



National Library
of Canada

Bibliothèque nationale
du Canada

Canadian Theses Service

Service des thèses canadiennes

Ottawa, Canada
K1A 0N4

NOTICE

The quality of this microform is heavily dependent upon the quality of the original thesis submitted for microfilming. Every effort has been made to ensure the highest quality of reproduction possible.

If pages are missing, contact the university which granted the degree.

Some pages may have indistinct print especially if the original pages were typed with a poor typewriter ribbon or if the university sent us an inferior photocopy.

Previously copyrighted materials (journal articles, published tests, etc.) are not filmed.

Reproduction in full or in part of this microform is governed by the Canadian Copyright Act, R.S.C. 1970, c. C-30.

AVIS

La qualité de cette microforme dépend grandement de la qualité de la thèse soumise au microfilmage. Nous avons tout fait pour assurer une qualité supérieure de reproduction.

S'il manque des pages, veuillez communiquer avec l'université qui a conféré le grade.

La qualité d'impression de certaines pages peut laisser à désirer, surtout si les pages originales ont été dactylographiées à l'aide d'un ruban usé ou si l'université nous a fait parvenir une photocopie de qualité inférieure.

Les documents qui font déjà l'objet d'un droit d'auteur (articles de revue, tests publiés, etc.) ne sont pas microfilmés.

La reproduction, même partielle, de cette microforme est soumise à la Loi canadienne sur le droit d'auteur, SRC 1970, c. C-30.

**ON DERIVING PERCEPTS AND PRODUCING MOVEMENT:
THEORETICAL STUDIES OF PERIPHERY-BOUND KNOWLEDGE PROCESSES**

by
Jean-Roch Beausoleil

Thesis submitted to the School of Graduate Studies
of the University of Ottawa
in partial fulfillment of the requirements for
the degree of Doctor of Philosophy

Ottawa Canada
August 1986

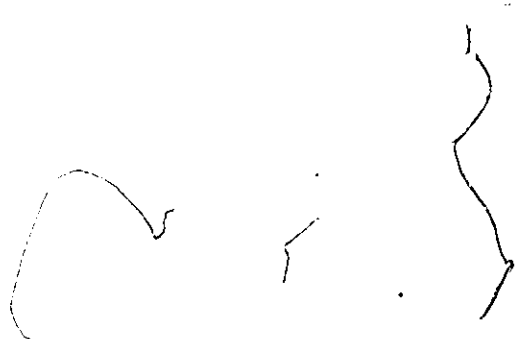
Permission has been granted to the National Library of Canada to microfilm this thesis and to lend or sell copies of the film.

The author (copyright owner) has reserved other publication rights, and neither the thesis nor extensive extracts from it may be printed or otherwise reproduced without his/her written permission.

L'autorisation a été accordée à la Bibliothèque nationale du Canada de microfilmer cette thèse et de prêter ou de vendre des exemplaires du film.

L'auteur (titulaire du droit d'auteur) se réserve les autres droits de publication; ni la thèse ni de longs extraits de celle-ci ne doivent être imprimés ou autrement reproduits sans son autorisation écrite.

ISBN 0-315-40683-6





UNIVERSITÉ D'OTTAWA
UNIVERSITY OF OTTAWA

à mes parents, en reconnaissance

ABSTRACT

At the root of Fundamental Psychology lie questions concerning human knowledge. The following work is set in the context of this understanding with respect to the computational management of *periphery-bound processes*, namely the derivation of visual percepts and the production of movement. We first develop the general argument behind the computational approach to knowledge processes, from which it follows that periphery-bound processes are also to be viewed as knowledge processes. These processes take the form of *grouping* in the visual case, and of *degrouing* in the movement case. Proceeding from there, we present an example of a *working* model of visual line detection instantiating the principle of grouping. This is followed by a more extended discussion of the principle of degrouing. It takes place, in the first instance, with respect to a hypothetical periphery of the visual kind. In the second instance, the computational demands and constraints associated with the development of a working theory of degrouing are further expanded in the context of movement control. These are initially studied in the light of past contributions to movement studies, stemming from what are described as the *performance-, system-, and technology-driven approaches* to the problem of motor control. For various reasons, these are found wanting from the point of view of their potential contributions to the foundation of a working theory of movement production (and consequently of degrouing) that would aim at a human level of performance. In the third instance, we present two models for the control of an arm: the first for a simple 'egyptian' type of effector, and the second for a humanoid effector. A distinction is introduced between two kinds of control schemes: the *extrinsic* type and the *intrinsic* type. The extrinsic type is the one favoured in Robotics, as well as, indirectly and to a lesser extent, by Neurophysiology and Psychology. We argue that a theory of movement control aiming for a human level of performance or production should be a model of the *intrinsic* type; however as yet no such model exists. Our first model of degrouing (or control) is of the extrinsic type, while our second, more ambitious, control scheme is of the intrinsic type. The intrinsic nature of our second model rests upon the fact that it is based upon a *representation of the effector as a global object of control*; this object is an *invariant* representing the complete movement potential of the effector. In that context, movement production and control are based upon the system's capacity to specify transformations of this representation's characteristics. A computational formulation of the problem of motor equivalence, or command invariance, is forwarded within the framework of this second model in terms of path control. The second model is shown to be functional through a simulation. Some possible extensions of the model are then briefly discussed.

ACKNOWLEDGEMENTS

I wish to express my gratitude and record my indebtedness to

- my supervisor Dr Claude Lamontagne, for his challenging attitude and inspiring support during many years, but above all for having created a vigorous and lively research context;
- my colleagues and friends Pierre Tremblay and Jean-Pierre Delage, partners in wonder;
- the National Research Council of Canada and the Ontario Ministry of Education, for financial assistance.

The greatest goals are those which are beyond the human conception of the time, and anyone who limits himself to the ends which he can clearly foresee is restricting himself unduly.

Kenneth Craik (1966, p. 90)

CONTENTS

0 The Computational Approach to Knowledge	1
0.0 Knowledge and computation	1
0.1 Knowledge representation	10
1 Periphery-bound Knowledge Processes I: Deriving Visual Percepts	20
1.0 Grouping	21
1.0.0 Strategic issues	21
1.0.1 Tactical issues	29
1.1 The helicoidal ring	32
1.1.0 The model	33
1.1.1 The simulation	36
1.2 Concluding remarks	40
2 Periphery-bound Knowledge Processes II: Producing Movement	41
2.0 The problem of control	42
2.1 Approaching movement control	53
2.1.0 Performance-driven research	56
2.1.1 System-driven research	74
2.1.2 Technology-driven research	89
2.1.3 Conclusion	100
2.2 Elements for a working model	101
2.2.0 Degrouping	105
2.2.0.0 Strategic views	107
2.2.0.1 Tactical points	115
2.2.0.2 Some basic constructs	120
2.2.1 Models for production	126
2.2.1.0 A simple extrinsic control scheme	128
2.2.1.1 An intrinsic control scheme	136
2.3 Extensions of the model	149
Conclusion: Grouping & Degrouping	153

Appendixes	156
Appendix I: Helicoidal Ring Simulation	158
Appendix II: Note on the 'hexconv' Function	197
Appendix III: The Simulated Arm	202
References	214

INDEX OF FIGURES

Chapter 1

Figure 1. Geometry of layer rotation	33
--	----

Chapter 2

Figure 1. The 'egyptian' arm	129
Figure 2. Control principle for the 'egyptian' arm	131
Figure 3. Problem of generalization	133
Figure 4. Extension of the 'egyptian' arm into 3-D	136
Figure 5. Geometrical basis of the 4 d.f.s/2 sffs effector representation	138
Figure 6. Screen display of the simulated arm and control scheme for the second model ..	140
Figure 7. Example of command	142
Figure 8. Path control (w.r.t. second control model)	144

Appendix I

Figure 1. Memory organization for simulation	160
Figure 2. Layer coordinate scheme used in projection tables	162
Figure 3. Tremor	164
Figure 4. Geometry of projection segment at layer μ	169
Figure 5. Areas of 'imperfection' in the circle	171
Figure 6. Stimulus circle on retina	175
Figure 7. Residue on a layer following circle analysis	176
Figure 8. Two stimulus line segments	177
Figure 9. Maximal mapping layer for segment B	178-180
Figure 10. Maximal mapping layer for segment A (1° tremor)	181
Figure 11. Maximal mapping layer for segment A (1,5° tremor)	182-184
Figure 12. Geometrical demonstration for the calculation of L	190

Appendix II

Figure 1. Cartesian coordinate system and hexagonal grid	200
Figure 2. Superimposition	201

Appendix III

Figure 1. The simulated arm	203
-----------------------------------	-----

Chapter 0

The Computational Approach to Knowledge

0.0 Knowledge and computation

At the root of Fundamental Psychology lie questions concerning human knowledge, its origin, nature, and evolution. The work herein reported also unfolded in the wake of these major issues, and although it has no pretension of breaking new ground at the level of generality usually associated with them, it should nevertheless be understood as an attempt, within specific problem contexts and according to a specific approach, at identifying - amidst apparently distant yet highly familiar knowledge phenomena - the most general principles involved and at formulating them in the most precise way possible.

The formulation of general explanatory principles is of course the basic goal of Science. One might be tempted to assume that the search for general principles underlying knowledge phenomena is a properly epistemological endeavour, and does not belong in Fundamental Psychology. But this would amount to ignoring the fact that modern Psychology grew out of a preoccupation with many problems that were first formulated by philosophers, in particular problems concerning knowledge. While Epistemology has come to be involved nearly exclusively with one type of knowledge, namely *scientific* knowledge, Fundamental Psychology for its part has largely remained attached to *common* knowledge and experience: perceiving, thinking, and acting.

