



IMAGE REGION ASAQUADTREE HYPER GRAPH

A. Panneerselvam* & M. Amudha**

* Professor, Department of Mathematics, PRIST University, Tanjore, Tamilnadu

** Research Scholar, Department of Mathematics, PRIST University,
Tanjore, Tamilnadu

Cite This Article: A. Panneerselvam & M. Amudha, "Image Region Asaquadtrees Hyper Graph", International Journal of Applied and Advanced Scientific Research, Volume 3, Issue 2, Page Number 8-12, 2018.

Abstract:

A suitable image illustration induces some smart image algorithms. Hyper graph theory could be a theory of restricted combinatorial sets, modeling a lot of issues of operational analysis and combinatorial improvement. Hypergraphs are at this time employed in several domains like chemistry, engineering and image process. We tend to gift a summary of a hyper graph-based image illustration giving much application in image manipulation nevertheless, the current hyper graph learning methods face issues, i.e., the way to produce hyper edges and the way to handle a large set of hyper edges. We present an summary of a hyper graph-based Image illustration and also the Image region quad tree Hyper graph (IRQH). With the IRQH it's attainable to make a replacement powerful image segmentation activity model. Though the new model is mathematically outlined and no human sensory system model is expressly used, our experiments on varied image distortion types indicate the efficient of proposed model.

Key Words: Hyper Graph, Image Region Quad Tree Hyper Graph (IRQH), Image Region

1. Introduction:

Intuitively, a graph represents a set of elements and a set of pair wise relationships between those elements. The elements are called nodes or vertices, and the relationships are called edges. Formally, a graph G is defined by the sets $G = (v, E)$ in which $E \subseteq v \times v$. We may denote the i -th vertex as $v_i \in v$, and the i -th edge as $e_i \in E$. Since each edge is a subset of two vertices, we may also write $e_{ij} = \{v_i, v_j\}$. In this book we do not consider graphs with self-loops, meaning that $e_{ij} \in E$ implies that $i \neq j$. We also do not consider graphs for which there are multiple edges of the same orientation connecting the same node pair. However, for specific graph encodings, self-loops and multiple edges can be useful

Every graph may be viewed as weighted. Given a graph $G = (v, E)$, vertex weighting is a function $\hat{w}: v \rightarrow R$, and edge weighting is a function $\hat{w}: E \rightarrow R$. To simplify the notation, we will use w to refer to both vertex and edge weighting. The weight of a vertex is denoted by $w(v_i)$ or w_i , and the weight of an edge incident to two vertices is denoted by $w(v_i, v_j)$ or w_{ij} . If $w_i = 1, \forall v_i \in v$ and $w_{ij} = 1, \forall e_{ij} \in E$, then we may consider the graph to be unweighted. If not otherwise specified, all node and edge weights are considered to be equal to unity. We treat $w_{ij} = 0$ as equivalent to $e_{ij} \in E$. Intuitively, an edge weight of zero is equivalent to meaning that the edge is not a member of the edge set.

Each edge is considered to be oriented and some edges additionally directed. An orientation of an edge means that each edge $e_{ij} \in E$ contains an ordering of the vertices, v_i and v_j . An edge e_{ij} is directed if $w_{ij} \neq w_{ji}$. A graph for which none of the edges are directed is called an undirected graph, and a graph in which at least one edge is directed is called a directed graph or digraph. A directed edge is represented by the notation $e_i!j$. A directed graph is more general than an undirected graph because it does not require that $w_{ij} = w_{ji}$ for each edge. Consequently, all algorithms for directed graphs may also be applied to undirected graphs, but the converse may or may not be true. Therefore, in this chapter we use digraphs to illustrate the most general concepts. The edge orientation of an undirected graph provides a reference to determine the sign of a flow through that edge. This concept was developed early in the circuit theory literature to describe the direction of current flow through a resistive branch (an edge). For example, a current flow through e_{ij} from node v_i to v_j is considered positive, while a current flow from v_j to v_i is considered negative. In this sense, flow through a directed edge is usually constrained to be strictly positive. Drawing a graph is typically done by depicting each node with a circle and each edge with a line connecting circles representing the two nodes. The edge orientation or direction is usually represented by an arrow. Edge e_{ij} is drawn with an arrow pointing from node v_i to node v_j .

