

AUTOMATIC SCALE ADAPTATION OF SEMANTIC NETS

K. Pakzad, J. Heller

Institute of Photogrammetry and GeoInformation, University of Hannover
Nienburger Str. 1, 30167 Hannover, Germany – (pakzad,heller)@ipi.uni-hannover.de

Commission III, WG III/4

KEY WORDS: Interpretation, Expert System, Knowledge Base, Model, Scale, Representation, Multiresolution, Generalization

ABSTRACT:

This paper deals with a methodology to derive object models for automatic object extraction in low resolution images from models created manually for high resolution images. The object models are represented by semantic nets, which describe landscape objects explicitly in terms of natural language. Starting from semantic nets for high resolution images the strategy is to first decompose them into parts, which can be handled autonomously. The object parts are then adapted, i.e. generalised, to smaller scale. The adaptation takes into account the object shape, radiometry, and texture. For the generalisation process “scale change models” are used, which describe how different types of objects evolve over scale mathematically. Finally, all object parts are fused and transferred to a semantic net representation. In this paper first results of the described methodology are presented. Focussing on line-type objects, such as streets, we describe how to create an object description with semantic nets using constraints, which have to be satisfied, in order to be able to adapt the nets to other scales automatically. In addition we show tests of the behaviour of some edge- and line-extraction operators through scale space. These tests are necessary to predict the scale behaviour of different object types. At last, we describe as an example for a particular object events during scale change observed in an image and their impact on a semantic net. This example demonstrates the suitability of the proposed kind of semantic net to follow the scale space events in digital images, and thus, its applicability in an automatic approach.

1. INTRODUCTION

Landscape objects appear differently in remote sensing images of differing resolution. While many object details are visible in high resolution images, in low resolution images many of them disappear or merge. Even the dimensionality can change. Where in high resolution images areas are observable, in low resolution images lines or even points might be found. This fact also affects an automatic extraction of landscape objects from digital images with different resolutions. For an automatic extraction from satellite and aerial images knowledge-based systems with an explicit knowledge representation, such as semantic nets, offer high flexibility and can easily be structured (Pakzad, 2001). This knowledge representation contains the object models, which describe the objects with all relevant parts and characteristics. As described above the models for the same objects have to be different depending on the resolution of the images. They are tailored to specific scales of aerial and satellite images. Decision about the best scale for object detection is mostly still made intuitively (Schiewe, 2003). In (Baumgartner, 2003), the representation of roads in a small and a large image scale is combined in a semantic net. However, the fusion of the two scales is solely used for increasing the reliability of the extraction results.

Existing approaches for explicit object models do not permit an automatic transfer to other scales. Hence, a new model is to be developed for each image scale manually. For the case of scale reduction, a description of object behaviour is possible though, as investigations of features in scale-space indicate (Witkin, 1986). The scale-space theory was formulated for a multi-scale representation of objects, depending only on one parameter for scale. Following this theory, with increasing scale parameter, i.e. lower spatial resolution, new details will not appear, but existent details will disappear and merge with each other (Lindeberg, 1994). The object representation in image data of

lower spatial resolution can therefore be predicted starting from high resolution. A methodology for an automatic adaptation of object models to lower spatial resolutions would make the manual generation of these object models for different resolutions redundant. Thus, a once created object model could be utilised for a wider range of applications and for diverse sensor types exhibiting a wider range of image scale.

This paper therefore presents an approach to derive object models for low resolution images from models created manually for high resolution images. Although the contents of scale-space theory were widely applied to many image processing tasks, e.g. for edge and line detection algorithms (Lindeberg, 1998), the connection to semantic net object representation for knowledge based image analysis is new.

Section 2 gives an overview about the general strategy of the procedure and briefly describes the different steps. Section 3 focuses on the composition of the semantic nets and suggests some constraints, in order to be able to handle the semantic nets automatically regarding scale adaptation. The semantic nets represent the high level processing of the image interpretation task, but also the low level processing, which is directly connected to the nets, has to be observed. Section 4 describes tests on the scale behaviour of some feature extraction operators, and section 5 contains an example for scale change events observed in a scene and their impact on the semantic net.

