

The Next Twenty Years in Information Retrieval: Some Goals and Predictions*

CALVIN N. MOOERS†

A HISTORICAL PERSPECTIVE

ALTHOUGH information retrieval has lately become quite a fad, I intend in this paper to stand back and take an unhurried look at what is going on, and try to predict where this field must go and what it must do in the future. "Information retrieval" is a name that I had the pleasure of coining only eight years ago, at another computer conference.¹ The name has come a long way since then.

In thinking about a definition of information retrieval and in considering the future of this field, we must take an evolving view. At the present time, information retrieval is concerned with more than the mere finding and providing of documents; we are already concerned with the discovery and provision of information quite apart from its documentary form. To encompass future developments, as we shall see, even this broad view of information retrieval will have to be modified and extended.

When we speak of information retrieval, we are really thinking about the use of machines in information retrieval. The purpose of using machines here, as in other valid applications, is to give the machines some of the tasks connected with recorded information that are most burdensome and unsuited to performance by human beings. At all times, it is important to remember that it is the human customer who uses the information-retrieval system who must be served, and not the machine. It makes a difference who is served, and this little matter is sometimes forgotten in computer projects.

To get a historical perspective of the introduction of machine methods to information retrieval, let us look back over a bit of history. I think that it can be said that the introduction of machine methods has followed the realization of a need, backed by pressure and means to do something about the need. Thus, although quite powerful mechanical methods could have been de-

veloped by the technology of the Hellenistic Era for the Library of Alexandria, other methods of retrieval, presumably based upon human memory, and the making of lists, were apparently considered quite satisfactory. The simple, though powerful, mechanical technique that could have been used at Alexandria is the method of perforated stencils invented by Taylor in 1915, which has sometimes more recently been called "peek-a-boo."² The British inventor Soper in 1920 patented a device which was an improvement upon Taylor's perforated stencils, and Soper described the use of his mechanism for information retrieval employing some truly advanced conceptions.³

Much more attention, however, has been given to the development of devices for scanning and selecting upon film strips. This work was apparently spurred by the perfection of motion picture films and cameras. Goldberg in 1931 patented one of the earliest film-scanning and photographic-copying devices.⁴ Independently, Davis and Draeger during 1935, in the early days of the American Documentation Institute, in connection with their pioneering work in microfilm documentation, investigated the feasibility of a microfilm scanner using a decimal coding.⁵ Apparently stimulated by reports of this work, V. Bush and his students at M.I.T. in 1938-1939 built perhaps the first prototype machine along these lines, a microfilm scanner with each frame of text delineated by a single decimal code for the subject, and a photoflash copying method. However, they were unable to interest any commercial or governmental organization in the device, and wartime distractions intervened soon thereafter, so the project was dropped. Two more "rapid selectors" based upon these same general principles have been built,^{6,7} but for various reasons neither of them has operated in a fashion that is consid-

* This work has been supported in part by the U. S. Air Force Office of Sci. Res. through Contract No. AF 49(638)-376. All opinions are those of the author.

† Zator Co., Cambridge, Mass.

¹ C. N. Mooers, "The Theory of Digital Handling of Non-Numerical Information and its Implications to Machine Economics," Zator Co., Cambridge, Mass., Tech. Bull. No. 48, 1950; paper presented at the March, 1950 meeting of the Association for Computing Machinery at Rutgers University, New Brunswick, N. J.

C. N. Mooers, "Information retrieval viewed as temporal signaling," *Proc. Internat. Congr. of Mathematicians*, Harvard University, Cambridge, Mass., vol. 1, p. 572; August 30-September 6, 1950.

² H. Taylor, "Selective Device," U. S. Patent No. 1,165,465; December 28, 1915 (filed September 14, 1915).

³ H. E. Soper, "Means for Compiling Tabular and Statistical Data," U. S. Patent No. 1,351,692; August 31, 1920 (filed July 23, 1918).

⁴ E. Goldberg, "Statistical Machine," U. S. Patent No. 1,838,389; December 29, 1931 (filed April 5, 1928).

⁵ R. H. Draeger, "A Proposed Photoelectric Selecting Mechanism for the Sorting of Bibliographic Abstract Entries from 35 mm Film," Documentation Inst. of Science Service, Washington, D. C. (now American Documentation Inst.), Document No. 62; July 27, 1935.