It is worthwhile to highlight the several desirable aspects of hyper graph based image segmentation: (1) different from commonly used pixel-wise similarity, hyper graphs consider patch based homogeneity, which is arguably more meaningful; (2) hyper graphs enable a user to directly specify soft constraints to be considered as edges (called hyper edges); (3) hyper graphs are suited to super pixels which are not defined on a uniform grid, and they two combined together both have favorable computational efficiency and enable a hierarchical segmentation scheme.

2. Hyper Graph in Image:

In this section we introduce hyper graphs, a natural generalization of graphs, and describe our procedure for constructing hyper graph based image models, followed by the hyper graph Laplacian matrices, which play a similar role for hyper graphs as the role of the graph Laplacian matrices for traditional graphs.

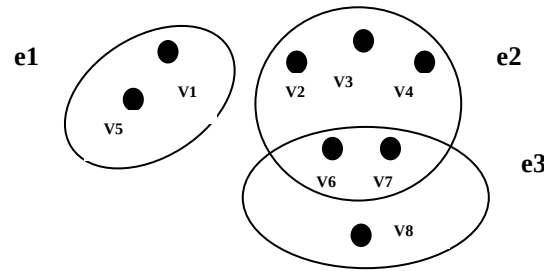


Figure 1: An example hyper graph with 8 vertices and 3 hyper edges. $V = \{v1, v2, v3, v4, v5, v6, v7, v8\}$,
 $E = \{e1, e2, e3\} = \{\{v1, v2\}, \{v2, v3, v4, v6, v7\}, \{v6, v7, v8\}\}$.

3. Hyper Graph:

A hyper graph generalizes a graph. In a hyper graph, edges can connect any number of vertices. Formally, a hyper graph is a pair (V, E) , where V is a set of vertices, and E is a set of non-empty subsets of V called hyper edges. Therefore, E is a subset of $P(V)$, where $P(V)$ is the power set of V . We present several useful definitions as follows. A hyper graph G can be represented by a $|V| \times |E|$ matrix H called the incidence matrix of G with $H(v, e) = 1$ if $v \in e$ and $H(v, e) = 0$ if otherwise. A hyper graph G can be a weighted one, where an edge e has a weight $w(e)$, and W is a diagonal matrix containing the weights of hyper edges. For a vertex $v \in V$, its degree is defined as $d(v) = \sum_{\{e \in E | v \in e\}} w(e)$. D_v is the diagonal matrix containing the vertex degrees. For a hyper edge $e \in E$, its degree is defined as $\delta(e) = |e|$, which is the cardinality of e as a finite set. D_e is the diagonal matrix containing the hyper edge degrees.

Thus, we see that hyper graphs go beyond traditional graphs in that an edge in a traditional graph defines a binary relation between two vertices, while a hyper edge defines a relation among a number of vertices. In image segmentation, a traditional graph edge models pair wise pixel similarity, while a hyper edge corresponds to an image patch.

4. Hyper Graph as Image Model:

Graph-based image segmentation methods are dependent on neighborhood graphs constructed from images based on the pairwise (distance, brightness, color) similarity. In our hyper graph-based approach, several hyper edges are constructed for each pixel containing nearby pixels and their weights are assigned based on the homogeneity of pixel colors in the hyper edge.