One of the first major issues that was dealt with in modern Epistemology was the problem of the 'value' of knowledge, more specifically the problem of deciding upon the truth of any kind of knowledge. It was recognized quite early (at least from Hume onward) that the quest for solid *logical* grounds for assessing the truth of knowledge (scientific *and* common) had to be abandoned (cf. Popper 1972, ch. 1), but the quest for general principles nevertheless survived, giving up the search for a logical criterion of 'absolute value' for knowledge not precluding the search for general principles touching upon the other major aspects of knowledge - its origin, nature, and evolution - nor having any bearing on the existence of these principles. This shift is quite apparent in the work of Popper who proposed a very convincing model of the evolution of scientific knowledge.

On the other hand, the objective of uncovering the inherent 'logic' (or dynamics) of common knowledge or experience, or of explaining it, was embraced by Psychology. As a science, Fundamental Psychology's task is not different, for instance, from Physics' task of explaining the inherent 'logic' (or dynamics) of matter. This should imply that no strong a priori expectation need be entertained about the nature of the general principles implicated in common knowledge, for

instance that they be of a logical nature, although of course belief in the *lawful* character of experience has been and remains undisputed and fundamental.

Now care must be taken to distinguish between the impossibility of providing a logical argument in favour of the *certainty* of any kind of knowledge and the fact that knowledge *can* also be logically derived, through deduction for example, or more generally through the application of 'computational' means, the absence of a general logical *criterion* of 'value' in no way affecting the belief that knowledge is governed, at least partly, by (possibly but not necessarily logical) general principles, in the same way that the discovery of visual illusions (i.e. that visual knowledge is not 'certain') has not undermined the belief that quite general principles are at work in visual systems. On the contrary, illusions are studied in the hope that they will provide clues to those very principles.

Within Psychology the problem of the 'value' of common knowledge was never a central issue, except that *for* Psychology, as for any other science, research involves the adoption of an epistemological standpoint. This is crucial, as in any science but particularly here, not only because endorsement of an epistemology orients the overall approach to problems but because it so happens that *these problems are also problems of knowledge*. This leads even more acutely to the problem of the relation between methodological assumptions and assumptions regarding the object of study. Now we have stated that common knowledge is an authentic object of inquiry, and that it can be, so to speak, 'treated' in the same manner as any other object of inquiry. However there *is* a difference in Fundamental Psychology as regards the above-mentioned relation: the adoption of an epistemology not only guides research but can also be 'used' as a general theory of knowledge, in other words as a framework for understanding the origin, nature, and evolution of knowledge phenomena, which are the object of inquiry. This possibility does not arise in other sciences where of course the object of study is *not* knowledge acquisition and knowledge manipulation. Generally speaking one would expect methodological (or epistemological) assumptions to be 'carried over' into those concerning the knowing object if only for the fact that knowledge is involved at both levels, and given the implicit assumption regarding the *unique* nature of knowledge at *all* levels, i.e. the assumption that throughout the great diversity of its manifestations there is something that is *common* to all knowledge phenomena. Thus one set of assumptions naturally overflows into the other. This 'biasing effect' explains the importance of a clear statement of methodological assumptions in Fundamental Psychology. Note however that even though it is *possible* to regard these epistemological assumptions as a general *theory*, the level at which it occurs prevents it from being testable. In other words it is only theoretical in the sense in which it constitutes a possible answer to the problem of the nature of knowledge in general, *not* in the sense of a scientific theory; this follows from the fact that the epistemology itself provides the rationale for scientificity, for which

it therefore cannot be tested. Consequently psychological theories of knowledge will, at best, be 'in line' with an epistemology.

Our view is the following. Since there are no certain scientific theories, there can be no sure method of arriving at such theories, only a method of ensuring the selection of the best possible ones. This is the method of conjectures and refutations (Popper 1963). It states that Science is essentially problem-driven and theory-driven and that an inductionistic method, i.e. one that would minimize speculation and maximize observation, is *less* effective than one that would first maximize the conjectural effort. This follows directly from a recognition of the 'arbitrariness' of conjectures, in that since there is no way of judging a priori the value of a generalization, or theory, the best method is one that would enhance a posteriori its chances of being criticized, corrected, or overthrown.

Our own bias will now readily be discerned: just as the major efforts reported in the present essay will tend toward theoretical issues of knowledge acquisition and knowledge handling, so the major presupposition behind the ideas put forth and developed in the present work is that common knowledge is based on and consists mainly of what could be called 'theories.' Following this fundamental assumption about common knowledge, namely that it is also theoretical in nature, or is a function of 'theories,' we will be attempting to explicitly formulate the foundation and dynamics, i.e. the 'logic,' of some of them.

But what are 'theories'? How are theories and 'theories' similar, how do they differ? Popperian epistemology states that most of what we know as science is, directly or indirectly, *theoretical*. Our basic assumption is simply that this is the case for common knowledge as well. Now Popper himself has argued both that the *nature* of all knowledge is theoretical and that the 'logic' or inherent dynamism of Science can be generalized to all knowledge ("from the amoeba to Einstein"). We do not intend to take up or add to that here, but we do want to make very clear the implications of such an assumption with respect to our work. This requires basically that we look a little more closely at the kind of relationship we suppose to exist between theories and 'theories.'

We suggest that 'theories' are *like* scientific theories, not that they *are*. For example it might be opposed that scientific theories are *explicit* and strive for ultimate clarity, while common knowledge (or more precisely that which accounts for its occurrence) seems largely 'subjective,' that it must be a set of unknown mechanisms or unspecified 'procedures.' But this is the whole point of the suggestion, for if we *knew* explicitly the mechanism behind, say, visual knowledge there would be no problem of explaining it. But we do 'know' something that allows us to 'apprehend' a visual world, and something that allows us to produce actions and sentences and thoughts: we have 'theories.' It is true of course for example that we need not be able to *explain* the sentence we utter, merely *produce* it, while a theory of language will try to account for it. But accounting for

language, or explaining it, also means producing it, i.e. a 'good' theory of language is one that, at bottom, will 'reach' the level of the subject's 'theory' of language. This seems to hold in fact for all theories; for example a physical theory is one that produces 'physical' phenomena through derivations based on models. Thus the imperative for explicitness that one observes in Science appears to be met only partly in common knowledge, but is not a critical requirement; in other words it is not because in large measure what we supposedly 'know' is hidden or implicit that it must be of a different nature, or obey a different dynamics, than the theories of Science.

'Theories' are also like scientific theories in their purpose. From the point of view of Science the finalities of knowledge are understanding, prediction, and control. Leaving aside the motivation to explain and to solve problems, and concentrating on the *uses* of scientific knowledge, one can see in them extensions of the 'biological' reasons for knowing or reasons touching upon the purpose of common knowledge: to detect and to act, in a word to adapt. Our suggestion is simply that detection and action have the same purpose as prediction and control. Note that this is *not* intended as a kind of 'recursive' resolution of scientific and biological imperatives: whoever would argue that the objectivity of scientific knowledge places it de facto apart from the inalienable subjectivity of common knowledge, and hence definitely undermines any extrapolative attempt from one to the other, is forgetting that the objectivity of Science is highly dependent upon a 'subjectivity,' namely that of man.

Finally, and most importantly, 'theories' and (scientific) theories will be argued to share a common 'essence.' By this is meant the idea that both kinds of knowledge depend largely on *computation*, more precisely that the latter uses it explicitly while the former 'uses' it implicitly, and that as a consequence computations can serve as basis to construct models of 'theories,' much like they are used to express or formalize physical models. We would maintain furthermore that the rôle of computation in Science is central, as central as that of models, and that what makes a model a *productive* one is simply computation, be it of a logical or mathematical character (or even, apparently, neither of these, but we shall argue later that these differences are not fundamental and can be subsumed under a single type of operation). This feature would account for the pervasive deterministic assumption: the belief that phenomena are caused is 'manifest' in models and theories through the fact that a *computational* chain links 'causes' to 'effects.' We conjecture both of these propositions to be true of common knowledge as well. Thus, for instance, would spring (or be derived) the impression of the inescapability (or objectivity) of percepts, in the same manner in which the fact that A implies C proceeds forcefully from the antecedant assumptions that A implies B and that B implies C. Both are a function of computational rules, a belief that escapes from tautology only because the theory of computation is demonstrably recursive; its nature is such that it allows for a non-ambiguous discourse about itself.

We have started by assuming that common knowledge "is based on and consists of 'theories'" and this has led to the further assumption that their implicit dynamics can be understood, or modelled, through computations. According to this view much of what we know - what we think, see, hear, do, and feel - results from computations within 'models' or 'theories.' This, as we have said, is a starting point. In order to become a possibly fruitful approach to problems of knowledge it must be placed clearly within the context of a much more precise statement of the rationale behind computational modelling. In a word it is one thing to state general principles, either at the level of the constituent elements of models (which we have not done yet) or of their 'means,' as we have called their computational aspect, but quite another to link them up in *working models* (which we have merely indicated), for just as an *unworkable* model is unsatisfactory, so is a simple statement of the general principle/s that must be called upon to make it workable. We will pursue, therefore, with an overview of a possible connection that can be established between the computational formalism and (models of) knowledge. This should lead to a clearer view of the problem as a whole, and in particular will open up onto a fuller development of the major problem of the present thesis, namely *periphery-bound* (particularly visual and motor) *knowledge processes*.

Now, having set down the general explanatory framework within which the research reported here evolved, it still remains to relate the rationale of the computational approach with a possible understanding of the problems of common knowledge (henceforth referred to simply as problems of knowledge), and in particular with periphery-bound knowledge, which we have dealt with. As the computational approach is closely tied in with a formalism, namely *the theory of computation*, this will require that we discuss in a little more detail the basic ideas behind it and consider their relevance to a theory of knowledge.

