

memory of aversive information, where the MRF is involved in STM and the hippocampus is involved in LTM. (It is assumed that STM is being measured 64 seconds after an FS and that LTM is being measured 24 hours after an FS.) That stimulation of the hippocampus after trials disrupts LTM in rats is consistent with similar results for cats (5). Thus, with the use of brain stimulation after trials it is possible to dissociate STM from LTM, a result that suggests that STM and LTM are operating independently. Further support for this view is provided by Milner (2), who found that patients with bilateral hippocampal lesions had good STM but failed to store new information in LTM, and by Warrington and Shallice (3), who found that a patient with a parietooccipital lesion had deficient STM but good LTM.

Furthermore, when acetoxycycloheximide (AXM), an inhibitor of protein synthesis, was injected before or immediately after a learning trial, the drug interfered with LTM but not STM (6). However, AXM that was injected 15 or 30 minutes after a learning trial did not interfere with LTM (6, 7). The fact that an inhibitor of protein synthesis interferes with LTM if it is given immediately after training but not if it is given 15 minutes later also suggests a parallel processing of information in STM and LTM. Finally, repeated presentations of FS and electroconvulsive shock (ECS) in short intervals (at least 0.5 second between FS and ECS) are sufficient to lead to LTM of the FS (8). Our experiment and the above-mentioned studies indicate that parallel rather than sequential processing of STM and LTM must occur for aversive information.

In another experiment, we tested a retention interval between 64 seconds and 24 hours to determine the time course of decay of STM and growth or consolidation of LTM. Subjects were 19 male Long-Evans rats, divided into groups for MRF stimulation ( $N = 6$ ), hippocampal stimulation ( $N = 7$ ), and no stimulation ( $N = 6$ ). The animals were prepared surgically, trained, and given the FS as in the earlier experiment. The animals received 5 seconds of brain stimulation 196 seconds after the FS; and 55 seconds after the offset of stimulation they were retested for 10 minutes for retention of the FS. The suppression ratios for brain stimulation 196 seconds after FS are shown in Fig. 2, with corresponding data (from the earlier experiment) for brain

stimulation 4 seconds after FS. For both groups the interval between brain stimulation and retest was the same (55 seconds). In the control group there was greater suppression at 256 seconds after the FS than at 64 seconds after the FS. However, at the 256-second retest there was suppression for the MRF group, but little suppression for the hippocampus group.

A two-way analysis of variance revealed that the mean suppression on the retests was influenced by site of brain stimulation ( $F = 14.1$ ; d.f. = 2/34;  $P < .001$ ) and the interaction between site of brain stimulation and time of retest ( $F = 8.88$ , d.f. = 2/34,  $P < .01$ ). A Newman-Keuls comparison test showed that at the 256-second retest, the MRF and hippocampus groups showed significantly less suppression than did the control group ( $P < .01$ ). Furthermore, for the control and MRF groups, suppression was significantly greater at the 256-second retest than at the 64-second retest ( $P < .05$ ); in contrast, for the hippocampus group suppression was significantly less at the 256-second retest than at the 64-second retest ( $P < .01$ ).

Thus, MRF stimulation that is applied 196 seconds after an FS produces marked suppression of bar-pressing at the 256-second retest. If it is assumed that MRF stimulation interferes only with STM processes, rapid growth and consolidation of LTM are indicated. Hippocampal stimulation that is applied 196 seconds after an FS produces

little suppression at the 256-second retest. If it is assumed that hippocampal stimulation interferes only with LTM processes, then this result indicates decay of STM. Thus, it appears that 256 seconds after an FS the animal's memory is primarily LTM, with some slight involvement of STM. The data also support a differential neural basis for independent processing of STM and LTM.

