

Comments on Martin Lam's article

Lighthill 17 years on: end of a shotgun divorce by Donald Michie, The Turing Institute, Glasgow, UK.

Seventeen years ago, criticisms of AI expressed privately by Sir James Lighthill (subsequently published, 1973) led to a shotgun divorce between the cognitive-science and knowledge-engineering wings of this fledgling field. Using the term “bridge activity” for the living head and body which he found uniting the two wings, he pronounced this part as unnecessary. The consequences which followed are history (Fleck, 1982). They are also part of the recollections of those who stayed in Britain through the ensuing “AI winter”. I shall discuss Sir James' criticisms and the lessons to be learned.

One lesson stands out immediately. I quote from the opening summary of Martin Lam's contribution to this issue:

It is a common Establishment practice in the UK to ask an expert, recognized for his/her sound views, to give a view on controversial matters of public interest.

In the specifically UK context the above should be interpreted to mean:

It is a common Establishment practice in the UK to ask an expert *in some other field*, . . . to give a view . . .

and here, surely, is the rub. In Western Europe or North America one could safely assume an opposite interpretation, namely that it is a common practice to ask an expert in a *relevant* field.

The discontinuance in 1966 of US support for machine translation illustrates the difference in practice. The US National Academy of Science's adverse recommendation was based on the work of a committee chaired by John R. Pearce. As Research Director of Bell Laboratories he was recognized for his far-reaching knowledge of the technical area in question, as also were others of his committee. So whether the final recommendation was right or wrong, the linguistics and computing communities could be confident that it had not been taken superficially or lightly. This also has a bearing on the issue of validity.

Validity of provenance. Before coming to the content of Lighthill's report, there is a question concerned with administrative framework. Was his report a proper basis for formulating Government policy?

Although I am far from suggesting that procedural rules should never be bypassed, a widely overlooked fact suggests that it was not a proper basis. As a matter of record, Lighthill's report was not (as sometimes stated) “commissioned by the SRC (Science Research Council)” —for the good reason that SRC's Council had already commissioned, received, endorsed and published a policy review from their own Computing Science Committee covering this very topic (SRC Engineering Board, June 1972, *see* References). Although Martin Lam's article in this issue refers to “the Lighthill Report of 1972”, the latter was a subsequent and supplementary document—commissioned privately and circulated informally among Government departments and agencies, and not published until April 1973 (*see* References). The then SRC Chairman Sir Brian (now Lord) Flowers makes plain in his Preface that after reading the SRC's official Policy Review in draft he commissioned Lighthill's on a personal basis. He also mentions that Lighthill's contribution was subsequently discussed within SRC in 1972. However, it was not referred to SRC's Council, so that formally SRC remained committed to its own Engineering Board's *Policy and Programme Review of Research and Training in Computing Science*, which had earmarked for increased support the

following topics of Long Range Research: (1) The Physical and Logical Structure of Computers, (2) The Theory of Computation and (3) Machine Intelligence.

This vehicle for building up AI on a national scale, launched by a panel of distinguished and knowledgeable advisers, was heading for implementation. There followed the sudden approach, along the same single-track line but from the opposite direction, of Lighthill's solitary locomotive flying the Chairman's colours. Somewhat hurriedly the official vehicle was shunted out of the way, with consequences summarized below.

Political effect. The ensuing period, referred to above as Britain's AI winter, may now in political and financial terms be nearing an end. But official attempts at revival, such as the Alvey initiative of the 1980s, remained under severe handicap from Lighthill's conceptual legacy. The concept lingered that warmth could only legitimately be bestowed on the "machine" wing, or on the "brain" wing as the case might be, leaving the middlepiece frozen. The latter corresponds to what we ordinarily call AI—the living bridge between cognitive neuropsychology on the one hand and computing technology on the other. The Alvey officials, and their academic advisers from outside AI, were so sensitive to this conceptual legacy that the new initiative was launched in scrupulous avoidance of the term "Artificial Intelligence". A neologism was coined, "Intelligent Knowledge-Based Systems" (IKBS), as a device, it may be, whereby adventurous computer science and software research could be furthered while steering clear of work which could *only* be justified in terms of the officially discredited bridge. Some areas of high importance and quality were fostered within this restriction. Planning and search studies, for example, can be justified from cognate activities in operations research, and logic programming can derive ample overt motivation without needing to appeal to AI. Other areas, however, such as machine learning, could not easily be accommodated without explicitly departing from Lighthill's legacy. In the event partial rescues were mounted by individual and commercial initiatives independent of Government funding, as will appear later.