As can already be surmised from the foregoing paragraphs, our aim is to arrive at detailed computational models of some knowledge phenomena. Because the problem it was designed to solve was a problem of knowledge, the formalism itself, through a generalization whereby it is interpreted as pertaining to *all* knowledge, can be seen as a sort of model for models, that is, it can be viewed as providing an explanatory context or foundation for specific models. It allows ideas and models to be precisely stated (implying, incidently, that ideas should be precise) and to become *formalized* models (as opposed to the *formalism's* model on which ultimately rests their meaning). In this way the commitment to a formalism is part of the overall conjectural effort; in our case at least, it also implies a commitment to the associated generalization. While this choice might be lengthily debated, and the formalism's model be questioned and criticized, it is not our intention to do so here, but merely to give sufficient insight into the theory of computation to reach a statement of the basic interpretation compatible with our assumption on the computational nature of knowledge.

Therefore we will now proceed to a discussion of the theory of computation, from which should emerge a clearer formulation of that assumption and the consequences pertaining to the nature of 'theories' and their explanation. In light of the interpretation, it is possible to consider the theory of computation itself as providing a first but highly specialized example of a 'theory.' This should cause no real surprise; that certain developments in modern Mathematics are considered highly suggestive with respect to problems of knowledge is probably intimately related to the fact that mathematical science seemed to have reached a point where it could rigorously ask and answer questions about 'itself.'

The philosophical expression of the computational assumption at which we arrived in the previous section could be felt to be unsatisfactory, if not ambiguous. This is no doubt due to an insufficient distinction between computation and the *theory* of computation, as it is unclear in that context how *what* we know (for example mathematical theories of the physical universe) is to be related to *how* we know (supposedly through 'theories'). Now "computation" is probably understood intuitively to refer to things like counting, arithmetic, or various types of calculations, essentially with numbers. If however "symbols" is substituted in place of "numbers," "computation" can be enlarged to include abstract operations like determining the differential of a function, and to encompass things like algorithms for the solution of a system of equations; it can even be shown to comprise methods of derivation and proof, and, finally, argued to cover *everything* we know. But to this we shall return later. In a word "computation" can be understood to mean a large part of what is *known* in Mathematics. Now counting, derivation, etc., have something in common: they are all *procedures*. In Mathematical Logic this notion of procedure (or effective decision procedure), of rule, operation, algorithm, or method, has been formalized and subsumed under the more precise concept of *computability*. This was achieved in three different but equivalent ways, namely recursion theory, the lambda calculus, and the theory of automata, only the last of which will be elaborated upon here.

To get a better grasp of how this was accomplished it is worth stating one of the original problems it was designed to solve, namely the Entscheidungsproblem: "to show that there can be no general process for determining whether a given formula u of the functional calculus K is provable, i.e. that there can be no machine which, supplied with any one u of these formulae, will eventually say whether u is provable." (Turing 1936, p. 259) The key elements in this statement with regard to the theory are "general process," "determining," and "machine," the first corresponding to rules, methods, or algorithms, the second referring to calculation, or decision, and the third to their formalization: computability is thus *defined* through a "machine," "automaton," or "Turing machine" as it is now called, which is specified unambiguously. To this we shall return in a moment. The very fundamental nature of this definition can be intimated from the fact that it

allowed Turing to describe a *universal* machine, that is, one capable of encompassing *all* definable machines: "It is possible to invent a single machine which can be used to compute any computable sequence. If this machine U is supplied with a tape on the beginning of which is written the S.D [Standard Description] of some computing machine M, then U will compute the same sequence as M." (ibid., p. 242) Although this was a *definition* of *computability*, it can be cast into the form of an hypothesis, namely that the notion of effective decision procedure is equivalent to that of calculability by a machine. Of course as such it cannot be proved, but to date no known mathematical rule, algorithm, or method can be demonstrated to be *undefinable* in terms of what are now called Turing machines. It might be worth mentioning that there *do* exist definable but uncomputable objects (numbers for instance, or problems, such as the one above), i.e. that the domain of the computable does *not* exhaust the domain of the specifiable; but the theory does claim to cover completely the notion of decidability or effective computability.

What is an automaton? How is it specified? This is readily answered by pointing out that there exists a class of devices that can be understood as embodiments of the universal Turing machine, namely computers. They are also known, through a simple association between information and the machine's symbol set, as information processing devices. As it turns out very few symbols are actually needed to achieve this realization, only two in fact - 0 and 1 - since "It is always possible to use sequences of symbols in the place of single symbols." (Turing 1936, p. 249), a method that had already been exploited in the Morse code, and used again in the construction of the ASCII code for example. Now symbols are but the *elements* of computation; the processing part can be quite as easily identified, a machine being specified essentially by the *rules* according to which symbols will be 'treated' by it: "The behaviour of the computer at any moment is determined by the symbols which he is observing, and his "state of mind" at that moment." (ibid., p. 250), that behaviour taking the form of a change of state and/or of certain actions, for example printing a symbol, or scanning the next one: thus it is defined as a 'programme' or 'table' setting out in detail the complete list of possible 'input symbols/state modification/output symbols' linkages, to which of course can then be associated a symbolic string, thus allowing 'programmes' to become elements within new 'programmes,' etc. As the embodiment of a *universal* machine, computers can emulate any specifiable machine, given an appropriate 'programme' printed on their "tape" (this is achieved by the processor, which we have just essentially described, which 'understands' binary strings - or symbols - associated with basic operations). Computers become information processors when, as is usually the case, the work they perform is either described from the point of view of the sets of their input and/or output symbols, that is, information, which can take a great variety of forms, from text to programmes to digitized signals of all sorts, or from the point of view of an informational goal to be achieved by processing, from sorting to compiling. Generally speaking,

processing then consists in 'acting' on the information, of performing *computations* upon it, and of producing certain results: we then understand the machine to 'know' what to do. It should be evident at this point that however complex or involved may be these tasks (think for example of the online processing of data required in a smooth spaceflight from liftoff to reentry) they can always be traced back or decomposed into much more elementary, and finally primitive, operations upon two symbols: this is the crux of the theory. This happens simply because these complex tasks are 'decomposed' into and 'rebuilt' from the elementary operations: they are programmed.

What was achieved by all this? Quite simply, and this is why it is called the *theory* of computation, a resolution or reduction (under the heading of "effective decision procedure" and then through its formalization) of a *diversity* of processes into *elementary* processes from which can then be *reconstructed* any and all cases in point. Now the Entscheidungsproblem was a specific instance of problems of *decision* within Mathematics, or of mathematical *knowledge*; as it turns out the answer was negative, so that it is not possible to *know* in general, at least according to *this* definition of knowing - in the sense of providing a description of a *method* to decide -, whether a given formula is demonstrable or not.

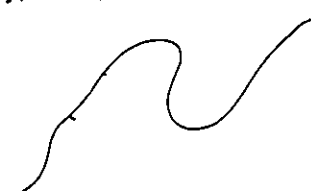
One may be wondering what all this has to do with the overall problems of knowledge. In order to address this issue we still need to provide an *interpretation* of the formalism, an interpretation which in fact can be intuited in the original paper (Turing 1936) and was later developed much more fully (Turing 1950). The psychological version of that interpretation is that the universal Turing machine is the formalization of a 'mathematical' *subject*, so that, for example, one can speak of its knowledge of numbers or of its ability to prove theorems, etc.; it has, or is in fact, a 'theory.' Of course this subject is highly abstract, in the sense of being far removed from the human or animal subjects that Psychology is usually concerned with. The generalization proposed by Turing (1950) consisted in extending *mathematical* (viz. computable) knowledge to *all* knowledge, or, in going from the mathematical subject to the 'natural' one: he argued that there was nothing special about everyday knowledge that could formally set it apart from the mathematical or prevent its conquest by similar means. However it is immediately apparent that the mathematical subject's knowledge is much better 'characterized' than the natural one's, but this points toward *theoretical* problems (of knowledge representation), not to fundamental difficulties with the basic generalization. It is true that the formalism, in the light of this interpretation, leaves room for the possibility of describing some instance of natural knowledge that would be uncomputable, but in order to do so we would still have to address the same theoretical problems; and even though we already *know* for example that most real numbers are *not* computable by a machine, this cannot count as a refutation of the basic assumption, for they are also uncomputable by *us*, while it is conceivable that a machine could *also* arrive at the same conclusion as we do. That is, uncomputa-

bility is decidable, as is also undecidability in fact.

More formally, Turing's generalization starts from the proposition that "what is decidable is computable" and extends into the proposition that "what is knowable is decidable," and therefore computable. This rests on an interpretation in terms of the subject. Thus we reach a clearer view of the idea previously set forth at the epistemological level: if to *know* is to *decide* and if to decide is to compute then what we know (i.e. 'theories') can be modelled through automata. What emerges as a consequence is a vision of the biological or human subject as a 'computer,' whereupon it follows that it knows essentially through computations, the nature and elements of which of course remain to be specified.

A 'theory' then is to be modelled by an automaton. That model becomes explanatory when it *productively* accounts for knowledge, the nature of these machines of course depending largely upon the type of processes that are supposed to underlie a given knowledge domain. The computational approach thus requires that symbols be identified and a set of operations be defined on them. This is the basic task.