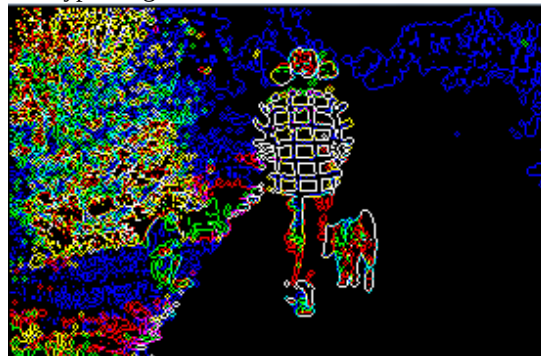


Figure 2: Fine scale super pixels overlaid on the man walking image

Our hyper graph image model utilizes the concept of *super pixels* from the computer vision literature, which are formed in a preprocessing stage to group pixels into larger sets, which are local, coherent, and preserve most of the structure necessary for segmentation at the scale of interest.

The motivations of this grouping are: (1) pixels are merely a consequence of the discrete representation of images; (2) the number of pixels is typically very high, which makes optimization on the level of pixels intractable. The normalized cuts algorithm is used to produce the super pixel map. In Figure 2, fine scale super pixels are plotted on top of the original image using the code available on line1. We observe from this example that the super pixels are roughly homogeneous in size and shape, which makes a *hierarchical* treatment of image segmentation possible, from the finest scale pixels to super pixels (hard label constraints) and then to hyper edges (soft label constraints) which merge the super pixels and finally to the learned meaningful segments. Although some structures may get lost when we use super pixels, they are usually minor details. The critical aspects of a hyper graph are the construction of hyper edges and the weight of each hyper edge. Ideally, hyper edges correspond to the set of super pixels that we want to combine together. We derive our image model by constructing multiple hyper edges at each super pixel, and assigning a weight to each hyper edge based on patch homogeneity such that the image dependent constraints for combining super pixels are incorporated in a systematic way to our hyper graphic image model. Our model has an advantage over the Image region quad tree Hyper graph (IRQH) proposed, since IRQH considers only one hyper edge at each location, which is more

prone to segmentation errors; our model, however, has multiple hyper edges at each location (super pixel) and assigns a weight to each hyper edge making the resulting segmentation more stable.

5. Image Region Quad Tree Hyper Graph (IRQH):

Image data is almost always sampled in a regular, rectangular lattice that is modeled with the graphs in the previous section. However, images are often simplified prior to processing by merging together regions of adjacent pixels using an untargeted image segmentation algorithm. The purpose of this simplification is typically to reduce processing time for further operations, but it can also be used to compress the image representation or to improve the effectiveness of subsequent operations. When this image simplification has been performed, a graph may be associated with the simplified image by identifying each merged region with one node and defining an edge relationship for two adjacency regions. Unfortunately, this graph is no longer guaranteed to be regular due to the image-dependent and space-varying merging of the underlying pixels. In this section we review common simplification methods that produce these irregular graphs.

One very well-known irregular image tessellation is the Image region quad tree Hyper graph (IRQH). A region quadtree is a hierarchical image representation that is derived from recursively subdividing the 2D space into four quadrants of equal size until every square contains one homogeneous region (based on appearance properties of every region of the underlying rectangular grid) or contains a maximum of one pixel.

