



Chapitre d'actes

2004

Published version

Open Access

This is the published version of the publication, made available in accordance with the publisher's policy.

Interactive segmentation with hidden object based annotations : towards smart media

Rytsar, Yuriy; Voloshynovskyy, Svyatoslav; Ehrler, Frédéric; Pun, Thierry

How to cite

RYTSAR, Yuriy et al. Interactive segmentation with hidden object based annotations : towards smart media. In: Proceedings of SPIE Electronic Imaging, Storage and Retrieval Methods and Applications for Multimedia. San Jose (USA). [s.l.] : SPIE, 2004. (SPIE Proceedings) doi: 10.1117/12.527000

This publication URL: <https://archive-ouverte.unige.ch/unige:47924>

Publication DOI: [10.1117/12.527000](https://doi.org/10.1117/12.527000)

Interactive segmentation with hidden object-based annotations: toward Smart Media

Yuriy Rytsar^a, Slava Voloshynovskiy^{*1a}, Frederic Ehrler^a and Thierry Pun^a

^aComputer Science Department, University of Geneva,
24 rue du General Dufour, CH-1211 Geneva 4, Switzerland

ABSTRACT

In this paper a novel “Smart Media” concept for semantic-based multimedia security and management is proposed. This concept is based on interactive object segmentation (considered as side information in visual human-computer interface) with hidden object-based annotations. Information-theoretic formalism is introduced that considers the human-computer interface as a multiple access channel. We do not consider an image as a set of pixels but rather as a set of annotated regions that correspond to objects or their parts, where these objects are associated with some hidden descriptive text about their features. The presented approach for “semantic” segmentation is addressed by means of the human-computer interface that allows a user to easily incorporate information related to image objects and to store them in a secure way. Since each selected image object carries its own embedded description, this makes it self-containing and formally independent from the particular image format used for storage in image databases. The proposed object-based hidden descriptors are invariant to changes of image filename or/and image headers, and are resistant to object cropping/insertion operations, which are usual in multimedia processing and management. This is well harmonized with the “Smart Media” concept where the image contains additional information about itself, and where this information is securely integrated inside the image while remaining perceptually invisible.

Keywords: interactive segmentation, side information, human-computer interface, multiple access channel, hidden annotation, Smart Media, labeling, semantics extraction, object-based retrieval, watermarking.

1. INTRODUCTION

The drastic increase of distributed public networks, Internet, digital imaging devices, telemedicine applications as well as the enormous amount of multimedia makes the storage and retrieval of images, videos and audio one of the most challenging problems of modern multimedia communications and management. The growing amount of on-line image collections and databases requires tools to efficiently manage, organize, and navigate through them. Taking into account the practical difficulties to control and to verify existing audio/visual communications in distributed and public networks, one can imagine additional ambiguity and insecurity in retrieval and management of multimedia information. The urgent necessity for efficient multimedia retrieval engines as well as for network security systems is recognized by the current multimedia industry and research. Additionally, the situation is complicated by the lack of accepted protocols, the relative simplicity with which the content can be modified (change of file name, modified headers of JPEG and JPEG2K graphical formats, object cropping/insertion from one image/video to another).

At the same time the metadata become the cornerstone of intelligent Web environments. Web-based applications increasingly rely on metadata (structures, nature and purpose of network resources) as well as on user requests in terms of domain ontologies. It should be also mentioned the important role played by metadata in semantics-aware multimedia processing. For example, most of existing content-based retrieval systems rely on semantic features, annotations (such as file names, headers, meta descriptors, indexing) and object-based media consideration as well as on their joint usage. The quality-of-service control in distributed networks and particularly in heterogeneous or time-varying networks can also be provided by metadata, describing information about the channel abilities and properties. Since metadata, annotations and other descriptors become an integral part of the image it is necessary to provide the appropriate security features against their intentional and unintentional modifications. For instance, even a simple modification of the image graphical format (from PNG to JPEG, for example) will unavoidably lead to a change of metadata content. A simple

*svolos@cui.unige.ch; phone: +41 22 379-7637; fax: +41 22 379-7780; <http://sip.unige.ch>

change of file name can trick the image search engine (Google, for example) and instead of the targeted data can display prohibited or age censored content. Unfortunately, the limited capabilities of current Web-based environments as well as format incompatibilities do not provide an adequate level of metadata attribute extraction and its sufficient protection.