Even today practitioners and commentators alike remain to a noticeable degree frozen in attitudes left over from Lighthill's era. These attitudes play down AI's intrinsic role as a vital bridge between the technology of computation and the biology of cognition. To the extent that today's knowledge engineers show listlessness in making co-operative use of what cognitive and brain science has to offer, we are seeing the legacy's residual effects even within the profession.

Validity of content. Lighthill's "ABC" offered a potentially useful classification of the broad melange of information sciences and technologies.

A Advanced Automation

B Bridge activity (also Building Robots)

C Computer-based CNS research.

He went on to praise A and C for their health and promise, but identified B as a transient and artificial linkage, doomed to failure.

This prediction, not borne out by subsequent history, was based on one rather startling idea, namely that to seek a bridge between A and C is not only without merit but is actually harmful. According to Sir James' later testimony, his self-perceived role was to warn against A–C cooperation and to discourage all AI activity which expressed enthusiasm for interdisciplinary involvement of this kind. In a published interview in *MI News* (Vol. 2, No. 5, February 1986, pp 3–4) he confirmed that he had seen his mission as "damping down the enthusiasm for a particular way of trying to go forward—by cooperation between biologists and computer scientists, and freeing them for the kind of development that was really necessary at the time on both sides". He did not clarify how the wholesale shut-down which followed his advice and the removal of livelihoods and working facilities from trained AI scientists was supposed to free them for an unspecified "development that was really necessary". His idea, I now believe, was that they should

look for places elsewhere, either in established automation groups or in established neuropsychology groups, according to individual orientation. In the event, most found places overseas.

Probably Sir James did not intend these consequences. He was plainly unaware, and was still unaware in 1986 as the *MI News* interview shows, that far from being a harmful excrescence upon AI, A–C bridging activity is at its heart. Surgical removal was life-threatening. As Professor Stuart Sutherland remarked at the time, Lighthill's B should more properly have been designated "Basic" AI and treated as suitable for solicitous nurture rather than for excision.

In the same interview, Lighthill made clear that the biologists he referred to in the above-quoted passage were neurobiologists. The ranks of British AI did not at the time of his report contain any neurobiologists. But there was a growing number of gifted and influential students of cognition and perception. Among them was the distinguished psychologist Professor Richard Gregory, who worked with us for some years in Edinburgh before leaving for Bristol in 1970. Circumstantial and contextual evidence suggests that Sir James was using the term "neurobiologist" in a very broad sense intended to cover a range of psychological as well as brain sciences.

Lighthill's A, B, C categories are in themselves useful and non-controversial. Problems begin when they are seen as a way of partitioning the field of AI itself, rather than identifying it as B with A-oriented and C-oriented flavours. In the latter case we would simply see AI as the field which exists wherever studies of computation and cognition intersect. The taxonomic problem, however, pales when one realises the precise nature of Sir James' central tenet—namely that co-operation between workers in A and workers in C is distracting for both and should be discouraged.

In this paper I suggest that if such a tenet looks nonsensical, it has that look simply because it *is* nonsense. The nature of computation and of cognition is such that they do intersect. Every worker in either is aware of this, whether or not he or she elects to work in the intersection zone. Artificial intelligence is our name for that zone, and it is in its essence a co-operative endeavour between knowledge engineers and cognitive scientists, i.e. between AI workers with computing backgrounds and AI workers with neuropsychology backgrounds. Each needs the other's specialist knowledge and insights. It can bring nothing but benefit if specialists of either side pay careful attention to models developed and tested by the other. I propose to support this position under five headings.

