aid in visual and contextual learning. For example, nouns in captions can help train object detectors; adjectives in captions could be used to train material recognizers; or written descriptions of objects could be used to train object recognizers. This workshop aims to bring together researchers in the fields of contextual modeling in computer vision, machine learning and natural language processing to explore a variety of perspectives on how these annotated datasets can be employed.

Scope: The list of possible topics includes (but is not limited to) the following:

- **Contextual Relationships for Recognition**
 - Scene, Object, Action and Event Recognition using context models
 - Learning structure and parameters of contextual models
- **Linguistic annotations to Assist Learning and Incremental Labeling**
 - Annotations for learning visual and contextual models
 - Rich language models of annotations.
 - Learning to recognize by reading
 - Utilizing other sources such as dictionaries
 - Modeling annotations with errors
- **Others**
 - Biologically motivated semantic models
 - Scalable learning from large datasets

ViSU'09: 1st International Workshop on Visual Scene Understanding

<http://vision.cs.princeton.edu/vcl-visu2009/ViSU>

Fei-Fei Li (Princeton)
Serge Belongie (UCSD)

One of the holy grails of computer vision research is achieving the total understanding of a visual scene, in a similar way that humans can do effortlessly. Great progress has been made in tackling one of the most critical components of visual scene understanding: object recognition and categorization. Recently, a few pioneering works have begun to look into the interactions among objects and their environments towards the goal of a more holistic representation and explanation of the visual scene. This workshop offers an opportunity to bring together experts working on different aspects of scene understanding and image interpretation and provides a common playground for a stimulating debate.

The Representation and Recognition Aspect

- What is total scene understanding? What defines *objects, things, stuff, context, scenes* and other high-level visual concepts (e.g., *activities and events*)?
- How is this problem related to the 'image annotation' problem from the CBIR/TRECVID community?
- Can we break down the task into individual component recognition? If yes, how? If no, why?

The Learning Aspect

- What are the challenges faced in learning for these problems? Are they the same or different from isolated object recognition?
- What can we do to exploit the huge amount of data on the web?

New Methods for Evaluation and Benchmarking

- Are the object recognition datasets still suitable for this problem? If not, what are the new datasets? What are the appropriate benchmark tasks and evaluation metrics?

Paper topics may include total scene understanding, visual recognition and annotation in complex real-world images, high-level understanding (events, activities) in single images, and all their related learning, representational, and dataset issues.

Important dates (same for both workshops):

Deadline for paper submission March 20, 2009
Notification of acceptance April 8, 2009
Camera-ready copies due April 14, 2009

Submission and Reviews

a) General instruction

Submitted papers must have a maximum length of 8 pages and must adhere to the same CVPR 2009 layout specifications as papers submitted to the main conference. All reviewing will be carried out double-blind by the Program Committee. Please refer to CVPR 2009 website (<http://www.cvpr2009.org/>) for instructions on the submission format. In submitting a manuscript, authors acknowledge that no paper substantially similar in content has been or will be submitted to another conference or workshop during the review period.

b) Submitting to VCL vs. ViSU

In the submission form, the authors will be asked to indicate which workshop (VCL or ViSU) the authors prefer to submit to. If no preference is indicated, the VCL-ViSU program chairs will make the decision.

Program Committees**VCL'09:**

Kobus Barnard (Arizona), Serge Belongie (UCSD), Alex Berg (Yahoo Research), Tamara Berg (SUNY Stony Brook), David Blei (Princeton), Larry Davis (Maryland), Sven Dickinson (Toronto), Pinar Duygulu (Bilkent U.), Alexei A. Efros (CMU), Mark Everingham (Leeds), Kristen Grauman (UT Austin), James Hays (CMU), Derek Hoiem (UIUC), Julia Hockenmaier (UIUC), Sanjiv Kumar (Google), Svetlana Lazebnik (UNC), Fei-Fei Li (Princeton), Ivan Laptev (INRIA), R. Manmatha (U Mass), Fernando Pereira (U Penn), Cordelia Schmid (INRIA), Ben Taskar (U Penn), Antonio Torralba (MIT), Tinne Tuytelaars (KU Leuven), Nuno Vasconcelos (UCSD), James Z. Wang (Penn State), Andrew Zisserman (Oxford)

ViSU'09:

Kobus Barnard (Arizona), Miguel Á. Carreira-Perpiñán (UC Merced), Tsuhan Chen (CMU), Trevor Darrell (Berkeley), Larry Davis (UMaryland), Pinar Duygulu (Bilkent U.), Alyosha Efros (CMU), Mark Everingham (Leeds), Pedro Felzenszwalb (U.Chicago), Vittorio Ferrari (ETHZ), David Forsyth (UIUC), Bill Freeman (MIT), Kristen Grauman (U.T. Austin), Abhinav Gupta (UMaryland), Jeremy Heitz (Stanford), Derek Hoiem (UIUC), Xian-Sheng Hua (MSRA), Dan Huttenlocher (Cornell), Frédéric Jurie (INRIA), Daphne Koller (Stanford), Lana Lazebnik (UNC), Bastian Leibe (Aachen), Jia Li (Princeton), Jitendra Malik (Berkeley), Greg Mori (SFU), Pietro Perona (Caltech), Andrew Rabinovich (Google), Deva Ramanan (UCI), Brian Russell (ENS), Silvio Savarese (UMich), Cordelia Schmid (INRIA), Erik Sudderth (Brown), Sinisa Todorovic (Oregon State), Antonio Torralba (MIT), Zhuowen Tu (UCLA), Nuno Vasconcelos (UCSD), John Winn (MSRC), Song-Chun Zhu (UCLA), Andrew Zisserman (Oxford)