RAYMOND P. KESNER\*

HUBERT S. CONNER

Department of Psychology, University of Utah, Salt Lake City 84112

#### References and Notes

1. J. L. McGaugh and D. R. Dawson, *Behav. Sci.* 16, 45 (1971); L. Weiskrantz, in *Amnesia*, C. W. M. Whitty and O. R. Zangwill, Eds. (Butterworth, London, 1966), p. 1; A. Cherkin, *Proc. Nat. Acad. Sci. U.S.A.* 63, 1094 (1969).
2. B. Milner, in *Amnesia*, C. W. M. Whitty and O. R. Zangwill, Eds. (Butterworth, London, 1966), p. 109.
3. T. Shallice and E. K. Warrington, *Quart. J. Exp. Psychol.* 22, 261 (1970); E. K. Warrington and T. Shallice, *Brain* 92, 885 (1969).
4. E. J. Wyers, H. V. S. Peeke, J. S. Williston, M. J. Herz, *Exp. Neurol.* 22, 350 (1968); E. J. Wyers and S. A. Deadwyler, *Physiol. Behav.* 6, 97 (1971); M. W. Wilburn and R. P. Kesner, *Exp. Neurol.*, in press.
5. J. H. McDonough, Jr., and R. P. Kesner, *J. Comp. Physiol. Psychol.* 77, 171 (1971).
6. R. Swanson, J. L. McGaugh, C. Cotman, *Commun. Behav. Biol.* 4, 239 (1969); S. H. Barondes, *Int. Rev. Neurobiol.* 12, 177 (1970).
7. A. Geller, F. Robustelli, S. H. Barondes, H. D. Cohen, M. E. Jarvik, *Psychopharmacologia* 14, 371 (1969).
8. R. P. Kesner, J. H. McDonough, Jr., R. W. Doty, *Exp. Neurol.* 27, 527 (1970).
9. Supported by NIMH grant MH-16918-03. We thank J. Denbutter for capable histological work and M. Garceau and J. Carli for aid in collection of the data.

\* Present address: Center for Advanced Study in the Behavioral Sciences, Stanford, California 94305.

23 August 1971; revised 2 February 1972

## Selective Dissemination

Schneider (1) has presented an interesting and complete account of the implementation of a system for the selective dissemination of information (SDI) for cancer-related literature. Unfortunately, enthusiasm for his own approach—use of an enumerative classification—has led him into some rather sweeping claims regarding the superiority of this method over a whole host of others, which he lumps under the general heading “keyword-based” and treats in a somewhat cavalier and irresponsible fashion. Among the “keyword-based” systems are those using uncontrolled keywords (humanly assigned) and those based on subject headings and thesauri, as well as systems operating on free text (that is, searching the natural language of a machine-readable text with a computer).

In actual fact, a retrieval or dissemination system, if properly designed, can function effectively via any of these methods. Under a certain set of conditions one method will be preferable to another, but all can be made to work. Schneider's criticism of existing systems (“imprecise indexing,” “a high level of ‘noise,’” and “occasionally provide a useful item of information to users” are among statements used) is exaggerated and highly subjective. Moreover, he fails to cite a single study to justify his criticism. Indeed, he chooses to dismiss lightly the results of the ASLIB-Cranfield Project (2), the most complete study of indexing languages yet undertaken, presumably because these results do not fit his own view of the universe.

Instead, Schneider refers to a recent

study at the Science Information Exchange (SIE) (3), the results of which tend to indicate that by searching on humanly assigned terms from an enumerative classification it was possible to get better results than when searching the free text of project abstracts. Unfortunately, Schneider has fallen into the trap (as did the SIE investigators) of comparing a good example of one kind of system with a bad example of another. The reason that better results were obtained with the classification approach is that the searchers were given a searching aid (that is, the classification itself), whereas the searchers in the free-text system were left to their own devices to think of all possible ways in which a particular concept might be represented in natural language. Schneider has overlooked the fact that some type of thesaurus, used as a searching aid, is as necessary for effective retrieval in a free-text system as it is in a controlled-vocabulary system. Many people overlook this despite the fact that at the University of Pittsburgh (4) a thesaurus was being used in conjunction with natural-language searching of legal text a decade ago.

I believe that free-text searching is becoming increasingly attractive for retrieval and dissemination activities, if only for the following reasons: (i) more and more bibliographic material is being made available in machine-readable form; (ii) computers are increasing in power, storage devices in capacity, and search programs in sophistication; and (iii) on-line search systems allow the user to browse within large bodies of free text in an interactive mode, a very valuable capability.