- Early A–C co-operation in pre-Lighthill AI in Britain.
- Operational validity of the Lighthill doctrine: effects of applying it.
- Knowledge engineers' blind spot in Alvey and in Japan's 5G project.
- Clearing the blind spot: the "superarticulacy" phenomenon.
- Frameworks of the future: can Government help?

First, however, I must make some brief observations on Martin Lam's contribution (this issue) to the discussion.

Lam's commentary. Lam plainly has constructive aims. I doubt however whether much is gained by repeating Lighthill's error that a distant observer is well placed to taxonomize material of which he lacks first-hand knowledge. In academic and industrial research and teaching, we are charged with systematizing today's profusion of information-processing activities with a core of clear scientific principles. Accordingly we do our best to discourage classifications based on fashion or the social utility of the day. Methodological discipline is not helped by Lam's deciding whether or not object-oriented programming, say, or intelligent robots, exemplify AI, by using external criteria such as whether practical applications yet exist, or which fields profit from them. Such an approach to classification is reminiscent of pre-scientific thinkers, or of those who lack professional training in a given field. Classification by surface similarities rapidly leads to a tangle of anomalous groupings—as if, for example, we grouped together bats and birds because both are warm-blooded fliers, or dolphins and sharks because both swim and have teeth. For some purposes this may be workable. But even a hobbyistic study such as train-spotting ordinarily depends upon a more principled, evolutionary and functional approach than Lam's classification by bits and pieces.

Perhaps for this reason, compounded by his omission of most of the field's recent advances, I found myself unable to make much of his commentary. So I have thought it better to proceed directly, rather than risk what might be a misrepresentation of his argument through failure properly to grasp it.

Early A–C co-operation. Richard Gregory is distinguished as a psychologist and physiologist of visual perception and as a creative engineer. From his participation in the Edinburgh project, the AI scene was enriched by contact with one of our generation's best biological minds. In the Acknowledgements of the 1971 paper in the *Computer Journal* "Tokyo–Edinburgh dialogue on robots in artificial intelligence research" (Barrow et al., Vol. 14, No. 1, pp 91–95) we stated:

A special debt is acknowledged to an overall philosophy contributed by Professor R. L. Gregory. His dictum "the cheapest store of information about the real world is the real world itself" was a major part of the original motivation, and the emphasis laid upon the role of internal models in Gregory's analysis of perception continues to be central to our work.

The FREDDY project's 1973 level of attainment won international renown, and is detailed in two papers in *Artificial Intelligence* (see References); also Ambler, Barrow and Burstall in S. Watanabe's *Frontiers of Pattern Recognition* (Academic Press, 1972, pp 1–29). During his 1972 visit to Edinburgh, Lighthill asked to see the robot. He did not, however, wish to have the system operated, or explained to him by the project leaders. In these circumstances it was perhaps inevitable that the ideas he gained of the objectives of such work were vague:

Incidentally it has sometimes been argued that part of the stimulus to laborious male activity in "creative" fields of work, including pure science, is the urge to compensate for lack of the female capability of giving birth to children. If this were true, then Building Robots might indeed be seen as the ideal compensation! . . . the view to which this author has tentatively but perhaps quite wrongly come is that a relationship which may be called pseudomaternal rather than Pygmalion-like comes into play between a Robot and its Builder.

My own assessment of the axeing of the FREDDY project which followed in February 1973 can be found in my book *Machine Intelligence and Related Topics* (Gordon and Breach Science Publishers, 1982, pp 219–220). The project, still remembered world-wide, could hardly be more eloquent testimony to the value of precisely the interdisciplinary approach which Lighthill felt was its flaw. Apart from Gregory's own inputs to the planning and earlier stages, the contribution of Harry Barrow, a trained psychologist, was central, to say nothing of Gregory's colleagues from the Cambridge psychology Department, James Howe and Stephen Salter, whose contributions, sustained to the end, were individually crucial to FREDDY's success.