Therefore, two main issues should be addressed in relation to the design of really secure and intelligent Web-based distributed environments:

- development of novel semantics extraction tools with human interaction in order to assist unsupervised multimedia processing;
- development of innovative techniques for semantic metadata protection against tampering or losses during communications (especially taking into account the growing amount of wireless communications and applications).

This paper advocates a possible solution of the above problems using a specifically designed human-computer interface (HCI) for semantics extraction, and a novel object-based data description tool based on digital data hiding technologies. The first version (to be extended) of the developed semantics extraction interface only uses object-based segmentation with adequate user feedback. The user assigned object descriptions will be integrated into bodies of each semantically segmented image object or region. This approach will provide more security and protection for image information as well as its complete imperceptibility based on robust watermarking. Thus, the images in huge digital databases will be indexed/labeled by their own visual content instead of being annotated by text-based keywords. Moreover the metadata will be completely hidden, thus preventing non-authorized access.

This paper is organized as follows. In the next section we describe an information-theoretic formulation of visual communications via a human-computer interface. In Section 2 we briefly introduce the interactive image segmentation for content-based image extraction as side information for visual communications. The hiding of object-based descriptors into the image is described in Section 3. Section 4 discusses some application scenarios of interactive segmentation with hidden object-based descriptors. Finally, Section 5 provides some concluding remarks.

2. INFORMATION-THEORETIC FORMULATION OF VISUAL COMMUNICATIONS VIA HUMAN-COMPUTER INTERFACE

Human-computer interaction is a very active, multidisciplinary area of research that includes image processing, psychology, information analysis and interpretation, etc. In order to design an optimal and efficient interaction it is important to model and establish the maximum amount of information that can be exchanged. Moreover, such model also has an important impact on the optimization and benchmarking of the human-computer interfaces. We address this problem from information-theoretic point of view considering the HCI as a visual communication channel. The key moment in our analysis consists in the consideration of a visual scene, composed of many objects, as a *multiple access channel* (MAC) [1]. In such a MAC, each object corresponds to a specific “user” that sends its information via the common channel. The observer in the HCI acts as the corresponding decoder that extracts and “decodes” the information sent by each “user”, i.e. carried out by each object within the scene. Obviously, all “users” interfere between each other, which constrains the “extraction” efficiency of the observer. The multiple access channel interpretation of a visual scene is shown in Figure 1.

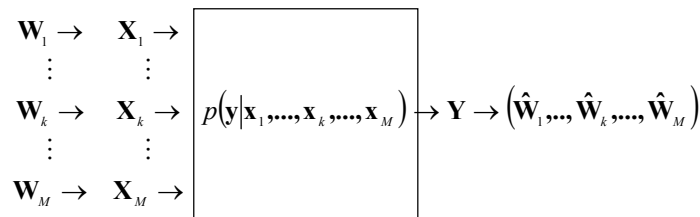


Fig. 1. Multiple access channel interpretation of a visual scene.

Here each index $\mathbf{W}_k$ ($1 \leq k \leq M$) denotes the “state of the object” in the corresponding codebooks $\chi_1, \chi_2, \dots, \chi_M$, where we assume that one observes M -objects within the scene $\mathbf{Y}$. In general, we have K codebooks $\{\chi_k\}_{k=1}^K$ for K possible objects that might appear in the scene. This describes all possible “variations” of the object parameters. Each

codebook $\mathcal{X}_k$ contains 2^{NR_k} independent codewords $\mathbf{x}_k(i)$, $i \in \{1, 2, \dots, 2^{NR_k}\}$ of length N , drawn from i.i.d. distribution $\sim \prod_{i=1}^N p_k(x_{ki})$. Thus, the codebook structure for k -th object can be presented as the following:

$$\mathcal{X}_k = \begin{bmatrix} x_1^k(1) & x_2^k(1) & \dots & x_N^k(1) \\ \vdots & \vdots & & \vdots \\ x_1^k(2^{NR_k}) & x_2^k(2^{NR_k}) & \dots & x_N^k(2^{NR_k}) \end{bmatrix}, (1 \leq k \leq K), \quad (1)$$

where each vector (codeword) $\mathbf{x}_k(i) = \{x_1^k(i), \dots, x_2^k(i), \dots, x_N^k(i)\}$ describes a particular set of stochastic parameters of the object (for example, shape and local variance assuming a parallel Gaussian channel object decomposition [2]).