⁶ R. Shaw, "Rapid selector," *J. Documentation*, vol. 5, pp. 164-171; December, 1949.

⁷ Anonymous, "Current Research and Development in Scientific Documentation," National Science Foundation, Washington, D. C., Rep. No. 3, p. 27; 1958.

ered generally acceptable, and neither is currently in actual use. At the present time, much attention is focused upon the Eastman Minicard machine, a cross between a rapid selector and a Hollerith punched-card machine.⁸

Card-sorting devices, such as those based upon the Hollerith card (the card used by IBM), as well as those based on other cards such as Perkins' marginally punched card, were recognized at an early date to be far too slow to cope with problems such as those of *Chemical Abstracts*. Within the past few years, there have been a number of instances of the use of electronic computing machines to perform information retrieval. As computing machines are presently designed, they are not matched to the job of information retrieval—they can do it, though not efficiently—and the situation of using a computing machine for this job is like using a bulldozer to crack peanuts. Oftentimes, if the information collection is small enough to allow the problem to fit upon the computer, there are easier methods to perform retrieval. If the collection is large, it does not have to be very large to tie up all the computer's memory capacity. It is clear that special computer-like devices will be called for if we are to perform efficient large-scale information retrieval.

Although we have been trying to build high-speed selecting machines for information retrieval over the past twenty years, since the date of Bush's machine, at the present time I do not think that it can honestly be said that we have done too well. We do not really have a machine which is an altogether happy answer to the problems of search and selection on collections ranging in size upwards from fifty or one hundred thousand items. The problem becomes even more unmanageable at the million point, since this size of collection requires reasonably high-speed processing and decision on a scanned record of something like 10^9 bits.

However, the hardware will be built—and is being built. But what about the classification terminology, the subject headings, the descriptors, and the like? One after another, various machine projects have foundered on this problem, especially those projects that have copied library classification decimal systems or made use in a detailed way of their indexing techniques. We should appreciate that new mechanisms deserve new methods, and that there is a consensus of opinion (although it is not unanimous) that the method of putting together independent idea-expressing terms and selecting upon their correlative occurrence constitutes the desired point of departure from the historic methods of the library.

A highly developed form of this point of view is the method of "descriptors," which was introduced and de-

veloped in theory in 1948–1950 in a number of papers in conjunction with a mechanical card selector.⁹ The descriptor method, which makes a great point of employing precisely defined terms composing a limited vocabulary, is a refinement of a number of earlier practices. The method was implicit in the work of Soper, it was toyed with and dropped by the librarian Bliss, and it was used in one fashion or another by a number of scientists and chemists with Perkins cards in the 1940's, e.g., by Bailey, Casey, and Cox.¹⁰ People seem to confuse descriptors with Uniterms. The latter might be described as a crude form of a descriptor system, originally making use of words lifted from titles and texts. The Uniterm approach, since it was introduced in 1951, seems however to be migrating both in concept and usage towards the descriptor methods, as is clear from many reports coming from projects where they claim to use Uniterms.

The problem of classification terminology or language symbols for machine retrieval is well toward a solution, even for complex and structured kinds of information. An example is the work on the coding of chemical structures for machine retrieval.^{11,12} However, it should be noted that considerable work on retrieval of structured information, especially for chemical compounds, has sometimes resulted in symbolism that is not completely suitable for machine use, as for example some of the methods considered by the International Union of Pure and Applied Chemistry.

THE PRESENT STATE OF AFFAIRS

Although we may soon have suitable machines for large-scale information retrieval and although the situation with respect to the language symbols of retrieval is in a reasonably satisfactory state (that is, ahead of the machines) we are not yet finished with our problems.

Presuming that we have a machine completely capable of dealing with a collection of one million—or even a hundred million—items, who will read these items or documents and assign the descriptors? Experience has shown that this is a difficult and time-consuming job. For example, in my experience in reading and coding patents, it takes me about fifteen minutes of reading, on the average, merely to figure out what the inventor is driving at. The Patent Office has some three million of such patents.

This is exactly the kind of burdensome job that should be turned over to the machine. In fact this problem is under active consideration and study in a number

⁸ A. W. Tyler, W. L. Myers, and J. W. Knipers, "The application of Kodak Minicard system to problems of documentation," *Amer. Documentation*, vol. 6, pp. 18–30; January, 1955.