Experimental robotics was but one of a rich variety of AI-related studies then thriving in Edinburgh's School of Artificial Intelligence. A densely documented report to SRC from my own Department of Machine Intelligence alone (*Experimental Programming 1973*, copies available from this author) lists nearly 80 technical papers, books and postgraduate theses for the period 1968–73, under headings including Computation Theory and Logic; Programming Language Design; Heuristic Programming; Robotics. During the same period Christopher Longuet-Higgins' Theoretical Psychology Unit was laying foundations of the highest quality for what subsequently became the Edinburgh Cognitive Science Department, and in Sussex the Centre for Cognitive Studies. If we add Bernard Meltzer's Department of Computational Logic, seed-bed of the world-wide development of Logic Programming and J. A. M. Howe's Bionics Research Laboratory which continued and developed what Richard Gregory had started, then truly it must be acknowledged that here was a research power-house indeed. The power came precisely from the very mingling of cognitive and computational ideas which Lighthill felt it to be his mission to "damp down". Indeed the recorded effects of the prescribed "damping down" treatment is sufficient in itself to confirm this view (see Fleck, 1982).

Operational validity. If an adviser's desire is to influence things in a certain direction, but when implemented in good faith his advice turns out to have the opposite effect, then in operational

terms the advice must be judged defective. Sir James' intent was undoubtedly to benefit the development of AI in Britain, slimming and purging it so that the sound parts might gain in health. But when assaying a cure, it is of the greatest importance to have indisputable evidence that sickness is present in the first place. Those who lived through these events, and also those who have studied the deeply researched chronicle and analysis by the historian of science James Fleck (*see* References), must be tempted to ask "Where was the sickness?" The old American saying seems pertinent: if it ain't broke, don't fix it!

The unblinkable fact concerning the Lighthill Report's operational validity must be seen in terms of dispersion and mass exodus. Trained British talent migrated overseas, not because our young people wanted to leave, but because the collapse of funding forced their hands. As often happens, the most gifted, and thus most saleable, workers are the ones who preferentially migrate, thus compounding the loss to the home country. So it simply will not do for Martin Lam to write "Given the world-wide scale of work on AI we can rule out the possibility that the Survey was self-fulfilling, i.e. that buds have not blossomed because they were nipped by Lighthill—or, more properly, by the scientific establishment, which could have disagreed with him." It may conceivably be that the starvation of an infant field in Britain does not harm the further development of this same field in America or Japan—although we must remember that in 1973 the UK accounted in sum for a rather substantial slice of the world's top quality AI work. But what comfort can that bring to British researchers, to British sponsors or to the adviser himself, since the surgeon's plan was that the scalpel should benefit the scientific health of the home community, in which ten years of home resources had been invested? In addition, it is improper for Lam to say that members of the scientific establishment "could have disagreed", thus implying that they did not. Some did, notably Sir Eric Eastwood FRS, Chief Scientist of the General Electric Company and a Council member of the SRC.

Knowledge engineers' blind spot. Almost every day, new interdisciplinary centres spring up in one country or another, involving the intertwining of specialisms drawn from a rapidly proliferating mix: knowledge-based programming, computational logic, automatic program synthesis, machine learning, computer vision, robotics, theory of knowledge, cognitive psychology, computational linguistics, psycholinguistics, human-computer interface, computer-aided teaching, neurocomputing, brain science, . . . the list goes on. Rather than weary ourselves with the retrogressive nature of the Lighthill doctrine, we should note that the world has meanwhile moved on and try to do the same. However, Lighthill's shotgun divorce has inflicted certain conceptual disablements, some of which continue to hold back full recovery. In view of the scope of interest of this journal, it seemed to me an interesting exercise to review just one specific area in which the post-Lighthill divorce of A from C resulted in a potentially winning insight being lost to the knowledge engineering community, with immensely costly ill effects for British information technology. For this I shall select a particular item from the repertoire of "what every cognitive psychologist and neuroscientist knows, but engineers don't", namely that the trained expert lacks introspective access to the brain procedures which encode his acquired skill.

