



PROGRESSIVE STUDIES OVER SPARSENESS OF LEAST SQUARES SUPPORT VECTOR MACHINES BASED ON HYPOTHETICAL ASSUMPTION

Machine Learning

Amitesh Kumar Singam

Postgraduate, IEEE Signal Processing Society.

Bharat Reddy Ramancha

Master of Science.

Venkat Raj Reddy Pashike

Master of Science.

Muhammad Shahid

IEEE, Ph.D.

ABSTRACT

No Reference (NR) Video Quality Assessment is the one which is most needed in situations where the handiness of reference video is partially available which is our Hypothetical assumption due to issue raised by reviewers since we used DCT Coefficients in our past research work. Our research work explores the tradeoffs between quality prediction and Video compression. Therefore, we implemented least square support vector regression algorithm as NR-based Video Quality Metric (VQM) for quality estimation with simplified input features based on DCT coefficients (Hypothetical assumption). We concluded that our proposed model overcame sparseness due to hypothetical Assumption.

KEYWORDS

VQM, MOS, LSSVM

INTRODUCTION

Based on Issue raised by reviewers in my past research work, one of the key characteristics of our work is not quality of experience (QoE) its User Experience (UX) for each test condition as observed by the end user. Quality of visual media can get degraded while capturing, storing, transmission, reproduction and display due to distortions which might occur at any of these stages. The true judges of video quality are humans as end users of the video services. The scientific process of evaluation of video quality by humans is called subjective quality assessment. However, subjective evaluation is often too inconvenient, time-consuming, expensive and it must be done by following special recommendations in order to produce reproducible and standard results. These reasons give rise to the need of some intelligent ways of automatically predicting the perceived quality that can be performed swiftly and economically.

Generation of Video Sequences

The encoding details of CIF and QCIF video sequences are mentioned in above tabular column. Out of 120 videos, the video sequences encoded at 15fps are temporally up scaled to 30fps by repeat frame method and QCIF videos are spatially up scaled to CIF using Bi-cubic Interpolation method with the help of Virtual dub. Finally, subjective analysis was conducted for 120 test sequences at temporal resolution of 30fps and spatial resolution of 352x288(CIF).

Mainly, the proposed method involves extraction of visual quality relevant bitstream parameters and building of a machine learning based model for quality prediction. These parameters were selected carefully to keep the complexity in control and to get reasonable coding information which can represent the coding distortions. Following is the description of the extracted parameters and the rationale of making a parameter a part of the proposed model.

Structural Information: In H.264, the 16x16 macro block (MB) can be sub partitioned into blocks of 4x4, 4x8, 8x4 or 8x8, 8x16, 16x8 depending on the coding mode being chosen for the sake of minimal error in prediction. These values provide an estimate of the structural information of a video. The occurrence of each of these alternatives was taken as a percentage of the total blocks.

Spatial and Temporal complexity: In H.264/AVC, a choice between 4x4 or 16x16 intra coding mode (and 8x8 for high profiles) is made based on texture complexity present in the video frame. The ratio of 4x4 blocks out of total blocks in an intra frame would give an estimate of spatial complexity. And the ratio of intra blocks out of the total blocks in an inter (B or P) frame would give an estimate of temporal complexity present in a video sequence.

Motion Contents: A set of motion vector (mv) based statistics were calculated to characterize the motion contents of a video. Ratio of zero mvs out of total mvs in a frame and average mv value in a frame were calculated. Moreover, two indicators of motion intensity were calculated. As in H.264/AVC difference in motion vector is sent and not the absolute value of mv; we also included this information as an estimate of change in motion.

Coding Distortion: Quantization is the main source of distortion introduced during the video encoding and QP value is the parameter which is used to steer the scale of quantization. Hence, average QP value used for encoding a video was included in the prediction model.

Feature Extraction of H.264 Bitstream Data

The feature extraction process for H.264 coded bitstream data was performed in two main steps. First the encoded video bitstreams were decoded using a modified version of JM reference software 16.1 in order to generate an XML file of coding parameters for each video sequence. These XML files contained video information at macroblock level such as quantization parameters, absolute and difference motion vectors, and the type of macroblocks. In the next step a Java program was developed to analyze and process the large XML data in order to extract 18 selected features of the coded videos at frame level. These features which are expected to have high correlation with the perceptual quality of the videos are summarized as follows.

Subjective Quality Assessment

In our past research work, we have considered the recommendations given by ITU-R BT 500-12[1] lab setup of our experiments. Particularly, the method followed was Single stimulus continuous quality evaluation (SSCQE) where a test video sequence is shown once without presence of any explicit reference, corresponds to the reality where users see only the processed version of videos. The tests were conducted in a lab set-up designed in accordance with standards.

DATA PREPROCESSING

After conducting the subjective experiments for each test condition to validate our hypothetical assumptions we need perform pre-processing which is based on User experience not quality of service in our case its Leave one out Cross validation based on least squares support vector machines.

Test Methodology of LSSVM Model

We employed leave one out cross validation technique for estimating the performance of predictive model, leave one out cross validation method is functions similar as K-fold CV, where K is equal to total number of data points. In each round of LOOCV, single data point is

used for validation while remaining data points are used for training and this procedure is repeated such that each of all data points should involve once for the model validation [2].