2. STRATEGY FOR SCALE ADAPTATION

This section gives an overview of the proposed strategy for scale adaptation. As shown in Fig. 1, the main input of the process is a manually created object model, represented as a semantic net, with the description of that object, which has to be extracted from images. The details of the object description are adjusted to a large start scale. Object parts, which are not observable at that scale, are also not represented in the object

model. The automatic creation of the semantic net requires some constraints in order to be able to adapt the net to another scale. These constraints are described in section 3. The second input of the process is the target scale, which has to be smaller than the start scale.

The first step is a decomposition of the semantic net. The goal of the decomposition is to identify parts of the semantic net, which can be processed separately. These parts can consist of single object parts or blocks of them, if the object parts influence each other during scale adaptation: E.g. two objects with a small distance to each other possibly have to be adapted together (depending on the target scale), because during the scaling process the small distance can disappear and objects can merge.

The next step is the scale adaptation of the decomposed parts themselves. The selected object parts and blocks have to be generalized. In this process different aspects have to be taken into account: It is possible that the object type changes as well as the object attributes. All objects of the same object type exhibit a comparable behaviour in scale change and can be extracted by the same group of feature extraction operators.

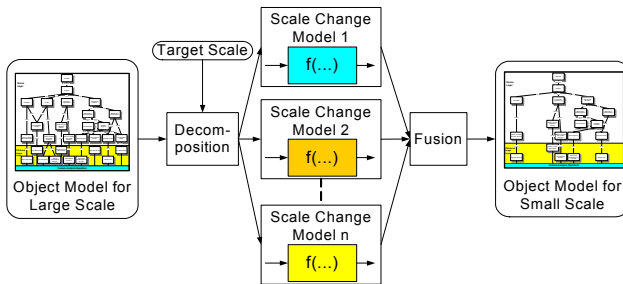


Figure 1. Strategy for Scale Adaptation

For the scale adaptation of the elements we intend to use scale change models. Scale change models describe the kind of change of a certain object type depending on the value of scale change. Different attributes, such as the grey value contrast of the object to the neighbours or the spatial measurements, are used as input parameters for the scale change models. The decision, of how to change the object parts or the blocks, has a direct connection to the question of how they can be extracted after scale change by using feature extraction operators. If it is necessary to change the feature extraction operators after scale change, the object type has changed. Hence, the scale behaviour of feature extraction operators has to be taken into account. In section 4 we describe first results of some examinations, aiming to determine the suitable scale range for certain feature extraction operators.

The last step is the fusion of the adapted object parts to a complete semantic net, which describes the object in the smaller scale and which can be used for an automatic object extraction in low resolution images.

3. COMPOSITION OF SUITABLE SEMANTIC NETS

The knowledge representation in this approach uses the form of the semantic nets of the knowledge based system AIDA (Liedtke et al., 1997, Tönjes, 1999). This system was developed for automatic interpretation of remote sensing images. Semantic nets (Niemann et al., 1990) are directed acyclic graphs. They consist of nodes and edges linking the nodes. The nodes represent the objects expected in the scene while the edges or links of the semantic net model the relations between these objects. Attributes define the properties of nodes and edges.

Two classes of nodes are to be distinguished: the concepts are generic models of the objects and the instances are realizations of their corresponding concepts in the observed scene. Thus, the knowledge base, which is defined prior to the image analysis, is composed of concepts. During the interpretation a symbolic scene description is generated consisting of instances. The relations between the objects are described by edges or links of the semantic net. Objects are composed of parts represented by the part-of link. Thus, the detection of an object can be simplified to the detection of its parts. The transformation of an abstract description into its more concrete representation in the data is modelled by the concrete-of relation, abbreviated con-of. This relation allows for structuring the knowledge in different conceptual layers, for example a scene layer and an image layer.

The concept of semantic nets for the extraction of particular objects enables many possibilities for the generation and composition of a particular object model, i.e. the representation of an object model for the extraction of an object can be realized with different semantic nets. Based on the goal of this research – adapting semantic nets automatically to a smaller scale – it is necessary to find rules for the generation of semantic nets.