Among the brain's various centres and subcentres only one, localized in the dominant (usually the left) cerebral hemisphere, has the ability to reformulate a person's own mental behaviour as linguistic descriptions. It follows that using the "dialogue acquisition" method the knowledge engineer can directly sample from the human expert only those thought processes accessible to this centre. Yet most highly developed mental skills are of the verbally inaccessible kind. Their possessors cannot answer the question "What did you do in order to get that right?"

This same hard fact has been belatedly driven home to technologists by the rise and fall in commercial software of the "dialogue-acquisition" school of expert systems construction. For example, the world's currently largest expert system, Brainware's BMT, installed for Siemens last year, comprises >30,000 rules. It was not, and could not feasibly have been, constructed by dialogue acquisition. In common with all the AI companies known to me in Britain, Scandinavia, W. Europe, USA and Japan which are actually making money (i.e. all debt and R&D investment

paid off), Brainware routinely uses structured induction as the main systems-construction engine. Alen Shapiro's (1987) monograph provides a fully worked academic exemplar of this methodology.

As a small but growing offshoot of the software industry, a new craft of rule induction is now the basis of successful commercial operations. In several cases known to me the payback to clients from such inductive learning exercises runs into millions of dollars per year. These tangible rewards accrue from relatively small outlays to commercial suppliers who are able to extract inarticulate thought processes from decision-data and serve them up to the client in articulate form. Production of installed code at rates in excess of 100 lines per programmer-day is the rule rather than the exception. This largely explains the profitability to be observed in this particular niche of commercial AI.

The contrasting disappointment which has overtaken the dialogue-acquisition school, after well publicised backing from government agencies of Japan (*see* Michie, 1988) and other countries, could have been avoided by acceptance of statements to be found in any standard psychological text. The chapter on cognitive skills in John Anderson's *Cognitive Psychology and its Implications* (1980) opens with the words:

- 1 Our knowledge can be categorized as declarative knowledge and procedural knowledge. Declarative knowledge comprises the facts we know; procedural knowledge comprises the skills we know how to perform.
- 2 Skill learning occurs in three steps: (1) a cognitive stage, in which a description of the procedure is learned; (2) an associative stage, in which a method for performing the skill is worked out; and (3) an autonomous stage, in which the skill becomes more and more rapid and automatic.
- 3 As a skill becomes more automatic, it requires less attention and we may lose our ability to describe the skill verbally.

In the light of point 3, failure of attempts to build large expert systems by dialogue acquisition of knowledge from the experts can hardly be seen as surprising. It is even less surprising when we have in mind that, in qualification of 2(1) above, many skills are learned by means that do not require prior description. A cognitive scientist might immediately volunteer a further piece of information, namely the following. In enquiring of a given skill how much it is likely to be subliminally encoded and how much available to introspection, just ask one question: what is the characteristic rate of the recognize-act cycle? If less than a second, you can be sure that the top performers can tell you nothing. If more than a minute, then some strategic and causal principles can probably be extracted from them to build into your model. But the tactical component will still be inaccessible to dialogue.

Yet popular notions run strongly counter to the relative absence from a learned skill's "automatization" phase of the ability to describe "how to do it". A concrete example may therefore be useful. Rather than illustrate with some cognitive skill remote from common experience. I have taken from M. I. Posner's *Cognition: an Introduction* an everyday instance which can easily be verified:

If a skilled typist is asked to type the alphabet, he can do so in a few seconds and with very low probability of error. If, however, he is given a diagram of his keyboard and asked to fill in the letters in alphabet order, he finds the task difficult. It requires several minutes to perform and the likelihood of error is high. Moreover, the typist often reports that he can only obtain the visual location of some letters by trying to type the letter and then determining where his finger would be.