Recalling now the earlier conjunction of symbols with information, and the subsequent illustration of the abstract concept of automaton through one of its major *technological* realizations, it will become apparent that it is equally possible, but this time in a conjectural manner, to construe *biological* realizations of automata by following through a further association between machine symbols and neuronal signals and between programmes and the micro-anatomy and physiology of nervous structures. The mathematical argument was first proposed by McCulloch & Pitts (1943) and later expanded by Kleene (1956) and Arbib (1960). Although such a development could be seen as pointing to a formal overture within Neurophysiology by providing a solid bridge across the 'neuro-epistemic' gap, it has not, unfortunately, paved the way directly to an understanding of nervous functioning as a whole, although it lies at the root of a line of thought that eventually led to what has become one of the major paradigms of Neurophysiology: single-cell recording, i.e. the recording of very *local* nervous signals associated with comparatively very *complex* recognition tasks. This only makes sense if it is assumed - following nerve-net theory - that the technique is a method for tapping into a level of *processing*, i.e. that the observed signal is incident upon a larger functional or computational whole identified as the nervous structure within which the recording takes place. We shall see in the next chapter an illustration of this reasoning which happens to coincide with an area of Neurophysiology where the single-cell technique has produced its most prominent results, namely vision. So although nerve-network theory provides a rationale for understanding nervous functioning, by showing how neuronal structures can assume computational tasks, it does not prescribe a method for embedding observational results within a computationally meaningful frame or model or theory, that is, a method for identifying the nature



of the processing in which these structures are involved. As such its level of generality is still very close to that of the original theory of automata. Now according to the view developed here nervous structures, along with their physiology, are largely to be understood as embodiments of 'theories.' But how are they to be discovered? How can the problem of their 'invisibility' be resolved? We believe some distance can be travelled along that road by a persistent effort at identifying the kinds of *problems*, and in particular the kinds of knowledge problems, that are supposed to be met by them, and pursuing the task, as far as possible, of the specification of methods for their solution. These methods give rise to the models.

0.1 Knowledge representation

To know is to decide: but what is decided and what is known? Most of our efforts having been concentrated on movement control and visual perception, we present here our view of the nature of knowledge problems as they arise in the context of these two domains, starting with a discussion of the most elementary knowledge levels that frames the problems in a quite constraining and precise manner, and leading to a tentative statement of the overall nature of knowledge processes in visual and control systems.

There are some immediately evident consequences that can be derived from the computational approach to problems of knowledge. These can be phrased in terms of overall requirements upon the features models or theories should demonstrate and pertain to what might be called their power. This involves, for example, speed of processing, by which is simply meant the idea that the relevance of models is to be judged, at least partly, but sometimes critically, by their capacity to produce results or arrive at conclusions in a time comparable to that of their natural counterpart. Thus an interesting perceptual or control model is not one that takes many minutes to decide what takes but a hundred milliseconds for a human subject. But even more importantly it involves the range or extent of the model's products with respect to that observed in human subjects. Thus an interesting model of control is not one that is limited to the emission of a single, however complex, movement, nor is a perceptual model that is limited to the recognition of a single, however complex, visual scene.

Of course here we were assuming that models' productions, be they of the order of percepts or movements, *are* comparable to the natural ones. But how is this possible? How can a computational model become a psychological one? Essentially, in the same manner as any model can: through a primary association between a set of concepts, thought to be critically connected with the phenomenon under study, and a set of symbols; the model then prescribes a set of (computational) relations, expressed as effective decision procedures, upon the symbols. The productive nature of models in that context stems from the fact that the postulated symbols (or associated concepts) are

primitive while the modelled phenomenon is not, but rather is derived from the model's activity. The problem of postulating such associations and then of defining within a model meaningful relations between the symbols are those of knowledge representation. Thus all knowledge is conceived as resting upon representations; and in particular in computational models of knowledge systems, this requires that representations be defined computationally. We will see shortly how this is possible in our case.

In the previous section we talked about a 'mathematical' subject's knowledge; in that context it is clear how the set of symbols and the set of possible relations among them is defined, most of the answers being found within Mathematics itself. Thus the problem of the mathematical representation of mathematical knowledge is easily solved right from the start. But this is a rather limited and highly specialized part of what we proposed to include under the heading of knowledge. However even if, on first encounter, the achievements that Mathematics can boast of, with respect to the representation of (mathematical) knowledge, could be felt at best only distantly related to the fields of common experience (seeing, hearing, moving, acting, thinking, and so on), and provide little in terms of clues as to what could serve as primary symbols and relational operations for vision, or just about any other instance of ordinary (non-mathematical) knowledge or experience, reflection shows this not to be the case at all: Physics, and in particular Optics and Kinematics, in its exploitation of geometry or the differential calculus, or more specifically in its mathematical theories and models, does provide elements, of a computational nature, that could be used as conjectural basis for perceptual and movement control representation: concepts and symbols relating to the physical world - in particular for objects and effectors in motion - could be judged relevant to the subject's knowledge and thus seen as pertinent with respect to the task of uncovering its nature. The relations conjectured to obtain among these concepts and symbols would then constitute the psychological model: from a mathematical subject we would now be defining a 'physical' one. Following this reasoning, it is therefore not surprising to find physical concepts playing (or having played) a rôle in perceptual and control models. In other words, suitably adjusted for the organism's peculiarities, knowledge of the world derived from Physics, coupled to the existence of a physical periphery with respect to which it is deemed relevant, is considered suggestive as to the way by which knowledge of the world could be derived from the theories and models of the subject; such models would work in a manner very similar to physical models. (Note that henceforth "theories" are replaced simply by "theories," which will implicitly refer to those supposedly held by or embodied within a *subject*, the objective, subjectless, use of "theories" being decidable contextually.)

This tendency is especially pronounced in the sensory and movement control, or periphery-bound, knowledge domains since they are linked to a physical periphery: eyes for vision and effec-

tors-for movement. Being the locus of the most primitive level of symbolization or computation, there apparently occurs at the periphery an informational binding with an identifiable subset of physical phenomena: light for vision and the development of force in effectors for control. This opens the door to a physical analysis of peripheral interactions. The problem of representation is solved at that level by seeking, on the one hand, the physical parameters of the events with which it is bound, as manifested, on the other hand, by a set of informational parameters characterizing those interactions. This apparently provides us with a critically important set of primary symbols upon which depend the answers to the overall problem of visual or motor representation. The periphery is thus the 'starting' point, and although it can be construed as a knowledge device, the *interface* (Lamontagne 1986) - or functional periphery - is one whose products are yet far removed from those inherent to the systems' highest objects or the level of knowledge usually associated with them: yet assuming one level of knowledge to be the functional *basis* for the other entails recognition of a sizeable gap between the peripheral or local, and the central or global, knowledge domains. But how is this 'distance' functionally spanned? How exactly does the periphery act as a *basis*? In our view there is but a single possible answer to these questions: production, i.e. hooking up the interface to a further knowledge structure, a system, that 'exploits' it in a productive manner. This leads to a very important refinement of the problem of representation to which we shall now turn.

We have stated that the sensory and motor peripheries, and more specifically their concomitant interface, act as starting point. From a computational point of view this is equally the case. (While this is evident for vision, it is not as clear for movement production; discussion of this point however will be postponed until chapter 2.) Now clearly, for the interface to act as functional basis it must connect to the whole that defines the function: in other words the fundamental assumption is that the interface's characterization, the physico-informational binding it operates, is *complete*, i.e. that any event, however complex, of the sensory or motor type, *can* be described by it. That is, any visual scene can be completely specified by a set of spatially ordered punctual retinal events, while any movement can be completely specified by a set of (quasi) rigid-body transformations. Of course this is *not* sufficient; it is merely a statement of the fact that "everything is there," *potentially*. This leads to a first formulation of the basic purpose of perceptual and movement control systems: their fundamental task as knowledge systems, in light of this observation, would consist in identifying, i.e. specifying, knowing, or producing, *actual* or current instances among the set of *possible* or potential ones that are determined by the interface. But how is this to be achieved computationally? And how does interfacial representation functionally relate to this process? How can a detection or control system *derive* or *produce* specific instances pertaining to the whole interface? These are the critical issues to be addressed in any attempt at building a computational model of detection or control or, equivalently, at solving the representation problem beyond the interface.

At this point it should be clear that, beyond the locality of individual interfacial entities, the *overall* structure of the interface is also an important factor in peripheral representation. Its impact however only becomes evident with respect to the system, upon which it 'imposes' a major constraint. This is due to the fact that instances *must* be known (i.e. computed) by the system in terms that are compatible with interfacial representation or characterization. As noted above, the type of description available at the interface is highly local in nature, while of course the type of description, of the *same* instance, found within a system is of a global nature, one that somehow lies beyond that level, and yet is bound by it. It follows that a global level of description must be a function of various *associations* of peripheral entities. This implies that interfacial description allows for the delineation of sets of interrelations and associations, upon which are directly founded the more global entities. Being complete, interfacial description cannot be complemented; it can only be transformed, meaning that while at the *retinal* level there is no way of specifying entities other than, or beyond, those allowed by the set of receptor cells, or at the level of *effectors* of specifying entities other than, or beyond, those allowed by the set of effector components, the whole detection or control systems are there to generate a set of 'new' entities that allow for a global description of the instance, that is, one pertaining to the periphery as a 'single' event. In short, interfacial representation is the most primitive level of computation upon which the *next* one depends, leading to a new level, upon which the *next* one depends, etc. And just as to the interface is associated a domain of variability, so to each new level will be associated a virtual set of events; on that view, representations breed representations. This in fact concurs strictly with the formalism's interpretation, and it corresponds functionally to what is meant by the *productive* nature of knowledge. Thus, the overall problem of perceptual and movement representation can be restated in the following way (cf. Lamontagne & Beausoleil 1984): it is the problem of imposing, in the context of recognized peripheral constraints, an overall 'significant' (i.e. significance-producing) strategy on the flow of computations, a flow that in the case of percepts is directed 'upward,' i.e. from the periphery to the center, and that in the case of movement is directed 'downward,' i.e. from the center to the periphery.