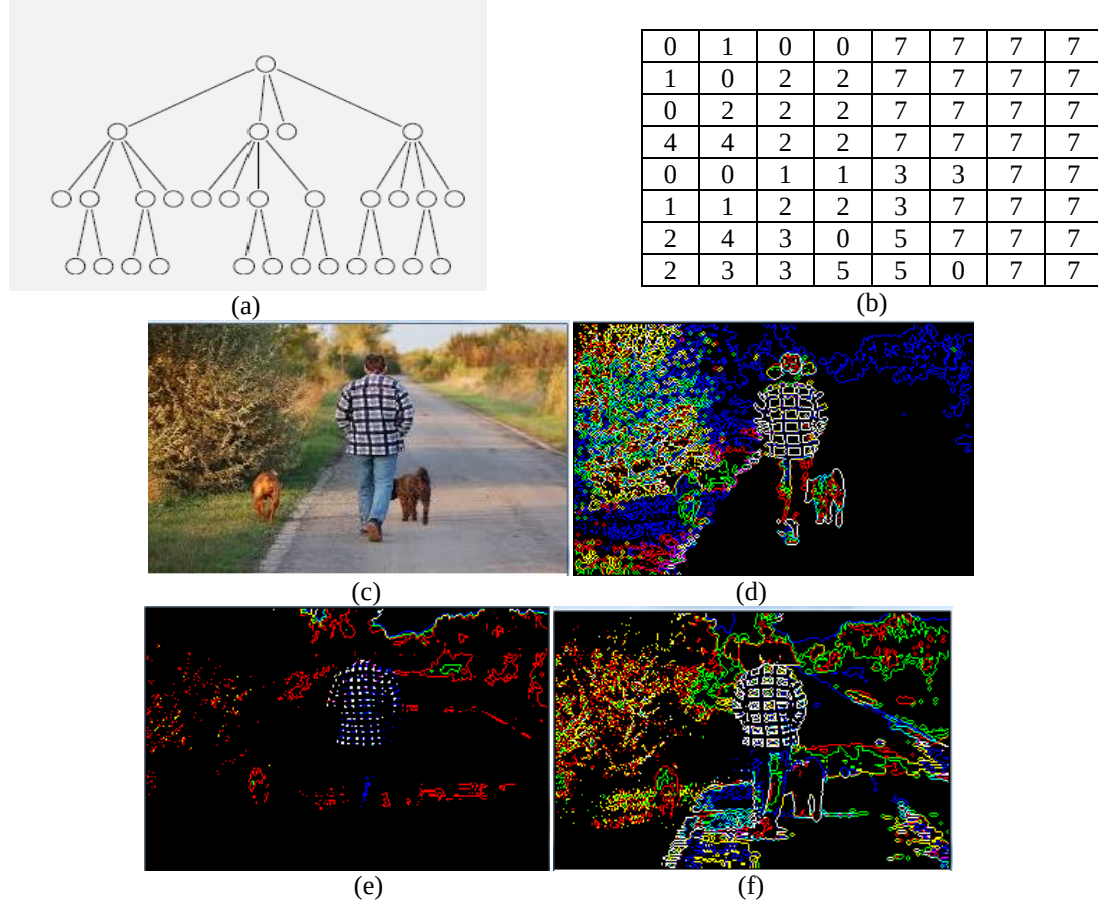


Figure 3 (a) An image with the quad tree tessellation, (b) the associated partition tree, (c) areal image (d) the region adjacency hyper graph associated to the quad tree partition, (e) and (f) two different irregular tessellations of an image

An Image region quad tree Hyper graph (IRQH) can be easily represented by a tree where each node of the tree corresponds to a square. The set of all nodes at a given depth provides a given level of decomposition of the 2D space. Figure 3 presents a simple example on an image (Figure 3(a)) with the associated tree (Figure 3(b)), as well as an example on a real image (Figure 3(c)). Quad trees easily generalize to higher dimensions. For instance, in three dimensions, the 3D space is recursively subdivided into eight octants, and thus the tree is called an octree, where each node has eight children.

Finally, any partition of the classical rectangular grid can be viewed as producing an irregular tessellation of an image. Therefore any segmentation of an image can be associated with a region adjacency graph. Figure 3(d) presents the region adjacency graph associated to the partition depicted in 3(c). Specifically, given a graph $G = \{V, E\}$ where each node is identified with a pixel, a partition into R connected regions $V_1 \cup V_2 \cup \dots \cup V_R = V$, $V_1 \cap V_2 \cap \dots \cap V_R = \emptyset$; may be identified with a new graph $\tilde{G} = \{\tilde{V}, \tilde{E}\}$ where each partition of V is identified with a node of $\tilde{V}$, that is, $V_i \in \tilde{V}$. A common method for defining $\tilde{E}$ is to let the edge