⁹ C. N. Mooers, "Zatocoding and developments in information retrieval," *ASLIB Proc.*, vol. 8, pp. 3–22; February 1956. This paper summarizes these developments.

¹⁰ C. F. Bailey, R. S. Casey, and G. J. Cox, "Punched-cards techniques and applications," *J. Chem. Ed.*, vol. 23, pp. 495–499; 1946.

¹¹ C. N. Mooers, "Ciphering Chemical Formulas—The Zatopleg System," Zator Co., Cambridge, Mass., Tech. Bull. No. 59; 1951.

¹² L. C. Ray and R. A. Kirsch, "Finding chemical records by digital computers," *Science*, vol. 126, pp. 814–819; October 25, 1957.

of places. It is not an easy task to give to a machine. It contains a great many aspects that would seem to require the exercise of real "intelligence." Fortunately, we already have one remarkable accomplishment which shows that this seemingly intellectual job is not completely incompatible with mechanization. I speak of the work by Luhn on his "auto-abstractor."¹³ By the method of Luhn, the computer takes in the text of an article, statistically picks out the unusual words, and then chooses sentences containing these words to make up the auto-abstract. If this process were terminated at the point of picking out the words, we would have Uni-terms. If the words picked out in this fashion could be replaced by standardized words having approximately the same meaning, that is, if the synonyms could be eliminated, then we would have descriptors. It should be noted that this kind of treatment of synonyms, which has been going on in retrieval for some years, has lately been given the fashionable name of "the thesaurus method." In the interests of precision in terminology, I should like to point out that there are significant differences in Roget's concept of a thesaurus and the set of equivalence classes of terminology that are required for retrieval. Indeed, this is precisely why I introduced the new terminology "descriptor" some years ago, that is, to give a verbal handle for a group of new conceptual methods with language symbols.

Such a take-off on Luhn's method would not be the final answer, because as Luhn has set it up, the machine is operating in an essentially brainless fashion. To do better than merely picking up words on a statistical basis, we would have to build into the method the capability of handling the equivalence classes of words and phrases. This gets us into language translation. After the statistical approach has segregated words of high import from the text, we need to translate these words into the standardized descriptor terminology for further retrieval. However, even building up the equivalence classes of the terminology is a burdensome job, and this too should be turned over to the machine. Not only should the machine build up these equivalence classes, but it should be made to refine its performance with respect to using these terms and getting the descriptors, and it should even be made to learn how to improve its performance.

Fano and others have suggested the use of statistics on the way people come in and use the library collection in order to provide feedback to help a machine improve its performance.¹⁴ While the suggestion is in the right direction, I think that this kind of feedback would be a rather erratic source of information on equivalence

classes, because people might well borrow books by Jack London and Albert Einstein at the same time. While this difficulty can be overcome, there is a more severe problem. Any computation of the number of people entering a library and the books borrowed per day, compared with the size of the collection, shows, I think, that the rate of accumulation of such feedback information would be all too slow for the library machine to catch up to and get ahead of an expanding technology.

In this respect, it is my speculation that a more powerful source of educational material for a machine is already available, and it should be tapped. Despite the admitted limitations of such material, the subject entries, the decimal classification entries, and the other content typed on catalog cards contains a great deal of ready information that can be used in teaching a machine how to assign descriptors to documents. Other collections, besides those in the libraries, also often provide a ready source of classificatory information that should be tapped. For instance, in the Patent Office, in each case record of each application for patent, there is a great amount of specific reference to other related patents, and this information, along with the assigned class numbers, is readily available for machine digestion without further high-level human intellectual effort.

In order to do these things, we shall need a machine with some rudimentary kind of "intelligence;" or more accurately, we shall need an "inductive inference machine" in the sense used by Solomonoff.¹⁵ An inductive inference machine is one that can be shown a series of correctly worked out examples of problems, that can learn from these problems, and that can then go ahead on its own (probably with some supervision and corrective intervention) to solve other problems in the same class. While an inductive inference machine can be quite capable at a given class of jobs, it need not have "brains" or "intelligence" in the general sense.

As I mentioned before, putting the descriptors on the documents—that is, delineating the information in the text by symbols for retrieval—is a form of crude language translation. It is crude because the machine does not need to worry about grammar in the target language, since the grammar of descriptors is nonexistent, or at most, is rudimentary. As I see it, machine translations of this kind for the purposes of information retrieval will be an area of early pay-off for work in inductive inference machines.