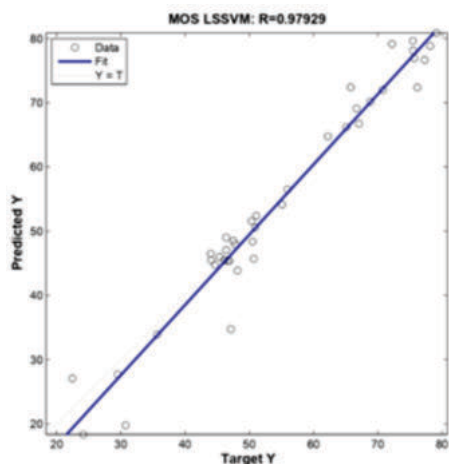
HYPOTHETICAL ASSUMPTION BASED ON CORRELATION MATRIX

The below correlation matrix illustrates the feature extraction of hypothetical assumptions as mentioned in above section

	X1	X2	X3	X4	X5	X6	X7	X8	X9	X10	X11	X12	X13	X14	X15	X16	X17
X1	1	-0.2	-0.19	-0.2	-0.17	-0.06	0.15	0.08	0.09	-0.05	0.11	0.14	0.03	-0.04	0.43	0.18	0.19
X2	-0.28	1	0.52	-0.53	-0.49	0.36	-0.17	-0.56	0.31	0.49	0.33	0.23	-0.84	-0.72	0.48	0.54	0.53
X3	0.2	0.52	1	-1	-0.62	0.26	-0.03	-0.23	0.23	0.27	0.24	0.15	-0.6	-0.5	0.65	0.65	0.99
X4	-0.2	-0.5	-1	1	0.62	-0.26	-0.03	0.22	-0.2	-0.27	-0.24	-0.15	0.62	0.5	-0.65	-0.64	0.99
X5	-0.15	-0.38	0.63	0.62	1	-0.69	-0.7	-0.3	-0.1	-0.16	-0.3	-0.28	0.69	0.74	-0.43	-0.54	0.61
X6	-0.06	0.63	0.26	-0.27	-0.69	1	0.29	-0.16	0.05	0.34	0.34	0.37	-0.6	-0.7	0.44	0.44	0.27
X7	0.16	-0.17	0.03	-0.03	-0.72	0.29	1	0.83	-0.1	-0.17	0.13	0.10	-0.2	-0.38	-0.12	0.07	0.03
X8	0.09	-0.56	-0.2	0.23	-0.32	-0.15	0.83	1	-0.25	-0.32	0.01	-0.06	0.21	0.08	-0.4	-0.2	-0.2
X9	0.09	0.3	0.22	-0.2	-0.05	0.05	-0.13	-0.16	1	0.70	0.73	0.57	-0.25	-0.23	0.48	0.67	0.22
X10	-0.05	0.49	0.27	-0.28	0.16	0.34	0.17	-0.33	0.71	1	0.71	0.67	-0.3	-0.35	0.63	0.73	0.27
X11	0.12	0.33	0.24	-0.24	-0.33	0.34	0.13	0.01	0.74	0.7	1	0.73	-0.3	-0.49	0.57	0.88	0.24
X12	0.15	0.2	0.15	-0.15	-0.28	0.37	0.10	-0.06	0.57	0.67	0.73	1	-0.2	-0.37	0.56	0.63	0.15
X13	0.03	-0.8	-0.6	0.6	0.69	-0.6	-0.2	0.29	-0.2	-0.29	-0.32	-0.17	1	0.9	-0.49	-0.6	-0.6
X14	-0.04	-0.7	-0.56	0.5	0.7	-0.7	-0.37	0.08	-0.25	-0.34	-0.46	-0.37	0.9	1	-0.5	-0.6	-0.5
X15	0.44	0.48	0.65	-0.65	-0.43	0.44	-0.11	-0.41	0.48	0.63	0.57	0.56	-0.5	-0.52	1	0.8	0.65
X16	0.18	0.54	0.64	-0.64	-0.53	0.44	0.07	-0.16	0.67	0.73	0.88	0.65	-0.64	-0.63	0.8	1	0.64
X17	0.2	0.53	0.99	-1	-0.62	0.27	0.03	-0.23	0.23	0.27	0.24	0.15	-0.64	-0.52	0.65	0.65	1

Table 4.3: Correlation matrix of Extracted Features

Statistical Analysis



The above plots illustrate correlation coefficient for predicted values or residuals after regressing with true values Y of MOS.

CONCLUSIONS

We concluded that our proposed model overcame sparseness due to hypothetical Assumption.

REFERENCES:

- [1] ITU-R Radio communication Sector of ITU, Recommendation ITU-R BT.500-12," 2009.
- [2] Suykens J. A. K, Vandewalle, J, Least squares support vector machine classifiers", Neural Processing Letters, 9 (3), 293–300
- [3] www.esat.kuleuven.be/sista/lssvmlab, Least squares support vector machine Lab (LSSVmlab) toolbox