

Encoding and Retaining Information in the Visuals
and Verbals of an Educational Movie

Patricia Baggett, Institute of Cognitive Science
and
Andrzej Ehrenfeucht, Department of Computer Science
University of Colorado

September, 1981

Technical Report #108-ONR

ABSTRACT

Viewers watching a narrated movie are simultaneously presented information in two media, visual and verbal/auditory. This study shows there is no competition for resources in an educational movie: when one is encoding information in one medium, one is not hindered from encoding information in the other. Even when the visual and linguistic information are presented sequentially, doubling study time, no more information is extracted than in an intact movie. College students are good dual media processors. In a sequential presentation, spoken narration first and visuals second is far inferior to visuals first and narration second. When the verbal material in a sequential presentation is read rather than listened to, order does not matter. Regarding retention, much information is extracted from linguistic material, but only half remains after a week. Less information is extracted from visual material, but it stays over a week. Practical applications are discussed.

Encoding and Retaining Information in the Visuals and Verbals of an Educational Movie

In watching a narrated movie, a person receives simultaneous information in two media, visual and verbal/auditory. This report examines how well college students encode the information in the visuals versus that in the verbals, of an educational movie, and how well they retain information from the two different media over a delay.

Work in information extraction from film and other dual media presentations, such as pictures and words, has been done by many authors, including Baker and Popham (1965); Dwyer (1968); Hochberg (1978); May and Lumsdaine (1958); Olson (1974); Peeck (1974); and Salomon (1979). However, no previous study has dealt with the issues which will be looked at here.

The study investigates two main topics. The first is a comparison of encoding and retention of visual versus linguistic information. The linguistic information will be studied in two ways, either auditorily, by listening to the film's soundtrack, or by reading it as written text. The second topic is the order of presentation of the visual and linguistic information. Three orders will be investigated: (1) synchrony, as in an intact movie with soundtrack; (2) the movie's visuals, played silently with the soundtrack turned off, followed immediately by the verbals with the visuals turned off; and (3) verbals with visuals turned off, followed immediately by visuals shown silently.

Information obtained from these different stimulus conditions, and from the conditions of visuals only or verbals only, will be compared to that of a control group which is given no stimulus presentation.

For convenience in terminology, the expressions linguistic information and verbals will be used as synonyms, and will mean either text or narration. Text will mean written text, taken verbatim from the film's soundtrack. Narration will

mean the film's auditory soundtrack. Visuals will mean the film's moving pictures, shown silently. Movie will mean narration and visuals in synchrony. The test used to evaluate the information obtained will be given either at zero delay, which will mean immediately after presentation of the study material, or after a seven day delay, which will mean a week after presentation of the study material.

Combining visual and linguistic information in all possible stimulus presentations, with tests at zero and seven day delay, yields 17 conditions. They are:

0. No information (group given no stimulus presentation).
1. Text - zero delay (T-0).
2. Narration - zero delay (N-0).
3. Visuals - zero delay (V-0).
4. Text first; visuals second - zero delay (TV-0).
5. Narration first; visuals second - zero delay (NV-0).
6. Visuals first; text second - zero delay (VT-0).
7. Visuals first; narration second - zero delay (VN-0).
8. Movie - zero delay (M-0).
- 9-16. Identical to 1 through 8, except the test is given seven days after the stimulus presentation.

The study time is different in different groups. Single presentations (groups 1, 2, and 3; and 9, 10, and 11) and synchronous presentations (8 and 16) have a study time of 11 min. Sequential presentations (4, 5, 6, and 7; 12, 13, 14, and 15) are studied for 22 min.

The study answers three specific questions:

- 1a. A movie presents visual and narrative information simultaneously. Does simultaneous presentation lead to poorer encoding of information presented by each medium (visual and narration) than when the information from the two media is

presented sequentially? Such a finding could be an example of competition for resources. It would mean that when a person is encoding information from one source, the person is hindered in encoding information at the same time from another source.

1b. Is there an increase of information extracted when the study time is doubled in the sequential presentations?

Information extracted in the movie condition, if it is less than in the sequential presentations, could be less for one or both of two reasons: (a) competition for resources; and (b) shorter study time.

2a. In the sequential presentations, does it matter whether the linguistic information is heard or read?

2b. In the sequential presentations, does order of input (visual first and linguistic second, or linguistic first and visual second) make a difference?