The semantic nets to be created should satisfy the following constraints:

- They should have a structure, which enables to analyse them automatically regarding scale behaviour. The structure should enable an automatic decomposition of the semantic net into suitable parts in order to treat them separately.
- They should describe the objects completely with all characteristics and attributes, that are important for the extraction of the objects in the starting scale.
- They have to be suitable for an automatic extraction of the objects from digital images. The semantic net should contain a refinement of the whole object into suitable parts, which can be extracted directly by using certain feature extraction operators.
- They should be easy to create by using standard language. The rules should direct the experts in creating the semantic nets, not complicate it. Otherwise advantages of semantic nets and expert systems, like the explicit knowledge representation, would be lost.

In this stage of our research, we focus on roads and road markings. As a further specification, the focus lies on objects which run for a long distance parallel to the road axis along the road. That means we deal with objects such as lane markings, but not with objects such as zebra crossings. We are able to describe such objects with the object types periodic and continuous stripes and lines. The semantic nets, which we create for such objects, contain only these four object types. We represent roads first by describing the entire stripe (the pavement) as the basis object and, as part of a road, the markings on it (see Fig.6). As a rule, only the objects at the bottom layer will be extracted. That means for Fig.7 the object will be extracted by finding its line markings.

The nodes of the semantic net represent objects. We define for every object of the semantic net the following attributes:

- Object Name: The description of the object by its name in standard language.
- Object Type: All objects of the same object type have a comparable behaviour in scale change and can be extracted by the same group of feature extraction

operators. For the description of streets we define four object types: *continuous stripe*, *periodic stripe*, *continuous line* and *periodic line*. This also affects the group of feature extraction operators, which are used to extract the object type. The difference between stripe and line is given by the width of the object. We define that lines are not wider than 2 pixels. For the extraction of lines in raster images other feature extraction operators will be used than for the extraction of stripes. Stripes would be extracted by finding the edges. The specification periodic or continuous indicates dashed or continuous lines and stripes respectively.

- Grey Value: This attribute describes the radiometric characteristic of the object.
- Extent: This value describes the width of the object and is independent from resolution, i.e. is stated in meters, not in pixels.
- Periodicity: For periodic object types this attribute expresses the ratio of the length of the lines/stripes and the extent. For continuous object types this attribute has no meaning.

The relations, which are used in the presented net in Fig.6, are part-of and spatial relations. Some of the part-of relations have additional attributes, as “central” or “left/right boundary”. The attribute “central” can be used to guide the object extraction. The attribute “left/right boundary” has an important function regarding the scale adaptation: part-of relations with these attributes describe the boundary parts of an object. These parts are important, if neighbouring objects exist. In that case the border parts of both neighbouring objects have to be considered, because they can affect each other as scale varies. The spatial relations play a key role in the object description, because they directly affect the necessary modifications for scale adaptation of the nets. It is essential that the position of all object parts are clearly specified by spatial relations. For the examined objects, the position perpendicular to the street axis has to be described. We use the relations left-of and right-of for this description. Furthermore, attributes are important for the description of the distances. Usually, it is not possible to describe the exact distance to another object. We therefore use ranges for distances here. The magnitude of the distances is also expressed in meters, independently from scale.

All nodes of the semantic net are connected to feature extraction operators, which are able to extract the represented objects. But the strategy of object extraction is to call the feature extraction operators only for the nodes without parts, corresponding to the nodes at the bottom of the semantic net. In Fig.6 the node “Roadway” contains a connection to a feature extraction operator, which is able to extract stripes with the given constraints. But as long as markings on the stripe are extractable in the target scale, the extraction of “Roadway” would be realized by the extraction of the markings.

Based on these semantic nets the scale behaviour of the defined object types can be investigated. Taking into account single object parts, the following behaviour is possible:

Nr	Before Scale Change	After Scale Change
1	<i>Continuous Stripe (CS)</i>	<i>CS or CL or Invisible</i>
2	<i>Continuous Line (CL)</i>	<i>CL or Invisible</i>
3	<i>Periodic Stripe (PS)</i>	<i>PS or CS or PL or CL or Invisible</i>
4	<i>Periodic Line (PL)</i>	<i>PL or CL or Invisible</i>

Table 1. Object Type Scale Behaviour of Single Objects

Unfortunately, these four possibilities are not sufficient for an automatic scale adaptation of the semantic nets. It is also possible that different parts affect each other during scale change. As an example, two stripes with a small distance apart will merge at a certain scale. Hence, neighbouring parts have to be analysed simultaneously in this case.