weight between each new node equal the sum of the edge weights connecting each original node in the sets, that is, for edge $e_{ij} \in \tilde{E}$

$$w_{ij} = \sum_{e_{ks}, u_k \in V_i, v_s \in V_j} w_{ks}$$

Each contiguous V_i is called a super pixel. Super pixels are an increasingly popular trend in computer vision and image processing. This popularity is partly due to the success of graph-based algorithms, which can be used to operate efficiently on these irregular structures (as opposed to traditional image-processing algorithms which required that each element is organized on a grid). Indeed, graph-based models (e.g., Markov Random Fields or Conditional Random Fields) can provide dramatic speed increases when moving from pixel-based graphs to super pixel-based graphs due to the drastic reduction in the number of nodes. For many vision tasks, compact and highly uniform super pixels that respect image boundaries are desirable. Figure 2 (e) and (f) presents some irregular tessellations based on super pixel algorithms.

Table 6: Log-average miss rate (%) on Caltech-Test of upper, left and full body parts with respect to data volume. All results are obtained by fine-tuning from NN which adopted image-level pre-training strategy

Fine-Tuning Data	Upper	Left	Full
Every 30th image	41.19	46.96	33.01
Every 10th image	38.25	46.36	30.45
Every 5 th image	36.63	41.59	23.83
Every 3 rd image	34.60	39.69	23.18
Every image	34.11	37.83	21.19

- ✓ Reasonable subset. Pedestrians are larger than 49 pixels in height and have at least 65 percent visible body parts. Reasonable subset is considered as a more representative evaluation than overall performance on all pedestrians and is the most frequent evaluation setting. Without special illustration, we use Reasonable as our default setting to compare performance.
- ✓ Partial and heavy occlusion subsets where pedestrians are larger than 49 pixels in height and have 1–35 and 36– 80 percent occluded body parts, respectively.