In further support of the free-text approach I would mention the following:

1) Systems do exist and operate on the basis of natural-language searching. Some have been functioning effectively for several years.

2) In the Cranfield studies (2), those systems based on natural language outperformed the others.

3) After a group of controlled comparative studies (5), the Institution of Electrical Engineers elected to use the natural language route in preference to the alternative approaches.

4) A system based on free text has the advantage of complete specificity since it is unlikely that we can index a document more specifically than by the words contained in it. Conversely, any controlled vocabulary inevitably sacrifices complete specificity.

5) In a full text system, indexing is completely "exhaustive" (that is, there is no information loss), which is untrue of any system based on humanly created surrogates.

6) If one has a thesaurus (to group related terms together, especially synonyms and near-synonyms) as a searching aid, it is possible to create (theoretically, at least) the same document classes at the time of searching that one would create, using a controlled vocabulary, at the time of indexing. This can be regarded as a "postcontrolled" as opposed to the conventional "precontrolled" vocabulary. Given complete specificity of text, and a searching thesaurus, the natural-language system has tremendous flexibility in allowing searches designed for maximum recall or for maximum precision.

7) There is an evident move toward natural-language systems in the United States at the present time. Even former bastions of the controlled indexing vocabulary appear to be weakening. As an example, Klingbiel (6) has recently stated: "Highly structured controlled vocabularies are obsolete for indexing and retrieval" and also: "The natural language of scientific prose is fully adequate for indexing and retrieval."

A humanly assigned enumerative classification appears to work well for Schneider's particular application, which is a relatively small-scale operation. But certain other methods would be likely to work equally well. We must not let our enthusiasm for one approach blind us to the possibility of others. In condemning free-text systems, for example, Schneider chooses to ignore a large body of operating experience and experimental evidence. He is also swimming against a very strong tide.

F. WILFRID LANCASTER

*Graduate School of Library Science,  
University of Illinois,  
Urbana 61801*

#### References

1. J. H. Schneider *Science* 173, 300 (1971).
2. C. W. Cleverdon, J. Mills, M. Keen, *Factors Determining the Performance of Indexing Systems* (ASLIB-Cranfield Project, Bedford, England, 1966).
3. D. F. Hersey, W. R. Foster, E. W. Stalder, W. T. Carlson, *Proc. Amer. Soc. Inform. Sci.* 7, 291 (1970).
4. J. F. Harty, *Proc. Bien. Amer. Ass. Law Libr. Inst. Law Libr.* 5th (1961), p. 56.
5. T. M. Aitchison and J. M. Tracy, *Comparative Evaluation of Index Languages*, part I, *Design* (Institution of Electrical Engineers, London, 1969); T. M. Aitchison, A. M. Hall, K. H. Lavelle, J. M. Tracy, *Comparative Evaluation of Index Languages*, part 2, *Results* (Institution Center, Alexandria, Virginia, 1970).
6. P. H. Klingbiel, *The Future of Indexing and Retrieval Vocabularies* (Defense Documentation Center, Alexandria, Virginia 1970).

15 September 1971

Lancaster has missed the basic point of my article which was to demonstrate that a hierarchical decimal classification could be used in addition to, or as a substitute for, current indexing methods almost all of which are based on the well-known methods mentioned in his first paragraph. When one presents a rather novel and complex point of view in a limited space, it is not possible to describe all the advantages (some were mentioned) of alternative methods, which obviously "can be made to work" to some extent. It seems rather ad hominem to use the words "cavalier" and "irresponsible" under such circumstances, or to pick three phrases out of context from a lengthy article and ridicule them as being exaggerated.

As a matter of fact, scientists I have talked to often make subjective statements regarding the high noise level and the low level of real "usefulness" (as distinguished from "relevance" or "interest") of data obtained from automated systems. An excellent, extensive study by Lancaster himself (1) might be cited as partial justification for statements I made about existing systems. He found an "overall precision ratio" (relevant articles divided by the total articles retrieved) of 50.4 percent for 299 searches made by the MEDLARS information retrieval system at the National Library of Medicine. This result means that 49.6 percent of the articles retrieved were judged to be of "no value" by the users. I believe this to be a reasonably "high level of noise." Out of 5278 abstracts evaluated by 103 scientists participating in my system for the selective dissemination of information (SDI), only 10 percent were judged to be "of no significant interest" and only 18 percent were "of no significant use" during the 1-year trial.