For a particular “object state” from the first codebook one has a set of parameters that represents shape and local features. These parameters are communicated to the channel and form a resulting composite scene $\mathbf{Y}$. The composite scene can contain M -objects represented by their corresponding codewords. The object composition within a scene can be described by all possible signal processing operations that include different distortions due to the image acquisition process (blurring, addition of noise) as well as geometrical transformations of objects shapes and locations possibly fitted by a class of projective transforms. Therefore, the goal of the decoder consists in the “extraction” of the object features from the observed scene $\mathbf{Y}$. In other words, the decoder should find a jointly typical pair of $\mathbf{y}$ and $\mathbf{x}_k(i)$.

Obviously, one is interested in the reliable extraction from the scene $\mathbf{Y}$ of as much as possible objects. However, a natural question arises: what is the limit of this visual communication channel in terms of objects per unit of observed scene. The classical definition of information defined in the number bits per image pixel is not completely valid here due to the semantic relationship between object features. Therefore, we define the *scene capacity* as the quantity reflecting the amount of objects that can be reliably communicated via a real scene $\mathbf{Y}$ as:

$$C = \max_{p_{\mathbf{x}}(\mathbf{x})} I(\mathbf{X}; \mathbf{Y}), \quad (2)$$

where the maximization is performed over all input distribution of object features.

To provide a basic intuitive explanation of this framework, consider a simplistic case when the scene $\mathbf{Y}$ is composed of the superposition of two objects with feature vectors $\mathbf{X}_1$ and $\mathbf{X}_2$, and contaminated by some acquisition noise $\mathbf{Z}$ (Fig. 2):

$$\mathbf{Y} = \mathbf{X}_1 + \mathbf{X}_2 + \mathbf{Z}. \quad (3)$$

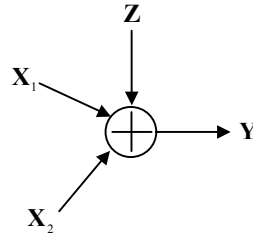


Fig. 2. Composite scene.

The task of the “decoder” is to reliably extract the feature vector $\{\mathbf{X}_i\}$ from the mixture $\mathbf{Y}$ (i.e., to reliably distinguish object 1 from object 2). Obviously, due to the noise and mutual interference between the objects the maximum rate R (or maximum amount of object features that uniquely describe the object) will be determined in each case as $R_1 \leq I(\mathbf{X}_1; \mathbf{Y})$ and $R_2 \leq I(\mathbf{X}_2; \mathbf{Y})$ for object 1 and 2, respectively. One can conclude that if the information about objects is available at the decoder, the rates can be increased to $I(\mathbf{X}_1; \mathbf{Y} | \mathbf{X}_2)$, if $\mathbf{X}_2$ is known, and $I(\mathbf{X}_2; \mathbf{Y} | \mathbf{X}_1)$ if $\mathbf{X}_1$ is known at the decoder. In this case, one can apply the analogy with the MAC and represent the capacity region of this channel as:

$$\begin{aligned}
R_1 &\leq I(\mathbf{X}_1; \mathbf{Y} | \mathbf{X}_2), \\
R_2 &\leq I(\mathbf{X}_2; \mathbf{Y} | \mathbf{X}_1), \\
R_1 + R_2 &\leq I(\mathbf{X}_1, \mathbf{X}_2; \mathbf{Y}).
\end{aligned} \tag{4}$$

This region is represented in Figure 3 and marked by bold lines. The filled rectangle shows the capacity region without availability of side information at the decoder, while the unfilled shapes represent possible gains in capacity that can be obtained using side information.