The purpose of periphery-bound knowledge processes therefore is to provide, in a manner respectful of interfacial (local) representation and overall structure, a global representation where essentially *many* peripheral events or entities are related to each other in such a way as to form *single* 'new' events or entities. For example we would wish a visual system to decide, or recognize, that an instance of stimulation, expressed as a matrix of receptor cells in various states of excitation, is an instance of a line, or a control system to decompose, or transform, an instance of reaching into a set of specific contraction commands in the various parts of an effector. In both cases it could be said that the systems *know* something about the interface's state through a process

of identification. (By definition there is no subject at the peripheral level, where events become objective; however we have no direct knowledge of these events, but the argument for objectivity can easily be transposed to the higher level of globality associated with the detection or control systems simply by observing that both levels share a common character of inescapability or determination. That, except for pure logical argument, this apparently ceases to be the case for the higher cognitive functions seems to be related to the fact that more avenues are open for interpretation for them than in the other cases, not to the fact that they are somehow undetermined or undeterminable.)

We have, so far, only hinted at a possible computational instantiation of representations. At the very low level of the periphery it was indicated that events are *characterized*. This was designated as the (interfacial) process of associating the physical parameters of peripheral events with informational parameters. Now a few remarks are in order at this point. First and foremost, it should be obvious that those informational parameters, or, more specifically, the values along those parameters conveyed at the interface, constitute the *computational* basis of representation: this is the process of characterization. Second, that it does not stop there, and furthermore that our understanding of the interfaces as genuine knowledge systems is derived from the belief that they partake of the same nature as higher levels of representation: thus the *principle* of characterization is generalized as the *modus vivendi* of the whole systems. On that view then, what systems do is to compute, along new dimensions of characterization based upon the interface's. This is inevitable, as we have seen: it means that 'information breeds information,' or, that all knowledge is based upon rule-governed or structured 'feeding' of information upon other 'informational' structures, a view that, in this context, is more clearly related to an equivalent formalization (recursion theory) and from which stems the recursive interpretation of the nature of knowledge. Thus a perceptual or control system hooks on to those informational parameters defined at the periphery. It can do nothing else. In this way the interface's determinism becomes clear: for it provides the most primitive computational parameters the systems will have to deal with. Given that this work is centered around questions of characterization, the problem of building a perceptual or control system boils down to the task of defining 'proper' new dimensions of characterization. And given, furthermore, that the systems' purpose is one of identifying peripheral instances (or at least parts of them) as wholes or unitary entities, then it follows that in order to achieve that goal, these new dimensions will somehow have to *relate*, as suggested earlier, the peripheral 'pieces' to each other in meaningful ways.

This provides a critical clue to the definition of these new dimensions: higher-level characterizations are functionally based on peripheral entities through procedures that compute relations *within and/or between particular dimensions* of the periphery's characterization. The main point here

is simply that at the interfacial level the existence of dimensions of characterization provides the systems with an element of generality which is exploitable by them: position, moment, amplitude, direction, extent, etc., all being more general than the set of specific values they momentarily assume at the periphery. More specifically, what is exploited is the fact that these dimensions are multi-valued, at least interfacially.

So far we have steered clear of an explicit discussion of the issue of the relation of the neurophysiological domain to these questions of representation and characterization, and therefore of the issue of the relevance of what has been advanced concerning these principles with respect to an understanding of the *knowing* brain, although it was indicated at the end of the previous section that the computational approach is consistent with what we know about neuronal function. Generally speaking, it was shown (McCulloch & Pitts 1943) that the brain can assume the rôle of a general-purpose machine: in the context of perceptual and movement control problems as we have discussed them so far, this means that *neurones* can theoretically be wired to compute any specifiable feature or attribute (characterization). That is, of course, assuming it is specified. But can this theoretical understanding be of any help in a search for the particular procedures that certain parts of the brain, i.e. the *observed* wiring patterns, supposedly embody?

The approach we tentatively favour with respect to this question, the nerve-network approach, consists of working in a context of minimal but non-negligible respect of neurophysiological data and concepts, as is quite obviously evidenced so far by our discussion of the periphery. This is a central feature of the nerve-net approach to computational modelling (as originally propounded by McCulloch & Pitts), which, even though it can never be guaranteed to hit the mark in neurophysiological terms, that is, lead to insights on the nature of the 'real' brain's functioning, is yet one that is hopefully leaning in that direction (and seems particularly well suited to achieve this), a course which cannot be argued to hold for other computational (e.g. robotic) studies of periphery-bound functions.

By "a minimal respect of neurophysiological data and concepts" is basically meant acknowledgement of the fact that the most important functional components of brains are neurones, along with the assumption that selected aspects of *neuronal* functioning are critical to a computational understanding of the *brain's* functioning. These relate to two fundamental tenets of Neurophysiology. The first one concerns the idea of the neurone as decision-maker or integrator and presumably derives from a tentative reconciliation, on the one hand, of the so-called all-or-nothing law with, on the other hand, the fact that nerve cells can receive, at any one moment, both excitatory and inhibitory inputs, the specific decision arrived at concerning whether or not a threshold has been reached. The second one, simply stated, is that anatomy is important. Although we have never seen it stated explicitly, a detailed study of the brain's structure, and in particular of the parts that connect

directly to the periphery, seems requisite to achieving neurophysiological understanding, by which we mean of course an understanding of the brain's business and of how it goes about it. Thus it seems to be assumed that structure (i.e. micro- and macroanatomy) is directly related to function, and that insight into function, if not derived from structure, will at least have to be consistent with it. (The concept of functional architecture as developed by Hubel & Wiesel (1962, 1979) is a recent instantiation of this idea, an idea that can in fact be traced back to the very beginnings of modern Neurophysiology - cf. Cajal 1894.) The integrative character of neurones along with their inherent connectivity (most neurones affect and are affected by other neurones) can suggest new types of integration at the level of the brain's nuclei, and from there to nuclear groups (e.g. the thalamus) and specialized structures such as cortices, and finally at the level of coordinations involving the whole brain. These are the basic tenets from which the nerve-network approach is defined. Again, it is well known that in that context neurones have been simplified into a discrete threshold device. Now the theory requires that effective procedures be specified through a structure or wiring pattern; it does *not*, unfortunately, provide ways of identifying a procedure from a given structure; that must be tackled through conjecture. The important point however is that by definition the requirements of the network approach conform minimally to important, albeit simplified, aspects of Neurophysiology.

It should now be clear how the required conjectural effort on the 'computational' nature of periphery-bound functions is related to the network approach; by bridging the gap between the formal computation domain on the one hand and the neurophysiological domain on the other it opens up onto a vast area of virtual possibilities: it defines a framework for explanation. And if that is combined with the similar, although undemonstrated, argument for the psychological domain, the enlarged framework now comprises the ingredients for potential solutions to the problem of the 'interaction' of mind and brain, at least as it is now conceived.

The apparent distance separating the locality and simplicity of interfacial events from the globality and richness of central events or of experience to which we alluded previously has as counterpart in the network approach the apparent distance between the simplicity of local nervous events and the richness and power that is supposed to be achieved by nervous structures. This seems to relate to the fact that nervous events are of a single type, irrespective of their place of occurrence; in other words the problem seems to be in understanding how a very local event can have a *meaning* or represent something obviously far surpassing it. This is where structure enters the picture: what gives meaning to otherwise identical events is the *context* of their occurrence and that context is the structure to which belong the active neurones. In the next chapters we shall come across specific instances of this principle. For now we might add, referring to the previous discussion of the theory of computation where it was mentioned that only two symbols were needed to account for an

infinity of potential symbols, that similarly here events can also be defined through a *set* of (possibly simultaneous)-neural events, to which, for example, could respond or be associated a *single* neurone. This is already the case for a simple AND gate for example, but the principle can also be extended into quite sophisticated template-matching devices, as we shall soon see. (Note how the function of an AND gate is related to its structure and how the nature of its output is *not* different than that of its input. In fact this holds for the whole CPU.) Conversely a single neural event can trigger a multiplicity of others, each of which, in turn, can do the same, etc. The local/global axis along which is defined the problem of representation is thus transposed into the local event/global structure dynamics of particular networks making up a system, where the richness of symbols is associated with the richness of structures and the complexity of events that can occur in them starting with the local peripheral ones. (It is worth noting how, by default, the locality of peripheral events can also be related to the structure of the interface, as it specifies a context where events cannot be related to others of the same nature; this is another way of stating their primitivity. Conversely there must be a point beyond which sets of (local or specific) neural events can no longer be 'fed into' a further structure since by definition a decision procedure must have a beginning and an end: this seems to suggest that the terminal point, or final stage, of processing, at any one moment, corresponds to a state of widely distributed activity in the brain.)