Imagine, then, a programming project which depends on eliciting a skilled typist's knowledge of the whereabouts on an unlabelled keyboard of the letters of the alphabet. The obvious short cut is to ask him or her to type, say, "the quick brown fox jumps over the lazy dog" and record which keys are typed in response to which symbols. But you must ask the domain-experts what they know, says the programmer's tribal lore, not learn from what they actually do. Even in the typing domain this lore would not work very well. With the more complex and structured skills of diagnosis, forecasting, scheduling and design, it has been even less effective. Happily there is another path to

machine acquisition of cognitive skills from expert practitioners, namely: analyse what the expert does; *then* ask him what he knows.

Clearing the blind spot. It is worth noting a second payback to the client, which could not be supplied by users of other technologies. I refer to the fact reported by Shapiro and Michie (in D. F. Beal's *Advances in Computer Chess 4*, Pergamon, 1986) and now a commonplace, that rule-structured representations of these subconscious skills can be delivered as a side-product. Such pocket how-to-do-it theories are much valued as transparency and documentation aids by the in-house specialists employed by client firms.

What is involved here is a phenomenon termed "superarticulacy" (Michie, 1986). A suitably programmed computer can inductively infer the largely unconscious rules underlying an expert's skill from samples of the expert's recorded decisions. When applied to new data, the resulting knowledge-based systems are typically capable of using their inductively acquired rules to justify their decisions at levels of completeness and coherence exceeding what little articulacy the intuitive expert can muster. In such cases a Turing Test examiner comparing responses from the human and the artificial expert would have little difficulty in identifying the machine's. For so skilled a task, no human could produce such carefully thought-out justifications! In my Technology Lecture at the London Royal Society published in 1986, I reviewed the then-available experimental results. Since that time, Ivan Bratko and colleagues (1989) have announced a more far-reaching demonstration, having endowed clinical cardiology with its first certifiably complete, correct, and fully articulate, corpus of skills in electrocardiogram diagnosis, machine-derived from a logical model of the heart. In a Turing imitation game the KARDIO system would quickly give itself away by revealing explicit knowledge of matters which in the past have always been the preserve of highly trained intuition.

We saw earlier that facts already well known to cognitive psychologists were digested by commercial knowledge engineers working out of reach of the Lighthill legacy. Unhampered by Alvey's narrow IKBS interpretation, they went on to revolutionize profit-making techniques in commercial knowledge engineering. The implications of superarticulacy, a phenomenon discovered on the computer-oriented wing of the field, and forcefully demonstrated in the KARDIO work, should now, it seems to me, be digested by the cognitive psychology and neuroscience communities. Surely there are potential laboratory tools here of great power. They could be used to prepare "X-ray photos", so to speak, of the structure of those "blind" skills which support high-speed decision-taking as perfected for example by the pilots of helicopters and combat aircraft. I have recently been engaged, with my psychology-trained colleagues Michael Bain and Jean Hayes-Michie in an experimental study of just this possibility, using successively more complex variants of the well-known "pole and cart" inverted pendulum problem. Preliminary findings have been positive and appear in *Knowledge-Based Systems for Process Control* (eds. Grimble, McGhee and Mowforth, 1990).

Cognitive psychologists use the term "cognitive skill" for intensively learned intuitive know-how. Knowledge engineers call it "compiled knowledge". Collectively, such skills account for the preponderating part, and a growing part, of what it is to be an intelligent human. The areas of the human brain to which this silent corpus is "compiled" are still largely conjectural. Detailed studies should be vigorously pursued, preferably as a combined operation involving those on both wings of our science who share a common interest in subliminal modes of thought, and the elucidation of their logical structures.