If inductive inference machines can be built at all, then it certainly should be possible for us to feed them with subject headings and classification numbers on the one hand, and with the titles of book chapters and section headings on the other hand, in order to teach the machines how to do at least some rudimentary kind of job of library subject cataloging. With librarians at

¹³ H. P. Luhn, "The automatic creation of literature abstracts," *IBM J. Res. Dev.*, vol. 2, p. 159; April, 1958.

¹⁴ R. M. Fano, "Information theory and the retrieval of recorded information," in "Documentation in Action," J. H. Shera, A. Kent, and J. W. Perry, eds., Reinhold Publishing Corp., New York, N. Y., ch. 14-C, p. 241; 1956.

¹⁵ R. J. Solomonoff, "An inductive inference machine," 1957 IRE NATIONAL CONVENTION RECORD, pt. 2, pp. 56-62.

hand to provide suitable intervention or feedback, the machine's performance should improve, and by further work, even the categories or descriptors which are used can be improved by the machine. Thus I feel that we can expect in the next few years to see some interesting results along this line.

GOALS FOR THE NEAR FUTURE

There are a number of other applications of machines for purposes of information retrieval of a kind that have not yet been seriously undertaken, and others that have not yet been considered. In my discussion I shall bypass treating such useful and imminent tasks as the use of machines to store, transfer, and emit texts, so that at the time that you need to refer to a paper, even in an obscure journal, you can have a copy in hand within, say, twenty-four hours. Neither shall I consider the application of machines to the rationalization and automation of library ordering, receiving, listing, warehousing, and providing of documents. Neither shall I consider the application of machines to the integration of national and international library systems so that at any first-rate library, you will have at your command the catalogs of the major collections of the world. These are all coming—but it should be noted with respect to them that the problems of human cooperation ranging from person-to-person to nation-to-nation cooperation are more serious than some of the machine and technical problems involved.

The first of the rather unusual applications of machines to information retrieval that I want to talk about can be introduced as follows. When a customer comes to an information retrieval system, he comes in a state of ignorance. After all, he needs information. Thus, his problem of knowing how to specify pieces of information that are unknown to him is a severe one. For one thing, the vocabulary of the retrieval system, and the usages of the terms in the system, may be slightly different from the language that he is used to. For another thing, upon seeing some of the information emitted according to his own retrieval prescription, he may decide that an entirely different prescription should be used. In short, the customer definitely needs help in using a machine information retrieval system, and this help should be provided by the machine.

An indication of what kind of system needs to be provided, and how it can be done, is given by certain of the simple sorted-card retrieval systems. Some of the sorted-card systems do very well in this respect, others do not. It has been common practice in Zatocoding systems, which use a simple schedule of a few hundred descriptors, to employ a descriptor dictionary system having many cross-references from words in the ordinary technical usage to the appropriate descriptors.⁹ Thus the customer can find his way into the system starting out with his own terminology. After the customer is referred to a descriptor, he finds there a carefully drafted scope note explaining the range of meaning attached to the

particular descriptor. In another tabulation in the descriptor dictionary, the descriptors themselves are grouped or categorized into fifteen or twenty groups, and each group is headed by a question pertinent to the descriptors under it. Thus, under a question "Are geometrical shapes involved?" would be found descriptors such as "round," "square," "spherical," etc.

These simple card systems provide another source of assistance to the customer because they are able to emit cards within a minute or less from the time the retrieval search is begun. Thus if the search is headed into the wrong direction, the customer, upon looking at the cards or documents, will immediately detect this fact, and can reframe his request to the machine before any further searching is done. It is deplorable, but true, that many contemporary proposals for machine systems may be so slow in providing feedback that the feedback time is measured in hours or days, with the consequent waste of machine sorting time and accumulation of human frustration.