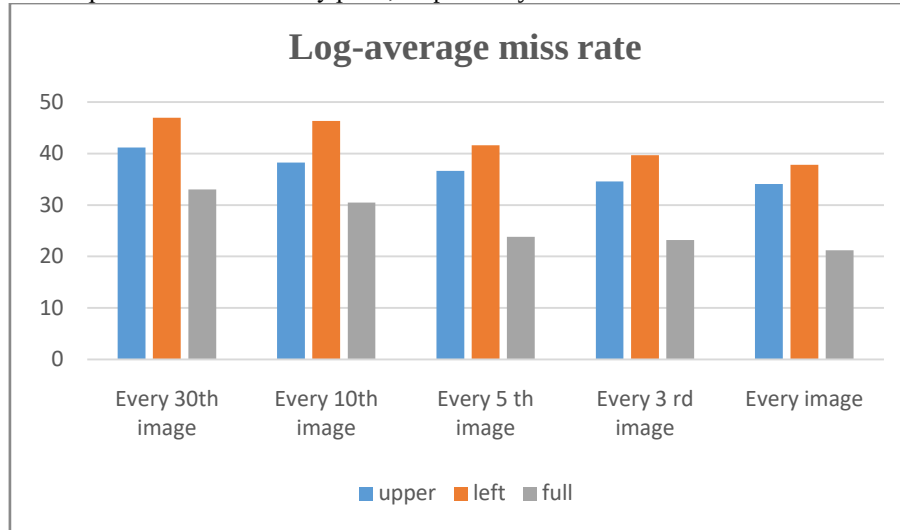


Figure 4: Graph for Log-average miss rate

5. Discussion:

First, our Image region quad tree Hyper graph (IRQH) method addresses two kinds of constraints: the hard ones are modeled by upper pixels, where we require that constituent pixels take the same label; the soft ones are modeled by hyper edges, where label consistency is encouraged. Based on our method, one can build an interactive interface, where a user supplies soft constraints (regions likely to come from an object) and the system takes them into consideration by adding new hyper edges in addition to those constructed a priori. Second, it is possible to combine active learning, which is the framework that allows the learner to ask for informative examples and in our case seeds, with hyper graph-based image region. There are sometimes regions in an image with varied texture and color, which can be very hard to segment given only few labeled pixels. Allowing a user to interactively provide some segment labels at critical points could greatly enhance the region quality. Third, the efficiency of our approach comes from the combination of quad tree and hyper graphs. The majority of computational time is spent on calculating a fine scale super pixel map, which there are methods in the literature to accelerate. Suppose N is the number of super pixels. The time needed for hyper graph construction is $O(N)$, and the linear coefficient is bounded by the maximum number of hyper edges at one super pixel, which is typically a small number. Empirically, the segmentation process when super pixels are available

takes from 0.5 to 2.5 minutes on a 1.5 GHz Pentium machine. The computation time depends on the size and the texture content of an image.

6. Conclusion:

In this paper, we have demonstrated the effectiveness of learning Image region quad tree Hyper graph (IRQH). In particular, we use hyper graph based interpolation and show its equivalence to an iterative procedure based on random walks on hyper graphs. Experimentally, we have achieved competitive results on the image with our method. We are planning to use more sophisticated probabilistic models to address the natural hierarchy of hard and soft constraints in image region.

7. References:

1. M. Belkin and P. Niyogi. Semi-supervised learning on Riemannian manifolds. *Machine Learning*, 56(1-3): 209–239, 2004.
2. M. Belkin, P. Niyogi, and V. Sindhwani. Manifold regularization: A geometric framework for learning from labeled and unlabeled examples. *Journal of Machine Learning Research*, 7: 2399–2434, 2006.
3. C. Berge. *Hyper graphs*. North-Holland, Amsterdam, 1989.
4. A. Blake, C. Rother, M. Brown, P. Perez, and P. Torr. Interactive image segmentation using an adaptive GMMRF model. In *ECCV*, 2006.
5. Y. Y. Boykov and M.-P. Jolly. Interactive graph cuts for optimal boundary & region segmentation of objects in n-d images. In *ICCV*, 2001.
6. A. Bretto and L. Gillibert. Hyper graph-based image representation. In *Graph-Based Representations in Pattern Recognition*, 2005.
7. Cavallaro, E. D. Gelasca, and T. Ebrahimi. Objective evaluation of segmentation quality using spatio-temporal context. In *ICIP*, 2002.
8. F. R. K. Chung. *Spectral Graph Theory*. Regional Conference Series in Mathematics, Number 92. 1997.
9. D. Zhou and J. Huang and B. Schölkopf. Learning with hyper graphs: Clustering, classification, and embedding. In *NIPS*, 2006.
10. O. D'Hondt, L. Ferro-Famil, and E. Pottier. The gradient structure tensor as an efficient descriptor of spatial texture in polarimetric SAR data. In *IEEE International Geoscience and Remote Sensing Symposium*, 2006.
11. O. Duchenne, J.-Y. Audibert, R. Keriven, J. Ponce, and F. Segonne. Segmentation by transduction. In *CVPR*, 2008.
12. L. Grady. Random walks for image segmentation. *IEEE PAMI*, 28(11):1768–1783, 2006.
13. L. Hagen and A. B. Kahng. New spectral methods for ratio cut partitioning and clustering. *IEEE Transactions on Computer-Aided Design*, 11(9):1074–1085, 1992.
14. M. Kass, A. Witkin, and D. Terzopoulos. Snakes: Active contour models. *International Journal of Computer Vision*, 1(4):321–331, 1987.
15. E. N. Mortensen and W. A. Barrett. Intelligent scissors for image composition. In *SIGGRAPH*, 1995.
16. X. Ren and J. Malik. Learning a classification model for segmentation. In *ICCV*, 2003.