3. What is the effect of delay on information obtained from different media? Is it the same for visual and linguistic information, or different?

The answers to these questions have practical applications which will be discussed, about how to present dual media educational material for good encoding and good retention.

METHOD

Subjects

459 students in introductory psychology classes at the University of Colorado in Boulder and the University of Denver participated in partial fulfillment of a course requirement. They were randomly assigned, in small clusters of two to eight, to one of the 17 groups described in the Introduction. Groups ranged in size from 25 to 30 people.

Materials

The film used is Plant Traps: Insect Catchers of the Bog Jungle, copyright 1954, distributed by Encyclopedia Britannica Films. It is 16mm sound and color,

11 min long, with 1270 words of narration. The film was chosen because it is about an interesting topic (carnivorous plants) with information that is new for most people; it is visually exciting with time lapse and extreme closeups; and it is appropriate, according to the distributor, for junior high through college age viewers.

Sixty three questions on carnivorous plants were written by the experimenter; 20 true-false, 17 multiple choice, and 26 short answer. (The questions and the film's narration, are available from the first author.) Percentage correct on these questions was the dependent measure. One question thus accounts for $1/63 = 1.6\%$ of the score.

Examples of the three types of questions are:

True or false: Sundew plants are more active toward live than toward dead prey.
(Answer: True).

Multiple choice: How fast can a healthy Venus flytrap shut? Pick the most accurate answer. (a) in less than 1/10 sec; (b) in less than 1/2 sec; (c) in less than 3 sec; (d) in less than 10 sec. (Answer: b).

Short answer: What attracts insects to the pitcher plant? (Answer: perfume).

Procedure

Subjects in all groups except the no information control group were told to watch the movie (or read the text or listen to the narration, or look at the visuals with the soundtrack turned off, etc.). They were told before their presentation that their task afterwards would be to answer 63 questions about carnivorous plants.

Subjects in the sequential presentations studied the input 22 min, 11 min for each of two modalities. Subjects in the single presentations and in the movie conditions studied it only 11 min. Results from the movie versus sequential presentations will answer question 1 in the Introduction.

Subjects given the text to read were told they could read it as many times as they liked in the 11 min period and that they could underline phrases or use whatever strategy they chose to learn the information.

Each subject was given a deck of 63 numbered questions, each on a 3 in x 5 in card. The order of the questions was the same for each person. The control group studied no input but was asked to try to answer the questions. Subjects in the other groups were given the questions after study and at the appropriate delay (zero or seven day).

Comparing scores from the sequential presentation groups will answer question 2 in the Introduction, and comparing scores from the zero and seven day delay groups will answer question 3 in the Introduction.

A few answers to earlier questions had to be given in the phrasing of later questions. Therefore, questions were placed in an envelope, and subjects were specifically instructed to take out the inverted deck, turn over the top card, question 1, and answer it on the card or leave it blank, and to return it to the envelope. They were then to turn over question 2, etc. They were told that once they had placed a question in the envelope, they could not return to it and change their answer. They were instructed that there was no penalty for guessing.

Subjects were run in small groups to enforce these instructions. Time to complete the task varied from 25 to 45 min.

RESULTS AND DISCUSSION

Scoring the 63 questions was done as follows. The 20 true-false and 17 multiple choice questions were objectively scored with full credit given for the right answer and no partial credit. Answers for the 26 short answer questions were decided on by the experimenter and written down with examples of variations in the answers and the amount of credit to be given for each variation specified. She and a second

experimenter then scored the 26 short answer questions blind. Where there was disagreement, in less than 5% of the cases, a discussion was held until agreement was reached.

A person answering completely at random would score 21.04% correct by chance. Chance level is calculated from the true-false and multiple choice questions only.

Table 1 shows the mean percentage correct on the 63 questions for each of the 17 groups, the standard deviation, and the number of subjects in each group.

Insert Table 1 About Here

To account for the data, we chose an additive model, which works as follows. To each group, we attribute some number of hypothesized features. A particular feature is therefore either present or absent for a particular group. Each feature has a numerical value, either positive or negative. Each group's percentage correct on the questions is the sum of the values of the features that are present in the group.

The actual values for the features are determined by the method of least squares (Hays, 1963). In the case of the scores given here, we did not a priori know what features to choose. The problem was to find a set of interpretable features that explain the data within experimental error.