Coming back now to the issue of characterization, one can see how the structure of networks, along with the properties of the constitutive neuronal elements, can be exploited to achieve specific computations. To pursue an earlier example, one can say that an AND gate computes or characterizes the 'andness' of the events to which it responds. In the next chapters we shall explore systems of detection and control that make explicit use of this principle. Note that it is possible to further connect a pair of AND gates to an OR gate, and so on, so that it is possible (as McCulloch & Pitts (1943) have shown formally) to compute any logical function on inputs; the point we wish to stress here however is simply that it is equally possible to interpret the OR gate's output as the characterization of a set of previous characterizations: this corresponds formally to the recursive interpretation of (structured) computations. In other words we are formally justified in adopting a 'synthetic' approach to computation where things are done piecemeal and gradually built up into more and more complex wholes; in fact, relating this now to the brain, and to the idea that nuclei for example are focal points of computation, it should not be too surprising that so few of their functions have been unveiled: for most of these nuclei receive their input from other nuclei and send their output to still others, each one, supposedly, carrying out a step in the complete travail which is, in view of the fact that the initial or terminal stages are unknown, consequently extremely difficult to identify in the vast majority of cases. This is true even when the fine structure is known: for example in the area of movement control the microanatomy of the cerebellar cortex was

essentially described long ago by Cajal (1912), but its purpose and functioning still largely eludes us (Pellionisz 1980). Except for the widespread occurrence of divergence and convergence, ensconced as anatomical principles, which suggests some functional importance be attached to them, we seem to be quite destitute of ideas on the (computational) purpose of observed wiring patterns, even where, as is the case for the cerebellar cortex, that structure is locally (relatively) simple and highly repetitive.

However this problem of relating a particular structure with a particular function, which seems to lie at the heart of much work in Neurophysiology, does not arise in the case of conjectural models, where the correspondence is by definition completely specified. For these network models, wiring or structure is *used*, just as the basic computational properties of the constitutive neuronal elements, as a tool for the solution of functional problems.

What are these problems? We have generally set them within the context of problems of knowledge and hinted at their more precise nature earlier in this section when we discussed the periphery. Concerning the general nature of knowledge, we suggested that representation was the key, to which was then attached the more precise idea of characterization as computational counterpart. This was then linked up with the neurophysiological domain, to which of course one is more or less strictly tied whenever periphery-bound functions are to be modelled as productive systems, through the network approach to modelling. We also proposed that the basic task of detection and control systems revolved about questions of identification, or of the creation (or production) of new entities bearing on the virtual pool defined by the interface. With every other computational node within a system (of which the interface is one) might also be associated a virtual pool from which the other nodes can be conceived as operating. This 'hierarchical' feature appearing to be at the heart of the analyses engaged in with respect to peripheral entities and of the systems as a whole, we might now introduce two general but 'typed' technical terms for speaking about it. In the case of detection or perceptual systems we will be speaking of their overall task as one of GROUPING, while in the case of control, and in particular of movement control systems, we will be speaking of their overall task as one of DEGROUPING. The 'type' of course refers to the view that one system is conceived essentially as operating upon *inputs* while the other is conceived as essentially producing *outputs*. We have, so far at least, attempted to maintain an 'untyped' view of the particular knowledge problems associated with periphery-bound or pre-conceptual (Lamontagne & Beausoleil 1984) representations, the communality perceived at the interfacial level being in fact the first indication of a supposed functional communality that presumably extends into higher-level sensory-motor and cognitive coordinations. However the work reported here has not yet reached that level, for it appears that achievement of a certain (so far unspecified) level of sophistication seems requisite for detection and control systems before they can start to meaningfully talk to

each other, or before final perceptual outputs can be coordinated with initial movement commands. Nevertheless it did unfold in the context of this wider concern.

The following two chapters are devoted to discussions of grouping (chapter 1) and degrouping (chapter 2), and concerned with the specification of systems that achieve them. In chapter 1, following a discussion of grouping in the context of vision, a simulated line detection system is presented as a working exemplification of the principle of grouping. This chapter is much shorter than the next one for reasons that will be explained there; however it is of central importance to the developments described in the next one. In chapter 2 we present a more elaborate introduction to the problem of degrouping, and specifically as it arises in the context of the control or production of movement. As this is the major problem we tackled, in chapter 2 we shall also be concerned with carefully placing our understanding in the light of the three major modern approaches to it. Finally, the basic elements for a computational model of movement production will be presented, along with a description of the model (and its simulation) as it now stands.

Chapter 1

Periphery-bound Knowledge Processes I: Deriving Visual Percepts

From the point of view of the conceptual evolution of the research herein reported, the importance of our concern with the perceptual side of periphery-bound knowledge processes derives mainly from the fact that it was the major theoretical introduction and springboard to our analysis of the problem of movement production (chapter 2). Not surprisingly, a considerable degree of conceptual overlap should be found between our general discussions of these two problems. However as suggested at the end of the previous chapter, even though the work on movement production proceeded and unfolded in the wake of a much larger and encompassing concern with periphery-bound knowledge, it has not, at the procedural or 'wiring' level, reached a point where both perceptual and movement systems could be brought together in a true sensori-motor system. Thus a sizeable measure of discrepancy will be also seen to exist between the specific perceptual and movement models or algorithms to be discussed here, although in a somewhat surprising manner there is at least one important meeting point between the perceptual model and the second control model.

The purpose of the present chapter is twofold. Firstly, we will introduce and outline the problem of grouping as defined in the context of visual perception; secondly, we will describe a particular line detection model, as reported in Lamontagne & Beausoleil (1982), and briefly expose some results of a simulation which we carried out. Full technical details of the implementation are to be found in Appendix I; for our purposes here we shall only be referring to, and extracting from Appendix I, those aspects that are directly relevant to the larger issues at hand, namely those connected with the 'extensions' of the model in the areas of visual experience, neurophysiological plausibility, and finally computational efficiency. Again, the line detection model itself is set in, and extracted from, the context of a much wider research problem (and model), namely that of visual spatiality and movement perception (cf. Lamontagne & Howe 1980; Lamontagne 1976). However we shall not be concerned with it here.

Thus, while essentially serving as a first concrete instantiation of the abstract ideas developed in chapter 0, the present one will above all serve to introduce the pivotal notions of grouping and degrouping on which such realizations directly depend; and while our concern here with the first, namely grouping, will mostly be limited to the conceptual ties it critically entertains with the second, namely degrouping, it should be clearly stated here that the main thrust of the rest of the present work is directed, albeit as an 'extension' of the problem of grouping, onto the problem of degrouping.

1.0 Grouping

Pursuit of the basic goal of computational modelling of periphery-bound knowledge processes, namely the development of models or machines that embody such processes, involves two broad categories of discussion and concern: we have labelled them the *strategic issues* and the *tactical issues*. They are used to organize the two essential but interdependent levels of preoccupation arising in the development of working models. If in the first category, pertaining to strategic issues, problems essentially revolve about the overall 'meaning' or 'purpose' to be sought by and instilled within *any* particular knowledge system, the achievement of a working model can clearly only occur through a complementary, but no less critical, preoccupation with issues of the second category, namely the set of tactical issues, that have to do with the detailed implementation of such directives. Now if it is by no means a simple thing to identify critical constraints allowing fruitful decisions about the overall functional characteristics of particular types of knowledge systems to be made, it is as equally difficult, once a tentative choice has been arrested, to proceed onward to the formulation of 'corresponding' procedures or algorithms. But again, as concerns the very elementary case that is the subject of the present chapter, and in view of the purposes stated in the above introduction, we will for the present merely be adopting extant views (while of course adding a few new touches of our own to a small part of them), the process of 'adapting' them to the area and problems of movement production being left for the next chapter. We will however be discussing briefly, at the strategic level, some questions on the relation between grouping and degrouping, keeping the detailed pursuit of the tactical issues of movement production for the next chapter. For the moment, the strategic decisions concerning the functional 'goal' of the visual system are directly based upon and derived from the original Gestalt ideas, in particular as expounded by Köhler (1947), but follow essentially the formulation in Lamontagne & Howe (1980) developed in the context of the problem of visual motion perception, while the specific procedure, i.e. the 'corresponding' computational tactic to be discussed in detail, along with some outgrowths of a simulation which we carried out, is derived from Lamontagne (1976) and Lamontagne & Beausoleil (1982).

1.0.0 Strategic issues

The concept of grouping proceeds from the strategic type of analysis of the problem of (visual) perception. Starting in the early twenties, the Gestaltists were the first to state and argue forcefully and clearly that the most fundamental, yet until then largely ignored, feature of visual experience is, quite simply, that it is organized. Köhler (1947, p. 71) conveys this point clearly in the following:

There is, in the first place, what is now generally called the *organization* of sensory experience. The term refers to the fact that sensory fields have in a way

their own social psychology. Such fields appear neither as uniformly coherent continua nor as patterns of mutually indifferent elements. What we actually perceive are, first of all, specific entities such as things, figures, etc., and also groups of which these entities are members. This demonstrates the operation of processes in which the content of certain areas is unified, and at the same time relatively segregated from its environment.

Discovering the laws of organization obtaining in the visual world, as the Gestaltists set out to do, entails that the organization first be described, or characterized. Accordingly it was proposed that the most general or, following Köhler, elementary quality (or "fact") of visual organization was grouping. (ibid., p. 85) Hence the laws of visual organization must bear upon grouping processes.

The Gestaltists are of course well known for various now classic demonstrations of these laws, but unfortunately all but forgotten as regards their equally if not more important attempts to transpose or redefine them into *working* models. Indeed Köhler's analysis for instance, to which we shall return shortly, as well probably as the whole gestaltist thrust in perception, was of such a nature as to naturally lead to and open onto such problems of 'implementation.' In this task, as far as it was recognized and pursued by them, the Gestaltists brought into play and largely favoured physical principles and metaphors that, again unfortunately, seemed not to have allowed for the development of working models for visual perception enjoying a degree of precision and power comparable to those associated with the physical phenomena to which they were otherwise systematically likening visual processing.