Taking into account object pairs, which might affect each other, the following possibilities can be found:

Nr	Before Scale Change	After Scale Change
5	<i>Cont. Stripe – any</i>	<i>CS or CL or Invisible</i>
6	<i>Per. Stripe – Per. Stripe</i> <i>Per. Stripe – Cont. Line</i> <i>Per. Stripe – Per. Line</i>	<i>PS or CS or PL or CL or Invisible</i>
7	<i>Cont. Line – Cont. Line</i> <i>Cont. Line – Per. Line</i>	<i>CL or Invisible</i>
8	<i>Per. Line – Per. Line</i>	<i>PL or CL or Invisible</i>

Table 2. Object Type Scale Behaviour of Object Pairs

This scale behaviour of the object types has to be used for the adaptation of the semantic concept nets. It is possible, that a given range in a concept net for a distance between two objects will lead in combination with a given target scale after scale adaptation to different possibilities for the new object type. In that case different possibilities have to be included in the new concept net as representation of one object.

The question, at which target scale the object type changes is directly connected to the scale behaviour of feature extraction operators. This problem is addressed in the next section.

4. SCALE BEHAVIOR OF FEATURE EXTRACTION OPERATORS

As described in section 3, the object type of a node in a semantic net is also determined by the feature extraction operator, which is bounded to the nodes of a certain object type and is used for its extraction. Objects of different types use different feature extraction operators. But as scale varies the object type may change, because the same operator is not able to extract the same object type successfully at all resolutions. In order to be able to predict from which scale on the object type has to change, an analysis is necessary about the scale range, in which the feature extraction operators are usable. The performance of three commonly used higher developed operators for edge and line detection are exemplarily examined - Canny, Deriche and Steger. The goal of this investigation is to analyse the behaviour of the operators in sensor data of different resolutions. For that in a first step the different sensors are simulated by creating synthetic images of different resolutions and applying the operators on them.

The Canny edge detector was developed as the “optimal” edge detector (Canny, 1986). Its impulse response shape closely resembles the first derivative of the Gaussian. The Deriche edge detector is based on the Canny operator, but uses recursive filtering and thereby reducing computational effort (Deriche, 1987). The Steger operator extracts lines in sub-pixel accuracy by using the first and second order derivative of the Gaussian (Steger, 1998).

The generated synthetic grey value image has a size of 1000x1000 pixels and also is composed of a bright line with two pixel width stretching over the entire image on dark background. An image pyramid was created from this synthetic image by Gaussian interpolation. The pyramid comprises 100 levels from the original image (labelled as pixel size 1.0) to the smallest image with the largest pixel size of 1000-fold,

corresponding to the lowest resolution. Each pyramid level image is then smoothed by a 3x3 Gauss filter. White noise with certain amplitudes is added afterwards. In these pyramid images, the feature extraction operators were run, varying the values for contrast (between line and background) and noise.

For the edge extraction algorithms, the results of the Canny and Deriche operators were followed by a non-maxima-suppression and thresholding. Finally, the skeleton was derived. The edge detection sequence is depicted in Fig. 2. The response of the Steger operator in the line detection algorithm was simply thresholded.

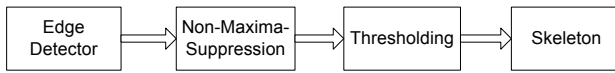


Figure 2. Edge Detection Algorithm

The edge and line detection algorithms were optimised for the smallest pixel size of the image pyramid, i.e. the resolution of the creation stage. The described parameters were maintained for the edge and line detection in all pyramid images. The images were all processed with the same procedure (same operators with the same parameter values) to ensure comparability of the results. In this paper we call one and the same operator with different parameters as different operators. Performance of the operators was obtained by recording the ratio of the actual edge and line length of the operator in the image to the expected well known length of the edges and line. All operators yield 100% performance in the first stage.