The problems of customer assistance are going to be severe with the large machine retrieval systems of the future, and these problems must be faced. The descriptor vocabularies are going to be large. Another possibility is that some of the machines will operate internally on vocabularies or machine code systems that are quite unacceptable for external communication to the human operators. There has already been some successful experimentation with symbol systems of this kind in coding chemicals.^{11,12} Such symbol systems work beautifully inside the machines, but people should not be forced to use them. For these reasons, in order to translate the customer's requests into forms suitable for the machine, machine assistance is going to be desirable. Holt and Turanski see this problem of processing the customer's request at the input to the machine as being very similar to the presently developing customer use of automatic programming for mathematical problems. The more advanced systems of automatic programming provide for a succession of stages of translation, with the symbolism at each stage moving further and further from the human word input to the abstract symbols and the detailed machine orders required for internal operation of the machine. In mathematical programming, the machine programs itself, and then carries out the program. In retrieval programming, the machine will form the proper machine prescription, and carry out the search. To my mind, there is an important difference. In retrieval, the machine should check back with the customer as it builds up the prescription in order to make sure that the search will be headed in the right direction; then it should search a sample of the collection and check again to make sure that the output being found is appropriate to the customer's needs. If we are to have larger and more complex machine retrieval systems, we must come to expect a great deal of back-and-forth man-machine communication during the formulation of a search, and as it is going on.

Quite another approach to handling the customer's input problem is advanced by Luhn, who suggests that the customer write a short essay detailing what he thinks is descriptive of the information he wants.¹⁶ The essay text would then presumably be handled in the fashion of the auto-abstract method (though Luhn is a little sketchy here on the details of his proposal), and the words selected from the short essay would be compared with words similarly selected from the document texts. When there is a sufficient degree of similarity, selection occurs. Although the Luhn proposal does put the load of translation of the customer's request upon the machine, it does not provide for customer guidance into the resources of the machine's selective language possibilities, or into the resources of the collection. Help in both of these directions would surely be of great assistance to a customer in extracting the maximum value from information in storage.

Another possibility is to use an inductive inference machine, because it is open to learning a great variety of tasks. It would be able to provide a generalized approach to the problem of customer assistance. But, however customer assistance is provided, I think it is safe to predict that we must build information retrieval systems with the planned capability to communicate back and forth with the customer so that he can better guide the machine in retrieving what will be useful to him.

RETRIEVAL VIEWED AS A PROCESS OF EDUCATION

If the machine aids the customer by guiding him in the use of the retrieval system, the machine is necessarily educating the customer. Let us take this viewpoint, and look upon a machine retrieval system as an educational tool. This viewpoint provides a number of new tangents to consider. We have seen how the customer can use some coaching by the machine in order to tap efficiently the information resources during the search process. But, as anyone knows who has had a large batch of documents sent his way, maybe the customer can also use some machine help in reading the mass of documents emitted from a retrieval system!

It is my prediction that some of the machine information retrieval systems of the future will go considerably beyond the tasks of mere retrieval or citing or providing document texts. I believe that some of them will also help the customer assimilate or read the output provided by the machine. This prediction is not at all fanciful, even though it is yet quite a way into the future. How far into the future it is we can only guess, or estimate by recalling that the Minicard follows a full twenty years after the first suggestions for a film selector, or that the widespread use of descriptors came about forty years after Taylor actually used something very much like them in information selection.

Machines can be very effective in teaching human beings. This is shown by the work of Skinner at Harvard where, in recent experiments, written modern languages and college mathematics have been set up on machine lessons.¹⁷ Essential to the process is rapid feedback, or communication between the machine and the human learner, so that the human knows immediately that he is on the right track, and so the machine can apply corrective action as soon as errors appear. Skinner's machines at present employ written materials prepared in advance by human beings, the machine performing on the basis of a fixed internal sequence of morsels of information of graded difficulty. However, machines need not be restricted to doing their teaching according to a preset sequence of lesson elements of this kind. In the same way that we are currently looking for techniques to allow machines to assign descriptors from texts, so can we contemplate the development of teaching procedures and machines whereby the machines by themselves will be able to pick out a graded sequence of information morsels from the documentary record retrieved and will then present them to the human customer.

Taking this view of a machine information center acting both as a retrieval device operating upon a store of information and as a teaching device for the human customer, we can see that the process of input request formulation and the process of giving out information will merge into a sustained communication back and forth between the customer and the machine. Of course, once the customer is on the track of documents containing information particularly pertinent to his interests, he will very likely desire to see the original text. This can be done, and a customer will have a choice of how much or how little of any particular actual document he wishes to read directly.