The features chosen are shown in Table 2. The presence of a feature for a group is represented by a 1 in the feature's column; the absence is represented by a 0.

Insert Table 2 About Here

The theoretical values derived from the least squares fit for the five features, and the names given to the features a posteriori are:



feature 1 = 37.52 (baseline)

feature 2 = 10.92 (linguistic recency)

feature 3 = 9.43 (linguistic)

feature 4 = 9.00 (visual)

feature 5 = -4.75 (penalty for spoken narration, except when in synchrony
with visuals)

These values are the amounts of a group's total percentage correct that can be attributed to each feature, when the feature is present in the group.

A group's theoretical value can be computed from the matrix in Table 2 and the feature values above. For example, group 5, narration-visuals, zero delay, has four features present, as shown in Table 2: Baseline, linguistic, visual, and penalty for spoken narration. Therefore, its theoretical value is $37.52 + 9.43 + 9.00 - 4.75 = 51.20\%$. (Its actual value is 50.19%.)

Table 3 gives the actual and theoretical values for each group score, and the difference between the two. Using one sample t-tests, none of the actual group

Insert Table 3 About Here

means is significantly different from its theoretical mean. Therefore, the hypothesis that each of the 17 group scores consists of the sum of the values of the features present in that group cannot be rejected.

Table 2 shows that all groups have feature 1 (baseline) present, for a value of 37.52%. Feature 2, linguistic recency, is present in all zero delay groups with linguistic input except NV-0. The value for linguistic recency is 10.92%. The linguistic feature, number 3, with a value of 9.43%, is present in all groups with linguistic input. Feature 4, visual, is present in all groups with visual

input, with a value of 9.00%. Finally, feature 5, a penalty for spoken narration (but not written text), is present in all groups with (spoken) narration except the movie groups, which have narration and visuals in synchrony. Its value is -4.75%. An interpretation of this particular assignment of features, and their values, will be given below.

The important question was which features to use to explain the data. The number of possible features that might have been chosen is 2^{17} , but only 5 were selected. Examples of two features not used in the analysis are:

- (1) A feature for the movie; this feature would have 1 in M-0 and M-7 and 0 elsewhere.
 - (2) A feature for delay; this feature would have 1 in groups 9-16 and 0 elsewhere.
- The reason for not using some features is not that they are not existent, but that their effect is negligible.

Finding the features presented in this paper was done by the following method. A computer package was prepared which allowed us to check how a given hypothesis (namely, a matrix as in Table 2, or a set of features) fit the data, and to modify the matrix (for example, introduce new features, find what new features give the best fit, or delete features which were irrelevant) to improve the fit. The package was written by R. Michael Perry and implemented on the VAX 11/780 under the UNIX operating system.

INTERPRETATIONS AND CONCLUSIONS

Answers to the specific questions asked in the Introduction will be provided in turn, and practical applications of the findings will be given.

Questions 1a and b:

There is no evidence for competition for resources between visuals and narration in the intact movie, or for an advantage in sequential presentations of doubling the

study time. On the contrary, subjects could both encode and retain visual and narration information occurring simultaneously in the movie even better than they could such information occurring sequentially, even though the sequential information was studied twice as long.

At zero delay, movie subjects scored 68.77%, the highest of any group. The best sequential presentation group with narration was the group with visuals first and narration second, VN-0. They scored 60.93%, significantly lower than the M-0 group, $t_{56df} = 3.81$, $p < .001$. This result shows that college students encode related visual and narration information better when it is presented simultaneously than when it is presented sequentially.

Table 2 shows the difference between the M-0 and VN-0 groups in terms of features. The latter group has feature 5, a penalty for spoken narration when it is not in synchrony with the visuals, whereas the former group does not. Feature 5 has a value of -4.75%. We think that the -4.75% is due to a decrement in encoding caused by misperceived phonemes in the VN-0 condition (and, as a matter of fact, in all listening conditions in which visuals are not simultaneously presented, as can be seen in Table 2). The movie's visuals, occurring either earlier or later than the narration, do not correct the misperceived phonemes. Evidence of such misperceptions was explicit in several answers to short answer questions in the listening conditions: "potion" was written rather than "portion," "foggy" rather than "boggy," "sunview" rather than "sundew." Such misperceptions are not found in the synchronous conditions. We suspect that, when visuals are presented simultaneously with spoken narration, the visuals help to disambiguate spoken words.