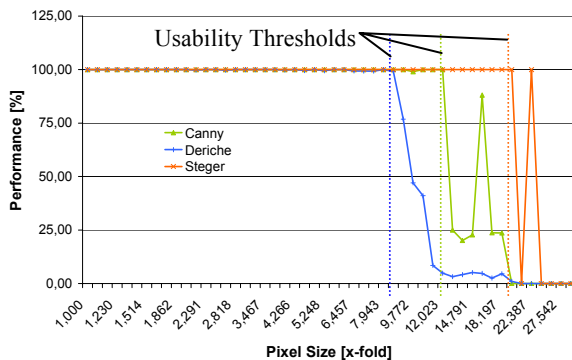


Figure 3. Performance of the Canny, Deriche and Steger Operator

The dependence of the feature extraction operator's performance on image resolution was analysed. The performance is gradually decreasing for lower resolutions. Fig.3 shows the performance curves of the operators in the highest resolution image of pixel size 1.0 with a grey value difference between the dark background and the white line of 240 and a noise amplitude of 3, corresponding to approximately 1% of the grey value range. This noise level was chosen to simulate a realistic noise impact on digital images.

As can be seen, the performance curves of the three examined operators behave quite differently. The difference in the shape of the performance curves is not only due to the operator itself, but is dependent on the chosen parameters in the implementation as well. The choice of the threshold values mainly determines at which pixel size the good performance breaks off. While the Deriche operator's best performance more or less ends abruptly, the performance of the Canny and Steger operators oscillates over a certain range of resolution before failing to detect features. The results of these two operators in these resolution ranges must be regarded as unreliable. Feature

extraction at these and lower resolutions cannot be carried out with these operators under the preference of the chosen parameters. The derived usability thresholds are marked in Fig3.

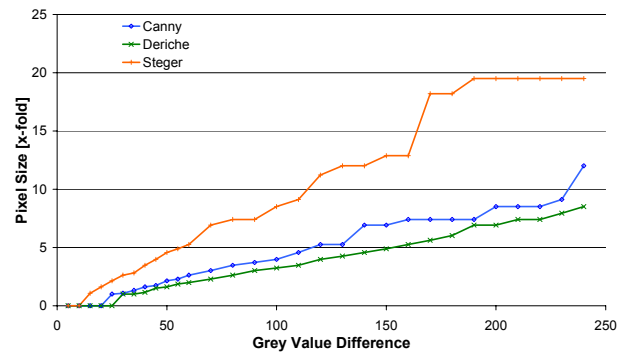


Figure 4. Usability Range of Canny, Deriche and Steger Operators with varying Contrast

Contrast in the image plays an important role in feature extraction and influences the operator's performance. To determine the resolution up to which an extraction with the presented operators with a given contrast is reliable, the limit for the operator performance was set to 98%. If the output of the extraction algorithm falls below this limit, the operator was regarded unusable for the respective image resolution on to any lower resolution, at least with the implemented parameters. Fig.4 depicts the performance limits for the three operators with a given noise level of 1% depending on the grey value differences between the dark background and the lighter line. A pixel size of zero means there were no operator responses even in the highest resolution of 1.00 because of insufficient contrast.

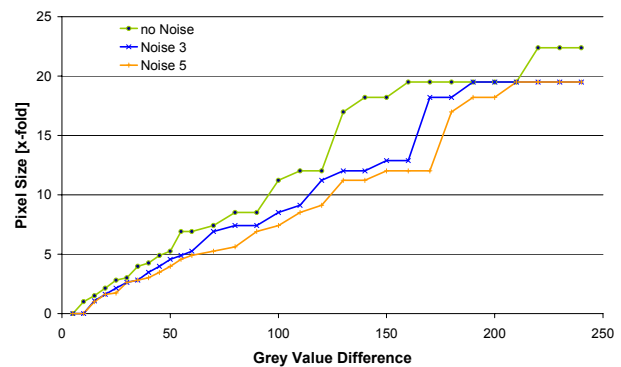


Figure 5. Usability Range of the Steger Operator with varying Contrast and Noise

Furthermore, feature extraction is also susceptible to prevailing image noise. Analyses were carried out to the operator's performance with three white noise amplitudes – 0, 3 and 5, corresponding to 0% - 2% of the 8-bit grey value range. With increasing noise level in the image the extraction performance declines. The sensitivity of the Steger operator to noise is exemplarily presented in Fig.5. The Canny and Deriche operators exhibit a very similar behaviour in varying noise. The performance of the Canny, Deriche and Steger operators is degraded by the influence of low contrast and high noise. The smaller the grey value differences and the higher the noise amplitude, the lower the image resolution at which the trustable performance of the operator breaks off and the smaller the resolution range for which the operator can be used.