Even though we have designed our strategic and tactical categories of concern in the study of knowledge systems with the computational approach in mind, it should not be understood that they are relevant only in that context. Indeed, as hinted at above and as argued in Lamontagne & Howe (1980) on epistemological grounds, where what are here labelled strategic concerns correspond (more or less) to what they called the meta-theoretical level of enquiry, they transcend the adoption of any particular formalism (although the converse is not necessarily true). The question of whether or not they pertain directly or can be generalized to any and all scientific endeavours, or are restricted to the domain of knowledge phenomena which happened to serve as the occasion and basic impetus for its formulation, is not one we shall take up here. But be that as it may it will no doubt be agreed that the gestaltist conjecture on the fundamental nature of perceptual processes and phenomena is at once more general than any particular (e.g. biological) instance, and so can serve as overall guideline in their study (including the computational studies), but specific enough to indicate the kind of task which any one perceptual system is supposed to accomplish. In our view this clearly makes of the conjecture a noteworthy strategic candidate and one, we might add here, that has so far neither been supplanted nor refuted. It is to our knowledge the *only* extant statement on the functional 'purpose' of perceptual systems.

We shall now proceed, following Köhler rather closely, to outline the fundamental problem of visual grouping. This will be concluded with a brief description of one of the few 'strategic' results situated in direct line with the gestaltist efforts with respect to this problem. But before we go on it might be worthwhile, in order to clear up a possible source of misunderstanding, to point out exactly how the nature of grouping is in fact conjectural. The initial observations of Köhler (1947) and others on the apparently simple fact that percepts *are* organized, and that this organization can largely be accounted for, from a descriptive point of view, in terms of groupings (the appearance of *objects*) and segregations (the appearance of *different* objects), etc., tends to overshadow the fundamental conjectural step taken when it is thence advanced that the 'purpose' of perceptual systems must exactly consist in the 'derivation' or 'formation' of the contents of visual experience *as they are experienced*. Thus is born the corollary idea of grouping as a process, which goes far beyond the idea of grouping as a qualitative feature, however fundamental or useful, of visual experience. Nowhere is this more clearly or forcefully expressed than in the following statement in which an explicit link is made between experience and processing:

visual experience corresponds to the totality of self-distributed processes in the visual sector of the brain, and [] all relations in visual space of which anybody can be aware rest on functional relationships within this totality. (ibid., p. 126)

Again this might appear to be quite simple if not plainly redundant, as assuredly it does as long as one does not then take up the challenge, already implicit within gestaltist writings, of 'proving' the theory by constructing a model or working system 'conforming' to the 'laws' of Gestalt. Yet however great its impact elsewhere, Gestalt theory appears to have left few traces to follow in the wake of this challenge, which was only 'resurrected' in the late fifties with the advent of Artificial Intelligence.

As stated above the starting point of Köhler's, as of Wertheimer's (1923) (cf. Köhler 1947, p. 85), analysis was ordinary human visual experience; more specifically Köhler's observation that "In most visual fields the contents of particular areas "belong together" as circumscribed units from which their surroundings are excluded." (ibid., p. 80), and Wertheimer's recognition of "the fundamental importance of spontaneous grouping in sensory fields." (ibid., p. 85), constituted the key elements that eventually led to the belief that "grouping [is] an elementary sensory fact." (ibid.) Here the word elementary, it should be noted, does not refer to the idea that grouping is a primitive feature of perceptual processing - implying that there is much more to it than simply grouping - but rather to the idea, lengthily argued against the then prevailing "empiricist" views in Psychology that saw grouping as the product of learning, that it is a *primary* and *irreducible* characteristic of the visual system's functioning.

The visual system of course refers to the eye and what Köhler calls the "visual sector of the brain." (ibid., p. 78) By positing, on the one hand, on strictly experiential grounds and independently of any explicit consideration of the system, the nature of visual processing to be of a certain type or to tend towards a certain 'pattern' or 'style' of analysis, one is inadvertently led, given, on the other hand, a recognition of the existence and potentially critical rôle of the system in visual processing, to 'extend' these notions into the realm of the (nervous) system's functioning. This is demonstrated, at least in Köhler's case, through many references to the visual sector. The attitude he adopts towards the brain, in view of his overall hypothesis on the nature of visual processing (i.e. dynamic self-distribution of local interactions), is revealed in the following extract.

The fact that organization depends upon relations among the local stimuli makes it entirely clear that sensory organization cannot be understood in terms of local processes as such. On the other hand, we have no difficulty in understanding the rôle which such relations play in organization, if we assume that the organization of sensory fields pictures the self-distribution of processes in corresponding areas of the brain. Dynamic self-distributions maintain themselves by interaction among the local events. But we have seen that in all parts of physics interactions depend upon the "condition-in-relation" as given in the various parts of a system. Since the same holds for visual organization, we have every reason to believe that this organization results from the self-distribution of certain processes in the visual sector of the brain. (ibid., p. 98)

Given "the absence of sufficient physiological evidence concerning these processes" (ibid., p. 78) however, one is left with the task of deriving and testing "hypotheses about their nature [] only from the facts of sensory experience" (ibid.). But Köhler's insistence on the *dynamic* nature of visual processing, along with the (at least for us) very convincing picture he draws of its overall style, even though it could be argued to have minimized the importance of a possible neurophysiological 'input,' should not, on the other hand, also have rendered superfluous or pushed to the wayside a preoccupation with detailed models of specific processes. It is only in the context of a sustained effort, or "a careful study," in that sense, we think, that could ever be reached a point where one could "tell [] quite specifically what physical processes distribute themselves in the visual cortex." (ibid., p. 98) For us here however that is of secondary importance. What is more central to the issue at hand is to grasp clearly the elements of the problem as outlined by Köhler.

The first key element of Köhler's view as just mentioned is the idea that visual processing should be thought of as a dynamic affair. This implies that whatever is brought forth to account for it should 'display' a form of inherent 'activity,' one that can descriptively (*re*)produce the kind of dynamism thought to obtain within the visual system. To Köhler and many other Gestaltists the dynamism of certain physical processes seemed highly suggestive with respect to the kind of activity that should have occurred in the visual sector of the brain. Of course we have since then

abandoned these notions, and replaced them by the more modern information processing metaphor, whose dynamism and inherent 'plasticity,' although of a somewhat different character than that of certain physical processes, is none the less potentially as rich, if not in fact more so. While at the 'superficial' level of computations the two in fact meet, physical dynamic theories being highly mathematical, and hence 'subsumed' under the theory of computation, and while a Gestaltist might have objected that the idea in itself is much more important than the fact that it be 'computable,' it could be countered that even though the theory of computation per se does not as readily lend itself to the kind of powerful metaphor for visual processing as the physical models do, those latter models have failed to produce a picture of even very elementary levels of visual processing (e.g. at the level of line detection). There seems to exist a fundamental discrepancy between physical and sensory *organization*, a difference that *might* be irreducible, but certainly, as concerns the *sensory* sector, an organization that might be better explained in terms of the dynamics of computation. Even though gestaltist models have largely been abandoned, Köhler's insistence on dynamism has certainly not been lost in the ensuing shift toward an informational view of visual processing; it remains in fact very pertinent.

The second key element of Köhler's formulation consists of a particularization of the problem of visual grouping in more explicit physiological terms, namely those provided by the retinal level of processing. In that form it will readily be recognized as lying much closer to those which have become quite common in current functional research in vision. If on the one hand the observation of visual experience led to the notion of grouping as the most pervasive characteristic of the visual *world*, the problem of reconciling that 'fact' with the constraints inherent to the visual *system* is one that quickly arises. But this is exactly the point at which the notion of grouping, from being essentially a phenomenal construct, must give way to becoming a functional one. This is simply because, with respect to the entities or objects of experience, corresponding "local [i.e. retinal] facts are entirely indifferent to any formal relations which may obtain among them." (ibid., p. 98), or again, because even if "the geometrical relations among the various points on the surface of an object are to a large degree reproduced in retinal projection [...] each local stimulus acts independently [so that] so far as retinal stimulation is concerned, there is no organization, no segregation of specific units or groups." (ibid., p. 95) Admitting both facts to be genuine can thus only lead to the problem of grouping, which can now take the form of the following question: how does a *multiplicity* of independent local retinal events having certain implicit formal (and more specifically geometrical, as defined by the structure of the interface) relations among themselves come to be associated together and constitute or eventually give rise to but a *single* event (or object)? Stated in this way, it seems evident that the key to grouping lies in the exploitation of those implicit 'retinal' relations. However it should be noted that not all characteristics or relations among visual

objects can easily be brought down to the level of relations between *local* retinal events, for whenever arises a question about relations among *constituted* visual objects, upon which the Gestaltists also insisted, it clearly cannot be handled at the interfacial level, where they do not yet exist, but must instead be referred to the 'geometry' of some other part of the system. We shall see in the next sub-section how such exploitation can be carried out for the case of line segments.

Working in the context of the problem of visual motion perception, Lamontagne & Howe (1980), both on epistemological and computational grounds, carried the gestaltist principle of grouping one major step further. In a statement reminiscent of Köhler's above on what might be called the 'closure' or sufficiency of the chosen approach with respect to visual processing, Lamontagne & Howe (1980, p. 27) define "the task of any visual system as being one of

1. *grouping* visual entities (starting with a.v.e.'s [atomic visual entities, or single retinal events]) into higher level visual entities, under well defined criteria,
2. *characterizing* the visual entities so obtained with well defined features bearing on each one of these entities *as a whole*,
3. repeating this process until visual entities are obtained which represent adequately the physical environment."

In what otherwise became a highly recursive formulation (in which of course can clearly be recognized the 'mark' of the theory of computation - cf. chapter 0) of the basic 'modality' or 'purpose' of visual processing, one can also observe the persistence of the concept of grouping as lying at the heart of vision. However here, as opposed to Köhler's formulation, the type of 'interactions' assumed to be responsible for the advent of visual entities are no longer of a physical but rather of an informational type. More specifically *grouping criteria* as well as various *characterizations* are hypothesized to enter into play at every step or level in the process of deriving percepts. (As we shall see in the next sub-section, and following Lamontagne & Howe (1980, p. 27 ff.), not all characterizations need have served as grouping criteria, although of course a grouping criterion can de facto become a characterization of the visual entity it helped to engender.)