Frameworks of the future: can Government help? In the guidance and support of basic research, including long-range applied (or "strategic") research, Britain has a widely respected record together with administrative mechanisms which have been tuned to this function over long periods of time. It seems to be at the fatal moment that an area is classed as applicable technology that its troubles at the hands of Government agencies begin. As far as concerns (machine-oriented)

artificial intelligence I believe that simple sociological facts, when combined with the evidence of the past 17 years, suggest that the insightful funding of innovative technology in an area of this type is not an art that civil service departments are likely to master. Far better, I would suggest, to cease the pretence that Government servants, established companies, entrepreneurial developers and University researchers are all brothers of the self-same breed, with common attitudes to such things as innovation, markets, accountability, risk, reputation, rewards, scientific truth and the rest. They are not, and (if they are all doing their jobs properly) cannot be. They like and hate different things, strive for different goals, have different hopes and fears, and different bases for mutual suspicion. Perhaps a Government department's help for technology R&D would be better exercised by manipulating from a higher level the motivational systems of the market, e.g. *via* well considered systems of tax breaks. But this still leaves unanswered the fundamental question. In Governmental ends-means analysis, tax breaks and the like can provide only the means. To determine ends requires us to identify dependable sources of scientific advice.

That we are still a long way from such a goal was evident from a recent forum on "Science and Government" of leading politicians and scientific advisors of the past and present (for overview by William Golden, *see* References; also reviewed by Susan Watts in *New Scientist*, 10th March 1990, pp 55–59). Meeting at the Weizmann Institute, Israel, in December 1989, the forum decided that no-one gives objective scientific and technological advice. Among the conclusions was that analysis should look at more than one option, and that factual findings, where possible, should be distinguished from value judgements.

However, as I wrote in the *New Scientist* in 1981 (23rd April, pp 241–242):

... we should not overlook a sprinkling of unsatisfactory results from this more cautious procedure. The British decision in 1912 to abandon all work on heavier-than-air flight, re-deploying resources into balloons, was the outcome of much patient work by Lord Esher's Sub-Committee of the Imperial Defence Committee. Documentary evidence had been sifted and expert witnesses cross-examined. Mercifully Esher had the courage to reverse himself in the following year, telling the Cabinet that he and his advisors had made a serious mistake.

As often in real life, the strongest moral to be drawn seems to be the old one that you can't win them all! But you can at least keep trying, and hope to do better, possibly even to advise better, next time.

References (only those references are here listed which lack adequate identification in the text)

- Ambler, AP, Barrow, HG, Brown, CM, Burstall, RM and Popplestone, RJ, 1975. "A versatile system for computer-controlled assembly" *Artificial Intelligence* 6 129–156.
- Ambler, AP and Popplestone, RJ, 1975. "Inferring the positions of bodies from specified spatial relationships" *Artificial Intelligence* 6 157–174.
- Bratko, I, Mozetic, I and Lavrac, N, 1989. *KARDIO: A Study in Deep and Qualitative Knowledge for Expert Systems*. Cambridge, MA: The MIT Press.
- Fleck, J, 1982. "Development and establishment in artificial intelligence" In: Elias, N. et al. (ed.) *Scientific Establishments and Hierarchies*. Dordrecht: Reidel, pp 169–217. Republished in: Bloomfield, B (ed.) *The Question of AI*. London: Croom Helm, 1987, pp 106–164.
- Golden, W, *World Science and Technology Advice to the Highest Levels of Government*. Forthcoming.
- Lighthill, J, 1973. "Artificial Intelligence: a general survey" In: Flowers, BH (ed.) *Artificial Intelligence: a Paper Symposium*. London: Science Research Council, pp 1–21.
- Michie, D, 1986. "The superarticulacy phenomenon in the context of software manufacture" *Proc. Royal Soc., London, A* 405 185–212. Republished in: Partridge, D and Wilks, Y (eds.) *The Foundations of Artificial Intelligence*. Cambridge: Cambridge University Press, 1990.
- Michie, D, 1988. "The Fifth Generation's unbridged gap" In: Herken, R (ed.) *The Universal Turing Machine*. Oxford: Oxford University Press.
- Shapiro, AD, 1987. *Structured Induction in Expert Systems*. Turing Institute Press in association with Addison Wesley.
- SRC Engineering Board Computing Science Committee, 1972. *Policy and Programme Review of Research and Training in Computing Science*. London: Science Research Council.