Lancaster also found that in 278 MEDLARS searches having "precision" failures (out of 302 searches studied in detail), 167 had failures due to indexing. In the same 302 searches, there were 238 searches where "recall" failure occurred (failure to retrieve an article that would have been judged relevant by the user), and 203 of these had failures attributed to indexing problems. This would seem to indicate a significant level of "imprecise indexing." Improvements in MEDLARS indexing methods since Lancaster's study in 1968 have probably reduced the noise level and the number of indexing failures to a significant extent.

It is natural for Lancaster to defend the results of the Cranfield project and to consider them definitive since he participated in the project (2). Rather than dismiss it lightly, I gave references to several papers that questioned the methods used and the significance of the results. In addition, I pointed out that the type of system I used, involving single-hit matching of categories from an extensive hierarchical classification with no post-coordination, had not been included in any major test of indexing methods. I stressed the need for additional tests of indexing languages which might involve less artificial test situations than the Cranfield project and would provide results more directly applicable to operating situations.

It may be misleading for Lancaster to imply indirectly in point 3 that selection of the natural language route by the Institution of Electrical Engineers (IEE) supports the superiority of natural language found in the Cranfield studies mentioned in point 2. In a letter to me Aitchison from IEE has stated: "Our experiments comparing the performance of the different index languages which are available on our files showed reasonably clearly that the thesaurus-based controlled-language was superior. However, our recall failure analysis showed the reasons for this, and we hope to reduce the performance difference between the controlled-language and the uncontrolled-language indexing (which for other reasons is much more attractive to us) by using a thesaurus for searching and profile compilation in conjunction with the uncontrolled-language." In addition, the first conclusion of Aitchison's report states: "In this evaluation, the controlled language is superior in performance to any of the other languages" (3).

As Lancaster points out, a classification does serve as an excellent aid to the searcher. It places information into logical hierarchies that parallel the thought processes involved in information retrieval. The ability to use a classification to pinpoint very small areas or "microprofiles" of information (in contrast to broad categories or "macroprofiles") and to organize data at the time of input greatly facilitates the retrieval process and improves its accuracy.

In contrast, Lancaster aptly describes a major problem with most free-text searching, namely, that the users to a

large extent are "left to their own devices to think of all possible ways in which a particular concept might be represented in natural language." An increasing number of free-text search systems are developing or have developed some type of aids to guide the user to possible search terms, although the lack of published data about such developments may explain why Lancaster selected a document published 11 years ago to illustrate this point. Many advocates of free-text searching originally felt that the use of natural language would eliminate the need to develop any type of vocabulary (see Klingbiel's statement which Lancaster quotes in point 7) and that this was a major advantage of free-text systems.

Although the free-text searchers in the Science Information Exchange (SIE) study did not use any word lists or thesauri, they were highly qualified literature analysts with 3 to 15 years of experience as information scientists. All had either an M.S. or a Ph.D. in a biomedical subject area (except for one with a B.S. in biology and an M.S. in library science). It is unlikely that they missed any major terms in formulating their free-text searches (which dealt with biomedical topics) or that their performance could have been significantly improved by the use of some type of vocabulary. Lancaster's claim that the free-text system evaluated by SIE was "bad" because no user aids were furnished to the searchers is obviously an overstatement.

Even when such aids are available, they present problems. Cuadra, a strong but objective advocate of on-line systems, has pointed out that the usual "teletype variety of terminal often creates the effect of a 'peephole' through which the contents of the files (and the authority lists or thesaurus) must be viewed serially by the new user, a little piece at a time. Even though the answer to a query may return in 6 to 12 seconds, it usually gives no indication of the 'conceptual distances or directions' it has taken in skipping from one place to another" (4). Under these circumstances the efficiency and value of on-line browsing which Lancaster stresses in point (iii) is dubious.