Performance in the M-0 condition does not differ from that in the sequential presentations when the linguistic material is text rather than narration. As mentioned before, subjects reading the text were allowed to use any strategy they

chose to learn the information. Nevertheless, there is no evidence for negative interference between visuals and narration in the intact movie.

The highest score for a sequential presentation group when the linguistic material was text was 66.75% for the group with visuals first and text second (VT-0). This score does not differ statistically from that of the M-0 group ($t_{50df} < 1$). That the M-0 group is similar to both the VT-0 and TV-0 groups is shown in Table 2. The features giving the best fit for M-0 are identical to those in the VT-0 and TV-0 groups. There is no difference in encoding between the M-0 and the two sequential presentation groups.

Turning now to retention over a week, the M-7 subjects score 54.17%, which is not significantly different from any of the four sequential presentation groups by two sample t-tests. The sequential presentation group scores are 50.63%, 52.77%, 57.33%, and 57.30% for NV-7, VN-7, TV-7, and VT-7, respectively.

Table 2 shows the difference in features of M-7 versus NV-7 and VN-7: the penalty in NV-7 and VN-7 for spoken narration when not in synchrony with visuals, feature 5, with a value of -4.75%. Still, the actual scores in the 3 groups are not significantly different. Table 2 also shows there is no difference in features in the M-7, TV-7, and VT-7 groups.

The final conclusion is that people retain simultaneously presented visual and narration information as well as they do such information presented sequentially, even when the sequential information is studied twice as long and the subjects are allowed to read the linguistic information as a text and study it any way they like. There is no evidence for competition for resources in encoding or retention for visuals and narration in synchrony. College students are good dual media information processors. An intact movie is an efficient means of transmitting information.

Question 2:

In the sequential presentations it does indeed matter whether the linguistic information is heard or read: spoken narration and written text interact differently with visuals. In particular, written text can be studied before or after the visuals, and the effect is the same. This can be seen in the similar percentages correct for text-visuals and visuals-text at 0 delay (66.25% versus 66.73%) and at 7-day delay (57.33% versus 57.30%). It can also be seen in the matrix in Table 2: groups 4 and 6 have the same set of features, and groups 12 and 14 have the same set of features. The difference between the 0- and 7-day delay groups is a single feature, linguistic recency, with a value of 10.92%. It is present at 0 delay and absent after 7 days.

Something very different happens for spoken narration at zero delay. When it is studied before the visuals, it is far inferior to when it is studied after the visuals. This is shown by the different percentages correct for narration-visuals and visuals-narration at 0 delay (50.19% versus 60.93%, $t_{50df} = 3.89$, $p < .001$). It is also shown in the matrix in Table 2: visuals-narration has a linguistic recency feature, while narration-visuals does not. This means that information in spoken linguistic material is encoded better when the visual material to which it is related is presented first, rather than second. When spoken linguistic material is presented before the visuals, the results are as poor as if the linguistic material were not presented at all. (Visual, zero delay = 47.89%; narration-visual, zero delay = 50.19%, $t_{55df} < 1$).

Framework for Interpreting the Auditory/Visual Interaction

Presented here is a brief overview of a theoretical framework which gives an interpretation of why there is a difference between NV-0 and VN-0, but not between

TV-0 and VT-0. We postulate a single conceptual memory in the form of a semantic network. Stimulus input creates a set of concepts (nodes in a semantic network). Concepts consist of many elements or components from different media, among them auditory and visual.

The differences in the NV-0, and VN-0, and M-0 groups could be analyzed in terms of how the visual component associates with the auditory component. A narrated synchronous film is input that hypothetically causes concepts with both visual and auditory elements, well associated, to be formed. The clear superiority of VN-0 over NV-0 would indicate that auditory components create good associations with visual components presented earlier. The poor performance of NV-0 would indicate that visual components do not create good associations with auditory components presented earlier.

The emphasis here is between auditory and visual. When the linguistic material is presented visually, as in the TV-0 and VT-0 groups, the difference is nonexistent.