The range of future possibilities is even greater when these ideas are combined with the possibilities inherent in mechanical language translation devices. Of course, we should expect that future information centers will be able to provide translation from one ethnic language to another of the texts that the retrieval system provides. Let us look further. As is well known, one of the problems in machine language translation is to provide sentences in the target language in the required form—that is, to provide a smoothly running, colloquial translation. For example, in going from German to English, we must rescue the verbs from the end of the German sentence and put them up where they belong in the middle of the English sentence. Any machine capable of doing a high-grade language translation must be able to arrange and rearrange idea units and word units to make acceptable text out of them. This being the case, it is reasonable to predict that the information morsels that a teaching machine would put out could be given as in-

¹⁶ H. P. Luhn, "A statistical approach to mechanized encoding and searching of literary information," *IBM J. Res. Dev.*, vol. 1, p. 309; October, 1957.

¹⁷ B. F. Skinner, "Teaching machines," *Science*, vol. 128, pp. 969-977; October 24, 1958.

put to a machine technique patterned on the output half of a language translator. The resulting textual output would be in the nature of a written article having at least some degree of acceptable style.

This means that you could go to an information center, describe a certain kind of information needed, have the machine assist you in making your request more definite, and then order it: "Provide me with an 800-word article, not requiring more than an undergraduate chemistry background, on the deterioration of polymers by sunlight." After a short, but decent, interval, the machine would come forth with such an essay.

There is an important corollary to this notion of the machine central being able to provide essay articles upon request. We are all aware of the considerable duplication of information found in the technical literature. The same, or very similar, piece of information is repeated in one article after another. Many articles are summaries of other articles, or are derivative upon other articles, and provide little or nothing that is new. If the output from an information system does not need to be in the form of a graphic image of the original text, or a type-out of the text itself, then it is possible to consider the storage of new information only in the machine. A machine could store facts alone, and only new facts; it would not store text. By eliminating the dependence upon the original text, and avoiding the duplication of the same information written over and over, it might be possible to secure considerable increase in the machine's storage capabilities.

Yet there are problems of a kind that will occur to any thoughtful individual. I do not think we want to throw away the original record which we have already—the printed books and articles in our libraries. Neither am I sure that we want to give up entirely our system of printed publication. But, putting these problems aside, let us do some more speculating. It might be possible for the scientist in his laboratory to feed his raw (or nearly raw) results directly into a machine for computation, checking for acceptability, correlation with earlier facts, and ultimate storage. Thus, instead of a scientific archive existing almost solely on paper, as we now have, it is possible that a part of the archive in the future will be in machine form. The only way that such a machine archive would be tapped would be by having the machine write a summary or article upon specific request.

When we are thinking about information machines of this kind, I wish to stress that we should not think in terms of some big single machine central. This is important. It would be foolish and expensive to build up

a single central "bottleneck." If such machine central information systems as I describe will be at all possible, they will be important enough to be set up at a large number of installations, quite in the same way as we now are making use of a large number of electronic computer installations. There will be both large and small information machines. Some of these machines will be in intercommunication with each other, while others will operate in isolation. At various times, the machine memory from one or several of the machines can be played out onto tape, and the tape record, containing a vast amount of information, can be incorporated into the memory systems of many other information centrals.

If machines can store and correlate laboratory facts, and can communicate with laboratory workers, we shall have to expect that the machines will find gaps in the information as a part of the correlation, and they will point out to the laboratory workers the need for further experimentation in certain areas. How far we can expect this kind of active feedback to extend is hard to guess. The present work with pattern recognition will ultimately lead to a kind of a machine eye, and we already have machine hands for the handling of radioactive materials. An information central machine system, aided by such receptors and effectors, would become, in effect, a laboratory scientist.

At this point I would prefer to terminate my speculations on the excuse that we are now perhaps more than twenty years into the future, the limit that I set for myself in this paper.

In summary, I think that it can be said that mechanical information retrieval has started rather slowly; it has taken from about 1915 or 1920 until now to become as popular as it is. At the moment, except for certain highly integrated small retrieval systems, we are yet only dabbling in the subject. We do not now honestly have any appropriate large-scale machine for collections involving millions of items. We are only beginning to get a widespread recognition of the capabilities of suitable retrieval language systems, and there still remains the problem of getting machines with internal digital operations that are as suitable for retrieval and information work as the operations of addition and multiplication are suitable for mathematical work.

In any event, it is useful for us to know what some of our future targets are likely to be. With such knowledge, we will be in a better position to steer our activities in the present. This is the excuse for the predictions—which I take very seriously—that are contained in this paper.