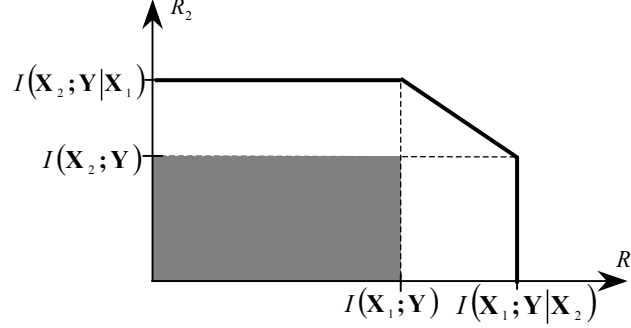


Fig. 3. Achievable rates for a composite scene $\mathbf{Y}$ with and without side information.

In the case of the HCI, the auxiliary (or side) information about object features can be obtained via direct interaction with the human being. This should lead to the separation of objects (features) and in the end to increased rates. Interactive segmentation is a possible way of providing such kind of information for interference suppression.

3. INTERACTIVE SEGMENTATION AS SIDE INFORMATION

The general problem of object recognition is difficult and often requires a large amount of computing resources, even for locating an object within a single image. In Figure 4 is showed the information-theoretic framework of visual communications with side information via interactive segmentation. We propose a human-computer interface for semantic extraction of the object(s) of interest from a large number of objects within the scene. It is natural that humans are much better than computers at extracting semantic information from images and at interpreting image objects and their relationships. The user's feedback will definitely simplify the complexity of the image search and allow more accurate retrieval of the selected objects from huge image databases. Moreover, only human feedback can provide an adequate correspondence between the content of multimedia data and its description (i.e. object and its annotation).

For that reason we propose to use the interactive segmentation results as some prior knowledge (side information) to improve the reliability of visual communication. The main idea of the interactive segmentation is to make available the side information about the objects of interest to the user (database); this approach will allow to significantly decrease the system complexity through the reduction of the search space dimensionality. Here, by interactive segmentation we mean selected region(s) or object(s) of interest outlined by the user in an image that he desires to find in the database; this selection is considered as side information about the object(s).

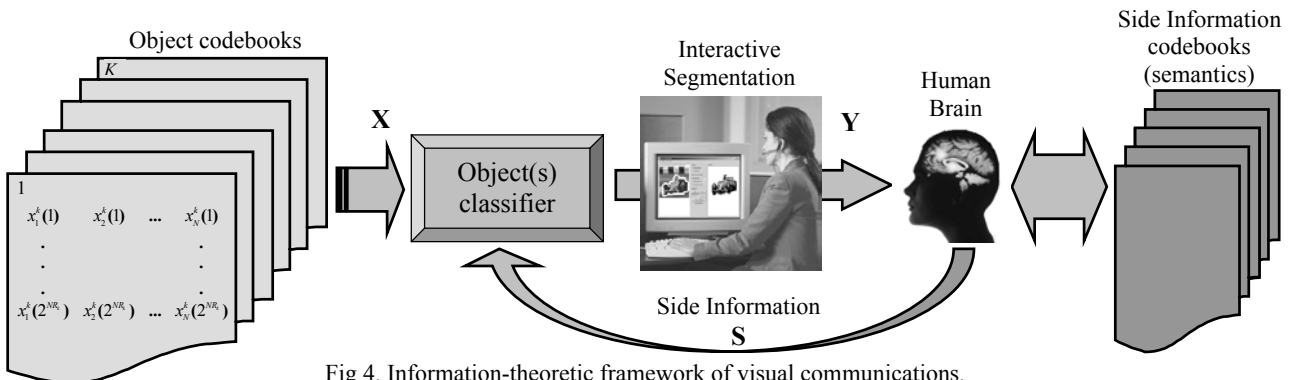


Fig 4. Information-theoretic framework of visual communications.

Image segmentation is well studied but still remains an ill-posed problem and the notion of “correct” segmentation is essentially dependent on the application. Thus, for object-based segmentation purposes, we want to extract the semantically connected object of interest as one segment. Here, we are not focused on particular segmentation algorithms (for example, [3-6]); we rather refer to the fact that successful multimedia retrieval critically depends on the chosen segmentation technique based on region or object content. The successful segmentation of the image is heavily dependent on the criteria used for the merging of pixels based on the similarity of their features, and on the reliability with which these features are extracted. An unsupervised stochastic segmentation is used as a first iteration [7] and the result is displayed for the user for further semantically meaningful segmentation in order to make an adequate correspondence of each image object to its integrated metadata content.