The results also fit with a single memory, dual processing hypothesis, in which visual information is processed by one unit (both visual linguistic and visual pictorial), and auditory linguistic information by a separate unit. When a person uses the same processing unit (as in TV-0 and VT-0, where the unit is visual) good associations are created, independent of order of presentation.

In the interaction between auditory and visual processing, it seems that auditory processing occurring later than visual (VN-0) brings in the earlier visual components in forming concepts. But visual processing occurring later than auditory (NV-0) forms concepts with visual components, without bringing in the earlier auditory/linguistic components.

This hypothesis could be tested as follows. During early occurring auditory input, some limited amount of visual input could be presented, to which the later occurring visual input could form associations. Or, people receiving the visual input second could be required to say what they are seeing, which would force the formation of auditory/linguistic elements. If the hypothesis is correct, both of these manipulations should improve performance in the NV-0 group.

Question 3:

A one week delay does indeed affect linguistic and visual material differently. Table 2 shows that a linguistic recency feature, with a value of 10.92%, is present in all zero delay groups with linguistic information except narration-visuals, zero delay. Its value is the highest of any feature other than baseline, and it disappears after a week. A significant visual recency feature was not observed. The linguistic feature which is present at both zero and 7-day delay has a value of 9.45%; the visual feature present at both delays has a value of 9.00%. This study shows that humans are good at storing lots of verbal information ($10.92\% + 9.43\% = 20.35\%$) for a short time, but that less than half of it (9.43%) lasts over a week. On the other hand, visual information, once encoded, is retained over a week.

There may, in fact, be a way to cause the information from the linguistic recency feature to last over a delay. If, during input, better visual/verbal associations could be presented, so that information from the two modalities would be more strongly knitted together, then the longer lasting visual material might be able to be used to retrieve the material from the verbal input.

Final Comments

This study has shown that there is no competition for resources when related information is presented in two media (visual and verbal/auditory) simultaneously.

Therefore, synchronous visual and verbal/auditory input is an efficient way to present information. It is 8% better than presenting the visual information first, followed by the spoken verbal information second, and better by far (18%) than spoken information first followed by visual information second. The advantage of a synchronous presentation, in terms of information extracted, is lost when one compares sequential presentations in which the verbal information is read rather than listened to, at least for the literate college students tested here. Finally, information from visual and verbal sources is encoded and retained differently. Lots of linguistic information is encoded, but only half of it is retained over a week. Far less visual information is encoded, but it all lasts over a week.

The material used in this study was a standard educational film containing scientific facts. We do not know if the results will generalize to other types of materials such as instructions or stories. We also do not know what the effect of a longer delay would be, nor whether different dependent measures, such as free recall or a test with visual material, would give the same results. But the findings of this study answer three important questions and have practical application. Namely, in a show and tell presentation, one should not tell first and show second. To improve encoding and retention, one should either show and tell in synchrony, or show first and tell second.

References

- Baker, P., and Popham, W. Value of pictorial embellishments in a tape-slide instructional program. Audiovisual Communication Review, 1965, 14 (4), 397-406.
- Dwyer, F. M. The effectiveness of visual illustrations used to complement programmed instruction. Journal of Psychology, 1968, 70, 157-162.
- Hays, W. G. Statistics for Psychologists. New York: Holt, Rinehart, and Winston, 1963.
- Hochberg, J. The perception of motion pictures. In E. Carterette and M. Friedman (Eds.), Handbook of Perception, Vol. X, Perceptual Ecology. New York: Academic Press, 1978.
- May, M. and Lumsdaine, A. Learning from Films. New Haven: Yale University Press, 1958.
- Olson, D. R. Media and Symbols: The Forms of Expression, Communication, and Education. The Seventy-third Yearbook of the National Society for the Study of Education, Part I. Chicago: University of Chicago Press, 1974.
- Peeck, J. Retention of pictorial and verbal content of a text with illustrations. Journal of Educational Psychology, 1974, 66, 881-888.
- Salomon, G. Interaction of Media, Cognition, and Learning. San Francisco: Jossey-Bass Publishers, 1979.

Footnote

This work was supported by Office of Naval Research Contract #N00014-78-C-0433 and National Institute of Mental Health postdoctoral fellowship #5 F32 MH07588-02 to the first author. Some of the results were presented at the 1980 annual meeting of the Psychonomics Society in St. Louis. We thank Agda Bearden for helping with data collection and scoring, and R. Michael Perry for implementing the computer package for data analysis. Requests for reprints should be sent to Patricia Baggett, Psychology Department, University of Colorado, Campus Box 345, Boulder, Colorado, 80309. This report is #108 of the Institute of Cognitive Science's Technical Report Series.