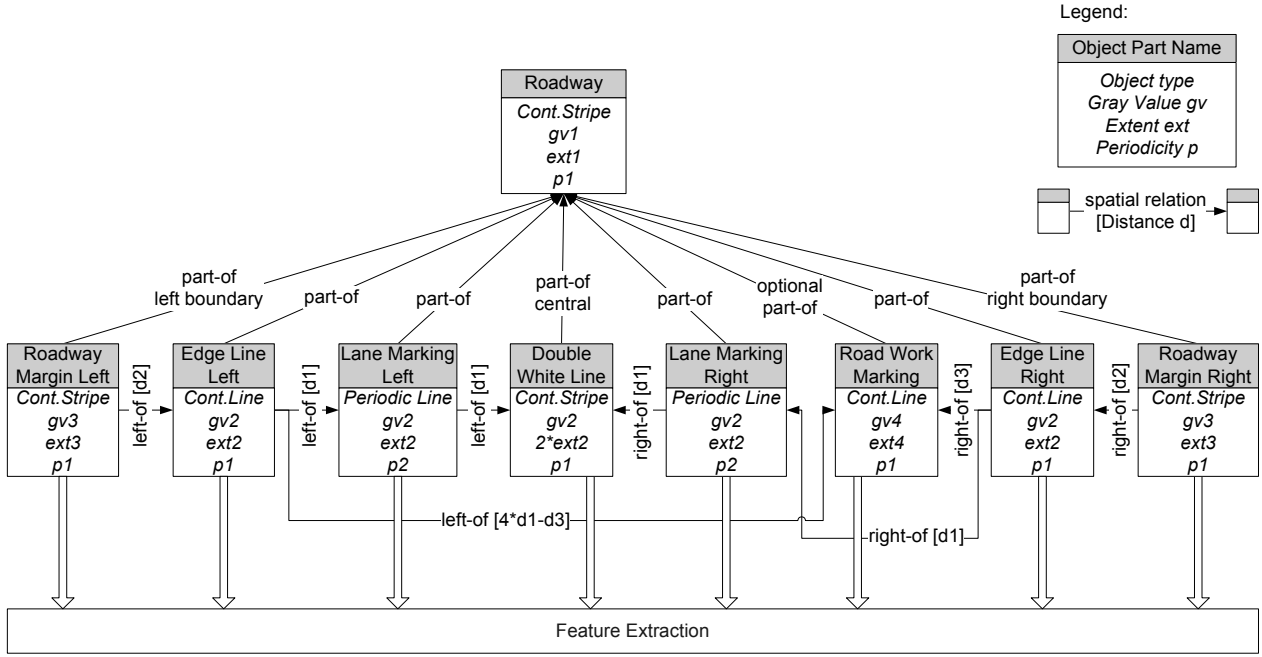


Figure 6. Concept Net for Dual Carriageway at Largest Scale, Generated for Images with Ground Pixel Sizes of 3.3 - 7 cm/pel

5. EXAMPLE FOR SCALE ADAPTATION

For an exemplarily chosen application of a dual carriageway we created a semantic net following the developed constraints as described in section 3. Fig.7 shows our test image, a cut-out from an aerial image with a ground pixel size of 3.3 cm. The goal of this section is to show that the proposed kind of semantic net is suitable to follow the scale space events in digital images, and is therefore suitable to be used in an automatic approach.

In the concept net, as presented in Fig.6, the roadway is modelled as a continuous stripe with certain ranges for grey values and extent, i.e. width. The roadway itself is composed of various parts, the road markings and roadway margins. While the road markings are of the object type periodic or continuous line, the object type of the margins is a continuous stripe. Attributes for grey value, extent and periodicity are assigned to the object parts of the roadway as well. The declaration of the spatial relations between these object parts is essential for the scale adaptation process, as previously described in section 3. Here, the distance d_1 represents the width of a single lane, d_2 corresponds to the distance between the outmost edge line and the roadway margins and d_3 locates the optional road work marking from the outmost edge line. All nodes of the net are connected to the appropriate feature extraction operators, but only the operators connected to the bottom nodes are used. In addition, the boundary object parts are labelled to facilitate the search for adjacent objects. With this information groups of objects can be formed, which have to be analysed in conjunction regarding scale space behaviour.