The HCI is intuitive enough and is organized in such a way that the user assistance allows to easily add/remove some objects or parts of objects from previous stages and select the objects of interest. Moreover, the user can even merge some objects that are classified as distinct according to their statistical properties but in fact represent parts of the same semantically connected object (see Fig. 5). Such “semantic” segmentation will assist to easily identify the desired object(s) of interest by a search engine.

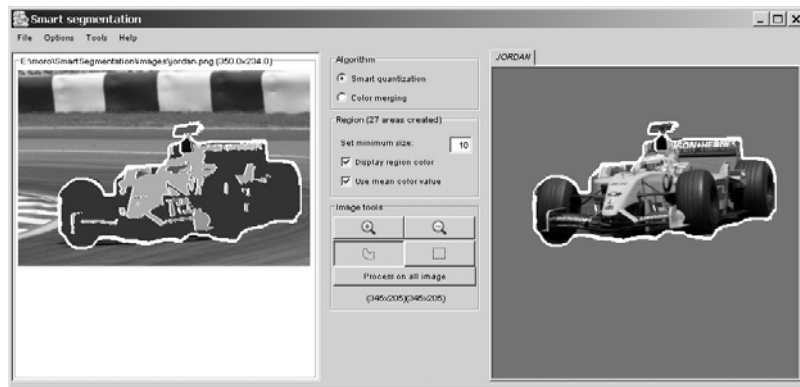


Fig. 5. The result of interactive “semantic” segmentation (left car image) and the object of interest interpretation (right car image) by HCI. The selected object (car) is labeled as “JORDAN”.

Since different people may interpret the same image content in completely different ways this may cause ambiguity and mismatches in the extraction process and in the management of multimedia information. The relevance of visual content significantly depends on the subjectivity of both a database provider and a user. For that reason, the HCI allows to label the image and its objects by their own semantic content, which is hidden into corresponding object, instead of being annotated by text-based keywords in image headers, tags or extra metadata. Generally, we do not consider an image as a set of pixels but rather as a set of annotated regions that correspond to objects or parts of objects, where these objects are associated with some hidden descriptive text about their features. Since each selected (or extracted) object of the image carries its own embedded description, this makes it self-containing and formally independent from the particular image format used for storage in image databases. This is the main idea of the “Smart Media” concept. This concept can easily provide “portability” of each object from one image to another, as well as makes it resistant to object insertion or/and cropping. Moreover, the image can be moved from one database to another without the need of any associated descriptions precisely because object features and descriptors are self-containing in the object. How to embed these descriptors is described next.

4. OBJECT-BASED HIDDEN DESCRIPTORS

The process of grouping low-level features into meaningful image objects and then attaching semantic descriptions to image objects of interest is still an unsolved problem in image understanding [8]. Much work has been done in the area of image feature indexing, especially dominant color indexing, shapes, textures, etc. The existing multimedia database indexing and searching systems [9-12] cannot automatically provide the description of multimedia content

using both low-level (dominant color, shape, texture) and semantic (objects, people, places) descriptors [13]. The proposed approach is based on the hidden indexing/labeling of image objects by using robust digital data hiding technique [14]. In this paper we are focused neither on the encryption and hiding of text descriptors into the image nor on the robustness of data hiding algorithms against some attacks [15]. Here, we rather discuss the possible ways where and how to integrate the annotations reliably into the image and then, how to extract it without errors. We intend to embed the semantically segmented map as well as the user assigned description into the body of each object as the robust watermark. Thus, the image is considered as the set of the "smart objects" separated by their boundaries; each object can be self-indexed, self-authenticated, self-synchronized, self-extracted and self-tamper proofed. It means that the image contains all necessary information about itself including descriptions of the object bodies, of their mutual allocation in images, as well as complementary hidden metadata. It should be noticed, that the integrated information may be based on a secret private key, with which only an authorized party can embed, modify and retrieve it, whereas the rest of the users can only decode it (symmetric protocol en/decryption [7]). Moreover, no additional header, attachment, tag or extra metadata are needed for further image indexing and identification.