Lancaster states that "some type of thesaurus, used as a searching aid, is as necessary . . . in a free-text system as it is in a controlled-vocabulary system." He is probably not referring to a simple alphabetical list that may reach 100,000 or more words (including all

possible synonyms, near-synonyms, and linguistic or stylistic variants such as spelling, spacing, hyphenation, plural forms, and so forth) that would be needed for searching a data base of moderate size. (This type of printed or on-line dictionary is often the only "aid" supplied to the user.) Instead, he is undoubtedly referring to a much more complex, structured thesaurus showing all related terms, including synonyms and near-synonyms plus broader terms and narrower terms, for each nontrivial word in the data base. However, this point is uncertain since in point 7 Lancaster seems to be opposed to structured, controlled vocabularies and to agree that they are obsolete after advocating an almost identical type of thesaurus in point 6.

Development and updating of such a thesaurus is a very complex task. Fenichel (5) points out that in the course of producing a machine-searchable data base at the Institute for Scientific Information (ISI) there are 4,000 new words out of 70,000 words in titles each week, that 2,300 out of every weekly batch of 7,000 titles have at least one new word, and that 150,000 new words would be added to the dictionary in 1971. (Some of these are undoubtedly minor or trivial variants of words already in the dictionary, but the computer treats each one as an entirely new word.) The problems of updating massive thesauri or alphabetic listings to keep them abreast of "new words" in titles are compounded manyfold when abstracts or other text are included, as in the "full-text" systems Lancaster refers to. The ISI handles this problem by an annual purge of more than 100,000 words that appear only once in the dictionary. This is clearly an unsatisfactory solution for a thesaurus used in a free-text search system.

In view of these problems, plus the fact that most machine-searchable data bases do not contain abstracts, theoretical generalities about using all the words in a document as index terms, as expressed by Lancaster in points 4 and 5, have little relation to existing, operational free-text search systems, many of which are based, in part, on humanly assigned index codes or terms. In addition, in my article I pointed out how the specificity of ideas is lost when thoughts are separated into individual index terms. I stressed the difficulty of retrieving complex concepts by combining several isolated keywords.

To balance Lancaster's enthusiasm for his simplistic view of on-line free-text searching, it seems appropriate to mention several additional points:

1) Although the ability to handle a large number of entries is improving, most operational free-text search systems are limited to a relatively small number of complete records (often well under 50,000). For many data bases this restriction means that only documents published over a very brief recent time period (2 or 3 years) can be searched.

2) As the size of the files increases, the complexity and cost of updating them and maintaining them on-line for long periods increases accordingly.

3) As the files grow larger and the number of users more numerous, computer response time increases. One of the largest on-line systems sometimes requires more than 40 seconds for the computer to respond to each command as the user tries to formulate a search. Other irritating and frustrating problems with a small-scale on-line system have been described in detail by Lancaster (6).

4) Hersey *et al.* in a more complete description of the SIE test (7) state that "the free text word retrieval approach is particularly susceptible to low recall of projects known to be pertinent." The user usually obtains some useful references, but he has no idea of the number of additional relevant documents that he missed. Lancaster found that 11 out of 45 on-line searches of a small epilepsy data base with free-text searching of 8000 abstracts, titles, and index terms retrieved less than 20 percent of the relevant documents, and 23 of the 45 searches retrieved less than 53 percent of the relevant documents (6). This partial recall is due, in part, to Cuadra's "peephole" phenomenon mentioned above.

5) There is some movement toward classifications in the United States. In order to deal with large data bases and to supply group SDI services, classifications consisting of categories or "macroprofiles" covering broad subject areas are being used with increasing frequency. The INSPEC Service at the Institution of Electrical Engineers in London uses such categories. The American Mathematical Society and the American Institute of Physics use

more detailed classifications (8, 9).

My article was directed mainly at demonstrating the high level of performance possible when enumerative classifications are used for selective dissemination of information and other automated information systems. I do not feel that they should be used exclusively for every information system. Instead, the best systems are likely to be those that use a combination of both detailed enumerative classifications and keyword or free-text searching for information retrieval. The American Institute of Physics has successfully demonstrated the feasibility of such hybrid systems (9).