The following semantic nets represent some instance nets of the adaptations to smaller scales of this particular roadway scene. The adaptations are done manually based on a visual inspection of larger ground pixel sizes, i.e. smaller scales of the image as seen in Fig.7. The adapted nets correspond to a selection of smaller scales. At these scales at least one event in scale space necessitates the adaptation of the previous net, which is appropriate for a larger scale.

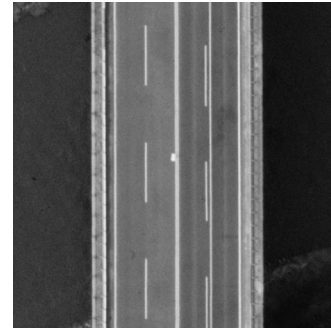


Figure 7. Aerial Image, Dual Carriageway

As scale is decreased, first the object type of the central object part, the double white line, changes from continuous stripe to continuous line, cf. Fig.8. Secondly, the road work marking merge with the neighbouring right lane marking due to the small separating distance between them. Since road work marking is an optional part, the term of the lane marking is maintained for the name of the resulting object part. The attributes of this modified object part, however, change. The resulting object type is continuous line.

With further decreasing of scale the edge line markings merge with the roadway margins, both left and right side. Fig.8 depicts the semantic net adapted to this scale. The nodes of the edge line markings are combined with the nodes of the roadway margins, resulting in new values for the attributes, grey value and extent. Since these new object parts are now located at the border of the entire object dual carriageway, they have to be labelled as boundary objects. Even though the feature extraction operators are not included in Fig.8, the connections to the nodes still exist and the operators are called for the extraction of the bottom object parts in the image.

When the scale gets so small that single lines are not detectable anymore, the lane markings vanish. Distances between the remaining object parts, double white line and roadway margins have to be modified, cf. Fig.9. The stripes of the roadway margins shrink to 2 pixels in width and thus, the object type changes to a continuous line.

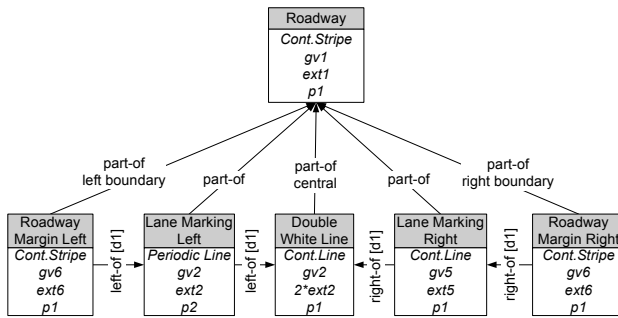


Figure 8. Instance Net for Dual Carriageway for 0.9 m/pel

In the next stage of the net adaptation process, the net consists merely of the roadway itself - the only feature that is still left detectable at that scale. The feature extraction operator connected to the roadway node will now be called to extract a continuous stripe, as the roadway has become the bottom node. At last, the object type changes from continuous stripe to continuous line before the line vanishes and there is no roadway or part thereof extractable in the example scene.

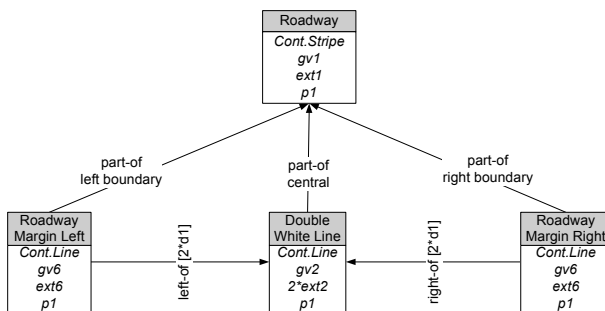


Figure 9. Instance Net for Dual Carriageway for 1.7 m/pel

6. CONCLUSION

In this paper an approach to derive object models for low resolution images from models created manually for high resolution images was presented. After an overview about the general strategy of the procedure we focussed on the composition of the semantic nets and suggested some constraints, in order to be able to handle the semantic nets automatically regarding scale adaptation. The prediction of the scale behaviour of object types requires investigations on the scale behaviour of feature extraction operators, which we presented for three operators. At last, we described an example for scale change events observed in a scene and their impact on the semantic net. This example demonstrates the suitability of the proposed kind of semantic net to follow the scale space events in digital images, and thus, its applicability in an automatic approach.