In Figure 6 is shown an example of visualization and extraction of the object-based descriptors for the racing car image by means of developed human-computer interface. The previously integrated information about the semantically meaningful object in the car racing image (here "JORDAN") is automatically extracted when the mouse pointer is focused on the object of interest. In this example the semantic content of the object body (car) is presented as an object-based textual descriptor. However, it can be replaced by other hidden multimodal media data depending on the application scenario. As an example, specific interfaces and web-browsers for blind people could associate with each image object its own voice or musical context [16]. In stegano- and cryptographic applications small amounts of textual or/and visual raw-data could easily be hidden into selected "smart objects".

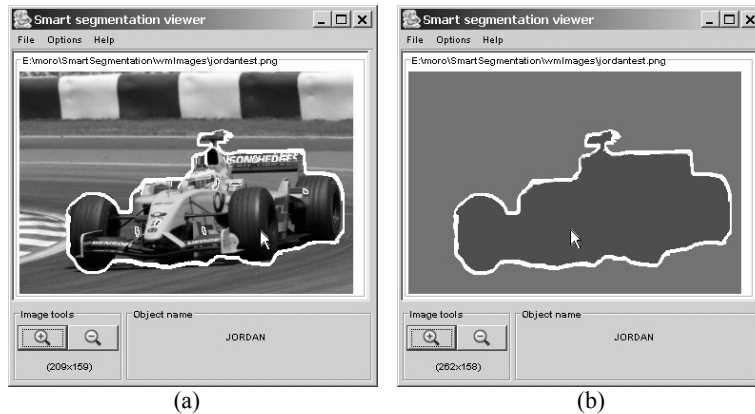


Fig. 6. Example of the visualization (a) and extraction (b) of semantically meaningful object of interest for the car racing image.

After the application of the segmentation technique only an authorized party can extract object(s) of the image to retrieve the hidden descriptors associated with the given object(s) by using a secret private key (symmetric protocol en/decryption). However, this information has to be more robust and secure to possible different image (object) changes and transformations, while remaining perceptually invisible. Thus, the synchronization of the resulting "semantic" segmentation is required for both the descriptor embedding and extraction processes. This can be achieved by the additional embedding into the same object of the raw data describing the contour of the object body. The invariance to different types of distortions can be achieved by providing additional embedding of hidden information into image objects, robust to such changes. Therefore, for robustness reasons the object-based descriptions can be additionally encoded by error correction codes ([17] or [18]) depending on the applications. To avoid possible multiple insertions of contradictory information into the same image object we consider several scenarios, which can be applied for the embedding of descriptors. As the first scenario, we propose to perform a preliminary detection before the insertion of text information into the selected image object. If the preliminary extraction of the object-based descriptor is successful no additional information will be embedded into the object. Another possible scenario is complete replacement, where the previously embedded information related to the object is completely replaced by new data of the user request. In such case, when it is necessary to add complementary information to the embedded description we propose to check the

“free space” for the location of additional raw pixel data. If this space is sufficient to perform the additional embedding, one can also completely replace the integrated information or leave it unchanged depending on the scenario.

The proposed approach of object-based hidden descriptors is invariant to any changes of image filename or/and image headers, and is resistant to object cropping/insertion operations, which are usual in multimedia processing and management. It is well harmonized with the “Smart Media” concept where the image contains information about itself. This information is not integrated inside image headers, and is perceptually invisible and cryptographically secure. One can consider this approach as a reliable joint distribution of visual and textual description (image and text) in visual communications.

5. APPLICATION SCENARIOS

In order to demonstrate how the presented approach can be applied in practice we consider the HCI and interactive segmentation results as side information for two application scenarios. The first one is a telemedicine application; which for confidentiality reasons is based on a symmetric protocol. A physician working with the MRI medical image of a patient after having used the interactive segmentation tool for the detection of possible tumor or clot of blood will be able to insert the necessary information directly into the image or into some parts of it. The region of interest in the MRI image will have its own label with hidden description connected to the extracted object, and the private secret key will be used for information embedding. Even if the patient or any unauthorized party recovers his MRI data he will be unable to extract this description without the secret key. Hence, the patient can give this image to another physician, who will be able to directly extract this invisible hidden information using the same segmentation tool and the proper private key. Therefore, time and money will be saved as well as security and confidentiality of this protocol will be provided.