JOHN H. SCHNEIDER

*Scientific and Technical Information  
Office, National Cancer  
Institute, National Institutes of  
Health, Bethesda, Maryland 20014*

## Interrelations of Humans, Dogs, and Rodents

In discussing the problems of evaluating tests of the toxicity and teratogenicity of 2,4,5-T, Sterling (1) stated that interpretation of human reactions from animal studies is complicated by the use of rodents as experimental animals because "rodents are much further removed phylogenetically from the human animal than are dogs or monkeys."

Of course, monkeys are unquestionably closer phylogenetically to humans than either dogs or rodents are. No Cenozoic common ancestors are known for the three orders, Primates, Carnivora, and Rodentia, and the relationships of the three (other than all being placental mammals) are not universally agreed on. A conservative approach shows all three lines converging somewhere in the Upper Cretaceous (2). The earliest known possible primates are from the late Cretaceous of Montana (3), and Primates were certainly well established by mid-Paleocene; a possible carnivore is reported from the early Paleocene of New Mexico (4); the earliest known rodents are from the latest Paleocene of Wyoming (5). From the fossil record, it appears fairly certain that rodents and primates are more closely related to each other than either is to carnivores. The latest study of the

## References and Notes

1. F. W. Lancaster, *Evaluation of the MEDLARS Demand Search Service* (National Library of Medicine, Bethesda, Maryland, 1968), pp. 41-43.
2. C. W. Cleverdon, F. W. Lancaster, J. Mills, *Spec. Libr.* **55**, 86 (1964); F. W. Lancaster and J. Mills, *Amer. Doc.* **15**, 4 (1964).
3. T. M. Aitchison, letter dated 8 December 1970; —, A. M. Hall, K. H. LaVelle, J. M. Tracy, *Comparative Evaluation of Index Languages*, part 2, *Results* (Institution of Electrical Engineers, London, 1970), p. 61.
4. C. A. Cuadra, *J. Amer. Soc. Inf. Sci.* **22**, 107 (1971).
5. C. Fenichel, *Proc. Amer. Inf. Sci.* **8**, 349 (1971).
6. F. W. Lancaster, *An Evaluation of EARS (Epilepsy Abstracts Retrieval System) and Factors Governing Its Effectiveness* (Graduate School of Library Science, Univ. of Illinois, Urbana, 1971), pp. 8-10 and 39-50.
7. D. F. Hersey, W. R. Foster, E. W. Stalder, W. T. Carlson, *J. Doc.* **27**, 167 (1971).
8. J. H. Schneider, *Science* **173**, 300 (1971) [reference (21) describes several systems that use macroprofiles in SDI-like announcement services]; *A Preliminary Annotated Bibliography of Papers Related to the Use of Classifications in Automated Information Systems* (National Cancer Institute, Bethesda, Maryland, 1971).
9. H. W. Koch, *Science* **174**, 918 (1971).

7 March 1972

earliest rodents has compared them particularly with certain Paleocene primates, and the conclusion was reached that an origin of rodents from primates at some time during the Paleocene was more probable than any other origin (5). If this derivation of the rodents is correct, their primate ancestor lived on the order of  $70 \times 10^6$  years ago, so that living rodents and primates are not very closely related; however, the latest common ancestor of primates and carnivores must have lived even earlier.

Therefore, tests on rodents should give every bit as valid indications of human reactions as would tests on dogs. This conclusion, of course, has no bearing on the validity of Sterling's other comments.

ALBERT E. WOOD\*

*Department of Biology, Amherst  
College, Amherst, Massachusetts 01002*

## References and Notes

1. T. D. Sterling, *Science* **174**, 1358 (1971).
  2. A. S. Romer, *Vertebrate Paleontology* (Univ. of Chicago Press, Chicago, 1966), fig. 316.
  3. L. Van Valen and R. E. Sloan, *Science* **150**, 743 (1965).
  4. G. MacIntyre, *Amer. Mus. Nat. Hist. Bull.* **131**, 115 (1966).
  5. A. E. Wood, *Trans. Amer. Phil. Soc. New Ser.* **52**, 1 (1962).
- \* Present address: 20 Hereford Avenue, Cape May Court House, New Jersey 08210.  
24 January 1972