Future work will deal with the specification of the exact steps of an automatic scale adaptation of semantic nets. Furthermore, we intend to work on extensions of the described semantic nets to new object types, and the impacts on the semantic net creation rules. In addition we want to work on the implementation of the nets into the knowledge based system GeoAIDA (Bückner, 2002, Pahl, 2003, successor of AIDA). Regarding the investigation of the feature extraction operators (section 4) an exact simulation of sensor data in different resolutions would require the incorporation of more complex models than we used. We assume that the used procedure is sufficient for our task. Yet, this assumption still has to be

verified by using real sensor images. Eventually, the comparison of our results with predictions from scale space theory is surely interesting.

7. ACKNOWLEDGEMENTS

This work has been funded within a project by the Deutsche Forschungsgemeinschaft under grant HE 1822/13.

8. REFERENCES

- Baumgartner, A., 2003. Automatische Extraktion von Straßen aus digitalen Luftbildern. *DGK, Reihe C, Dissertationen*, No. 564, München, 91 p.
- Bückner, J., 2003. Ein wissensbasiertes System zur automatischen Extraktion von semantischen Informationen aus digitalen Fernerkundungsdaten. *Schriftenreihe des TNT der Universität Hannover*, Band 4, ibidem-Verlag, Stuttgart, 162 p.
- Canny, J., 1986. A Computational Approach to Edge Detection. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 8(6), pp. 679-698.
- Deriche, R., 1987. Using Canny's Criteria to derive a Recursively Implemented Optimal Edge Detector. *International Journal of Computer Vision*, 1(2), pp. 167-187.
- Liedtke, C.-E., Bückner, J., Grau, O., Growe, S., and Tönjes, R., 1997. AIDA: A system for the knowledge based interpretation of remote sensing data. *3rd. Int. Airborne Remote Sensing Conference and Exhibition*, Vol. II: pp. 313-320.
- Lindeberg, T., 1994. *Scale-Space Theory in Computer Vision*. Kluwer Academic Publishers, The Netherlands, 423 p.
- Lindeberg, T., 1998. Edge Detection and Ridge Detection with Automatic Scale Selection. *International Journal of Computer Vision*, 30(2), pp. 117-154.
- Niemann, H., Sagerer, G., Schröder, S. and Kummert, F., 1990. ERNEST: a semantic network system for pattern understanding. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 12(9), pp. 883-905.
- Pahl, M., 2003. Architektur eines wissensbasierten System zur Interpretation multisensorieller Fernerkundungsdaten. *Schriftenreihe des TNT der Universität Hannover*, Band 3 ibidem-Verlag, Stuttgart, 145 p.
- Pakzad, K., 2001. Wissensbasierte Interpretation von Vegetationsflächen aus multitemporalen Fernerkundungsdaten. *DGK, Reihe C, Dissertationen*, No. 543, München, 104 p.
- Schiewe, J., 2003. Auswertung hoch auflösender und multi-sensoraler Fernerkundungsdaten. *Habilitationsschrift, Fachgebiet Umweltwissenschaften, Schwerpunkt Geoinformatik, Hochschule Vechta*, 159 p.
- Steger, C., 1998. An Unbiased Detector of Curvilinear Structures. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 20(2), pp.113-125.
- Tönjes, R., 1999. Wissensbasierte Interpretation und 3D-Rekonstruktion von Landschaftsszenen aus Luftbildern. Dissertation, Universität Hannover, *Fortschritt-Berichte VDI, Reihe 10*, No. 575, VDI-Verlag, Düsseldorf, 117 p.
- Witkin, A., 1986. Scale space filtering. In Pentland, A. (Ed.): *From Pixels to Predicates*. Ablex Publishing Corporation, New Jersey, pp. 5-19.