As the second application scenario, we used several images from an official Formula 1 car racing website [19] (Fig. 7 a,b) as a part of the database for image retrieval. For every image a local unsupervised region-based segmentation was performed [20]. Then, we selected and merged together small-segmented objects into one semantically meaningful object (car) by using the human-computer interface (white contours). The marked “small objects” were then labeled by a corresponding textual description (the number of captured video frame, the lap, the number of car-leader and time stamp) and integrated into the objects by using a symmetric encryption procedure (Fig. 7a). This information was used for further object indexing and search in content-based databases.

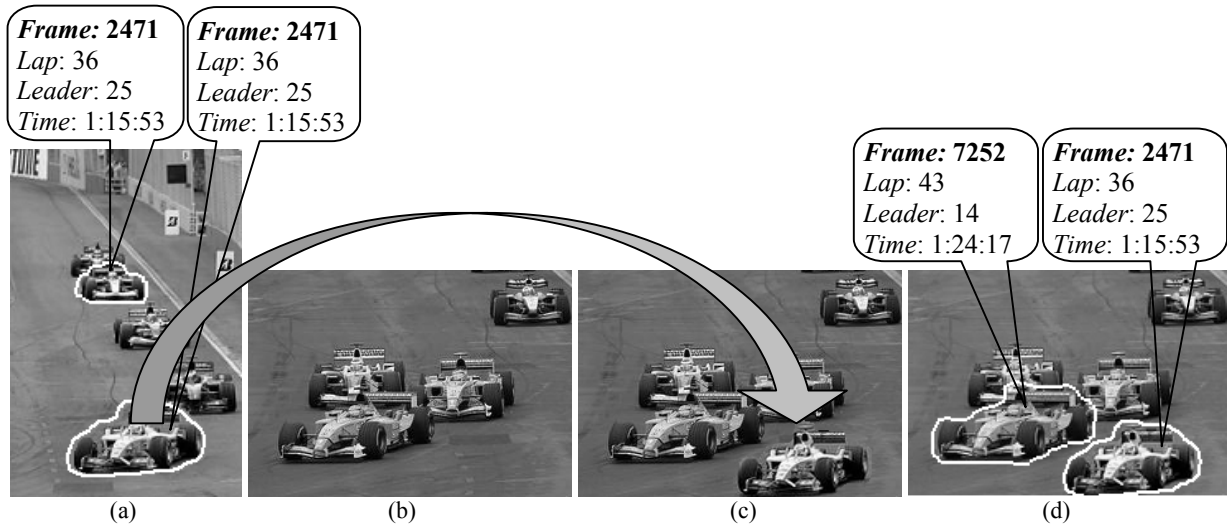


Fig. 7. Hidden object-based descriptions for the car racing images: (a) original image with integrated descriptors; (b) target image; (c) target image with inserted object; (d) the resulting image with extracted hidden descriptors.

For testing reasons one of the marked cars was copied and inserted into another image (see Fig. 7c). Since the descriptor was integrated and hidden into the image object, the information was simultaneously “transported” with this object to the target image (Fig. 7b), while the image quality of the object was completely preserved. During the retrieval

stage the same hidden object-based description was successfully extracted from both images by using the same secret key. Since the identical object (car) is present in both images, the only possible way to distinguish and retrieve the right information from an image database is based on the proper hidden descriptor of this object, but not on the visual data itself. By the corresponding textual description we mean for example the number of captured video frame, the lap, the number of car-leader or time stamp (Fig. 7d), which can be used to prove the originality of the photograph and of its content. In general, it may be an information about the object, its identification number, date and time of image creation, ownership, spatial size of given object, short information about the image to which this object belongs, etc. Here we were not focused on copyright protection or tamper proofing, we rather demonstrated the easy way in which the object information can be kept reliably and properly with complete preservation of the image quality.

6. CONCLUSIONS

In this paper a “Smart Media” concept was proposed for content management and reliable communications. The information-theoretic framework of visual communications via human-computer interface was discussed. An important contribution of this work is to allow a user to easily incorporate semantics into selected image objects in order to provide their additional protection and security. Besides, the presented concept can provide the “portability” of each object from one image to another, as well as made it resistant to object insertion or/and cropping. This paper advocates a possible solution of the above problems using specifically designed human-computer interfaces for semantics extraction and novel object-based data description tools based on digital data hiding technologies.