Table 1
Mean Percentage Correct on 63 Questions for 17 Groups

Group	Mean Percentage Correct	Standard Deviation	Number of Subjects
0. no information	38.44	6.38	26
1. text-0	56.69	10.78	26
2. narration-0	53.46	12.13	28
3. visuals-0	47.89	8.62	27
4. text-visuals-0	66.25	5.33	27
5. narration-visuals-0	50.19	12.02	30
6. visuals-text-0	66.73	9.11	25
7. visuals-narration-0	60.93	8.92	29
8. movie-0	68.77	6.25	29
9. text-7	45.14	9.30	30
10. narration-7	42.19	6.70	25
11. visuals-7	45.09	8.10	26
12. text-visuals-7	57.33	10.08	26
13. narration-visuals-7	50.63	10.52	25
14. visuals-text-7	57.30	8.57	27
15. visuals-narration-7	52.77	8.38	27
16. movie-7	54.17	7.92	26

Note: 0 = zero delay; 7 = 7-day delay; text = written text; narration = auditory soundtrack. Groups 4, 5, 6, 7, 12, 13, 14, and 15 had sequential input presentations, e.g., group 4, TV-0, read the text first and then saw the visuals with the soundtrack turned off. These groups studied input twice as long as groups 1, 2, 3, 8, 9, 10, 11, and 16.

Table 2
Matrix of Features, Their Names, and Their Values

Feature 1	Feature 2	Feature 3	Feature 4	Feature 5
8 = Baseline	L _R = Linguistic Recency (except NV-0)	L = Linguistic	V = Visual	N _P = penalty for spoken narration except when in synchrony with visuals
Values of Features	37.52	10.92	9.43	9.00
Group				-4.75
0. no information	1	0	0	0
1. text-0	1	1	0	0
2. narration-0	1	1	0	1
3. visuals-0	1	0	1	0
4. text-visuals-0	1	1	1	0
5. narration-visuals-0	1	0*	1	1
6. visuals-text-0	1	1	1	0
7. visuals-narration-0	1	1	1	1
8. movie-0	1	1	1	0
9. text-7	1	0	0	0
10. narration-7	1	0	0	1
11. visuals-7	1	0	1	0
12. text-visuals-7	1	0	1	0
13. narration-visuals-7	1	0	1	1
14. visuals-text-7	1	0	1	0
15. visuals-narration-7	1	0	1	1
16. movie-7	1	0	1	0

Note: 1 means a feature is present in a group's score; 0 means it is absent. *The NV-0 group has no linguistic recency feature.

Table 3
Actual and Theoretical Values for Each Group's Score, and the
Difference Between the Two

Group	Actual Score	Theoretical Score	Difference
0. no information = B	38.44	37.52	.92
1. text-0 = $B+L_R+L$	56.69	57.43	-.74
2. narration-0 = $B+L_R+L+N_p$	53.46	52.69	.77
3. visuals-0 = B+V	47.89	46.95	.94
4. text-visuals-0 = $B+L_R+L+V$	66.25	66.86	-.61
5. narration-visuals-0 = $B+L+V+N_p$	50.19	51.20	-1.01
6. visuals-text-0 = $B+L_R+L+V$	66.73	66.86	-.13
7. visuals-narration-0 = $B+L_R+L+V+N_p$	60.93	62.12	-1.19
8. movie-0 = $B+L_R+L+V$	68.77	66.86	1.91
9. text-7 = B+L	45.14	46.51	-1.37
10. narration-7 = $B+L+N_p$	42.19	41.77	.42
11. visuals-7 = B+V	45.09	46.95	-1.86
12. text-visuals-7 = B+L+V	57.33	55.95	1.38
13. narration-visuals-7 = $B+L+V+N_p$	50.63	51.20	-.57
14. visuals-text-7 = B+L+V	57.30	55.95	1.35
15. visuals-narration-7 = $B+L+V+N_p$	52.77	51.20	1.57
16. movie-7 = B+L+V	54.17	55.95	-1.78

Note: Each group's theoretical score is the sum of the values of the features that are present.