ACKNOWLEDGMENT

This work was partially supported by the Swiss National Center of Competence IM2 – Interactive Multimedia Information Management and SNF Professorship grant No. PP002-68653/1. The authors are thankful to Oleksiy Koval and Marc Kalberer for many helpful and interesting discussions.

REFERENCES

1. T.M. Cover and J.A. Thomas, *Elements of Information Theory*, Wiley, New York, 1991.
2. T.A. Ramstad, S.O. Aase, and J.H. Husoy, *Subband Compression of Images: Principles and Examples*, Elsevier, New York, 1995.
3. N.R. Pal, S.K. Pal, “A review on image segmentation techniques”, *Pattern Recognition*, **26**, **9**, pp. 1277-1294, 1994.
4. Y.J. Zhang, “Evaluation and Comparison of Different Segmentation Algorithms”, *Pattern Recognition Letters*, **18**, **10**, pp. 963-974, 1997.
5. Y. Rui, T.S. Huang, S.-F. Chang, “Image retrieval: Past present and future”, *Journal of Visual Communication and Image Representation*, **10**, pp. 1-23, 1999.
6. F. Kurugöllü, B. Sankur, E. Harmanci, “Multiband Image Segmentation Using Histogram Multithresholding and Fusion”, *Journal of Image and Vision Computing*, **19**, **13**, pp. 915-928, 2001.
7. Y. Rytsar, S. Voloshynovskiy, and T. Pun, “Metadata representation for semantic-based multimedia security and management”, *Workshop on Metadata Security OTM'03*, Catania, Italy, Nov. 3-7, 2003.
8. Y. Tao and W.I. Grocky, “Object-based image retrieval using point feature maps”, *Proc. of the 8th International Conference on Database Semantics (DS-8)*, pp. 59-73, Rotorua, New Zealand, January 1999.
9. V.E. Ogle, M. Stonebraker, “Chabot: Retrieval from relational database of images”, *Computer*, **28**, **9**, pp. 40-48, 1995.
10. J.R. Smith, S.-F. Chang, “VisualSEEK: A fully automated content-based image query system”, *Proc. of ACM Multimedia '96*, 1996.
11. T. Gevers, A.W.M. Smeulders, “Pictoseek: combining color and shape invariant features for image retrieval”, *IEEE Trans. on Image Processing*, **9**, **1**, pp. 102-119, 2000.

12. H. Müller, W. Müller, D. McG. Squire, S. Marchand-Maillet, T. Pun, "Performance Evaluation in Content-Based Image Retrieval: Overview and Proposals", *Pattern Recognition Letters (Special Issue on Image and Video Indexing)*, Eds. H.Bunke and X.Jiang, **22**, **5**, pp. 593-601, 2001.
13. Y. Xu, E. Saber, A.M. Tekalp, "Object Segmentation and Labeling by Learning from Examples", *IEEE Trans. on Image Processing*, **12**, **6**, pp. 627-638, 2003.
14. S. Voloshynovskiy, F. Deguillaume, O. Koval, T. Pun, "Robust digital watermarking with channel state estimation: Part II Applied robust watermarking", *Signal Processing*, 2003. (*accepted*)
15. F. Deguillaume, S. Voloshynovskiy, T. Pun, "Secure hybrid robust watermarking resistant against tampering and copy-attack", *Signal Processing*, **83**, **10**, pp. 2133-2170, 2003.
16. P. Roth, T. Pun, "A multimodal system for the non-visual exploration of digital pictures", *Interact 2003 "Bringing the Bits together"*, 9th ICIP TC13 Int. Conf. on Human-Computer Interaction, Zurich, Switzerland, Sept. 1-5, 2003.
17. Block Turbo Codes home page: <http://www-sc.enst-bretagne.fr/turbo/principale.html>
18. R.G. Gallager, *Low Density Parity Check Codes*, Monograph, M.I.T. Press, 1963.
19. Australian Grand Prix 2003 postcards website: <http://www.formula1.com/gallery/images/2/Sunday/1949.html>
20. Y. Deng, and B.S. Manjunath, "Unsupervised segmentation of color-texture regions in images and video," *IEEE Trans. on Pattern Analysis and Machine Intelligence (PAMI '01)*, **23**, **8**, pp. 800-810, 2001.