Of course visual motion perception, as a perceptual sub-problem, has to be fitted and dealt with within the confines of the very wide 'framework for processing' defined above. The whole of Lamontagne & Howe (1980) is devoted to a discussion of various "macro-issues" that arise and have to be critically addressed while pursuing a computational theory of motion perception. The pivotal theoretical question raised in that context centers on "how much grouping should be achieved before, and how much grouping should be achieved after, motion detection?" (ibid., p. 110) While it is intuitively evident that *some* perceptual object (a product of grouping) should exist in order for motion to be assigned to it, it is nevertheless far from clear *what* that object should be and *how* it is to be derived, especially when one has given himself as ultimate reference human visual (motion)

experience. With respect to the above three steps, and, again, in the context of visual motion perception, that question arises somewhere between 1 & 2 and 3; that is, at what point in processing will an adequate representation be reached so that motion can then be computed? An "adequate representation" of course is an eminently vague expression; however, making it more precise, even for the 'restricted' domain of motion perception, is a non-trivial task. The solution arrived at in Lamontagne & Howe (1980) entails dropping the until then unquestioned (even by Gestaltists) assumption of a strict correspondence between *physical* and *perceptual* entities; it proposes the following strategy:

the system will group a.v.e.'s [atomic visual entities] right up to the level of a SINGLE global visual entity, or VISUAL OBJECT, before motion (for "final representation purposes") is computed, this visual object embracing either many physical objects in the observed scene, or only one, or only a part of one [...]; this computational scheme has been labelled the SINGLE VISUAL OBJECT SCHEME (SVOS). (ibid., p. 110)

Three immediate implications on the matter of implementing the SVOS are then set out: they pertain to a system's ability to 'sort out' a given scene, so that "different parts of the scene (or stimulus environment) can be selected and "isolated" for processing" (ibid., p. 111); to a system's ability to 'zoom' in and out of a scene, so that "the current outcome of the motion detection process can be used to define the next moment's input layout" (ibid., p. 112); and finally to a system's ability to 'steer' itself so that "the "grouping strategies" controlling the visual background/visual object split can be readily changed in the flow of analysis." (ibid.)

Just as, at the strategic level, the whole of visual processing can be argued to rest on a single principle, namely grouping, at the level of computational tactics, and in view of the strategic decisions of the sort we have just outlined, one can be similarly concerned with implementing them with as few operators as possible. It is with one of these operators, or *functional primitives* (cf. Lamontagne & Beausoleil 1982), formally paralleling the small set of symbolic operators found in recursion theory, that we will be dealing in the next sub-section. As can be surmised, their function consists of grouping and characterizing. The fact that the nature of the entities thus derived can naturally form the input to further, and hence broader or more encompassing (with respect to the initial atomic visual/retinal entities), levels of analysis constitutes the tactical counterpart and complement to the implicit recursiveness of the general strategic course set down above. The operator in question, namely the HELICOIDAL RING, is thus extracted from a much more comprehensive model of motion perception in which it served as a building-block. Again, we shall only be concerned with one specific 'application' of the operator, that of line segment detection and characterization. For further actual and potential applications of this computational 'primitive' to

motion perception see Lamontagne (1976) and Lamontagne & Beausoleil (1982).

Before proceeding further, we shall pause to examine briefly, in the light of what has so far been advanced on the strategic issues in vision, the question of the relationship between grouping and degrouping. This will serve as a first introduction to the problem occupying the whole of the next chapter. As hinted at at the beginning of the present chapter, it will also allow us to outline the important conceptual lineage between grouping and degrouping. It should however be kept in mind that although our work on the problem of movement production was, at the strategic level, largely inspired by our knowledge of those strategic results reported above, they served mainly as starting point for the analysis, and that the results we have achieved, and the new problems we have uncovered on the way, have not yet allowed for the unification of these two concepts in terms that would now make it possible to speak of a single theory of periphery-bound knowledge processes. Yet we believe that the fact that some ground could be covered on the motor side of periphery-bound processes in terms that had already been shown relevant for its sensory side indicates that the path of unification may yet be successfully travelled.

Just as a visual object is experienced as a *single* entity while on the other hand stemming from a *multiplicity* of retinal events, so in much the same manner, a purposeful movement of, say, the arm is experienced as a *single* (voluntary) event while being composed of a (potentially) *large number* of peripheral events. (It should be pointed out that whereas in the case of vision peripheral or local events are clearly understood to refer to retinal events, the case of the arm presents a possible source of ambiguity with respect to the exact nature of 'its' local events. Let it merely be said here that the peripheral events we were referring to above are those component movements of the mobile *parts* of the effector targeted by a given command: thus for example a reaching movement (as a whole) implies, in the general case, that 'smaller' movements will develop in the hand, the forearm, the upper arm, and possibly even the shoulder.) Now clearly the presence in both of these instances of a 'numerical' discrepancy between central or experiential events and those non-experiential ones that are bound up with them at the periphery can certainly be highly suggestive as to the existence, on the movement side, of a problem whose 'logic,' if not logistics, is closely related to the one thought to lie at the heart of the perceptual-visual side.

Now just as groups of orderly arranged retinal stimuli 'compose' into a visual melody, just so sets of orderly orchestrated 'smaller' movements will 'compose' into meaningful movements of the whole effector. And just as in order to perceive the former a system will exploit the "formal relations which obtain among them" (Köhler 1947, p. 98), in order to produce the latter, one can conjecture that the movement system will take advantage of whatever 'constant' geometric features can be established to hold among the parts of the effector. An evident level at which this can be intimated in the case of the arm is of course the level of the fixed connexions between the effector

components. It remains to be seen however just how this can be integrated within a movement system. This problem will occupy a central place in chapter 2.

So far we have presented a picture where a movement command would 'correspond' to and partake of the same global nature as a visual object. In essence this picture, so far at least, is a 'negative' of the strategy adopted in the perceptual case, the main difference in processing lying in their respective 'direction' or 'motive' rather than in their fundamental fabric. That is not, however, a drawback; on the contrary it can easily be argued that in order for visual (or perceptual) and motor processing to eventually 'meet' or to 'talk' to each other, as they certainly must do, a meeting ground must be laid for them; the fact that the analysis of movement production is approached in a manner that apparently owes more to perceptual than to any other kind of analysis does nothing to rob the problem of its individuality: in fact it has been argued that by waving aside the direction of processing one can distill the fundamental problem as being one of 'associating diversities,' a task that can be construed to be a fundamental task of cognitive systems. (cf. Lamontagne & Beausoleil 1984) The fact that independently for *both* the sensory and movement domains and dependently at *both* the experiential and peripheral levels (corresponding) diversities can be recognized will probably be questioned by no one. That is why the problem of degrouping has been defined to consist of "traduire la diversité globale en diversité locale." (ibid., p. 189) Evidently the nature of global diversity is different from that of local diversity, but only from the point of view of levels of processing; relating the two, or translating from one to the other, or establishing a 'semantic' relation between them, is of course the purpose of both sensory and movement systems.

We shall presently leave the strategic terrain and turn to a discussion of some of the important tactical issues connected with the realization into (specific) working models of the large-scale strategic decisions that have been discussed above. Although this discussion will immediately be followed by a presentation of the perceptual model we have studied, it should be mentioned here that, following the case of the strategic views outlined above, the tactical views next developed will also largely carry over into the next chapter's treatment of movement production; in section 2.0 for instance a visual metaphor will be used in the initial discussion of the problem of control.

1.0.1 Tactical issues

Proceeding from the strategic to the tactical domain, a new set of issues arises. Having defined, at the strategic level, *what* overall 'purpose' or 'direction' the systems to be built must follow, new questions, of the tactical type, pertaining to *how* it is to be achieved naturally emerge. The most important of them have in fact already been touched upon in the previous sub-section; however we will discuss them here in a more systematic fashion, as these issues go beyond any specific instantiation, such as the one presented in the next section. It will be seen that while a certain

degree of precision exists with respect to some of these tactical decisions, some grey areas become manifest as they are tentatively pushed into as yet computationally uncharted regions of visual processing.

As stated previously, grouping, and eventually degrouping, are to be accomplished through mechanisms of characterization. Now the *creation* of a visual entity entails that some criterion or condition of existence of that entity, in a word a decision procedure, be 'applied' to a given set of signals. The 'satisfaction' by a set of signals of a given condition provides immediate grounds for a first characterization of the resulting entity, a characterization that is called *primary* since it *essentially* describes the 'outcome' of the computation. However in general the created entity can be more fully characterized, and again according to the satisfaction of certain criteria; now these features are called *secondary* characterizations since they are non-critical with respect to the existence of the entity to which they are ascribed. However their computational rôle, far from being secondary, is crucial to the pursuit of analysis, the primary characterizations *by themselves* constituting a computational dead-end. It is only in (possible) *conjunction* with the secondaries that the primaries can hope to contribute to further, as well as higher, levels of analysis. For example in the line detection model it will be seen that while the existence of the constituted objects, namely lines, depends upon the condition of linearity obtaining among its constituting elements, namely retinal 'point' events, these resulting entities, as is well known, can also be provided with a (visual) position, orientation, and length. From then on, one can conceive how combinations of these characterizations could be used to pursue the analysis into the analysis of curves and various shapes, etc.

But what about grouping itself? What is the mark of a grouping operation? Quite simply it is an operation in which for the creation of a *single* new entity or visual object there must have entered at least *two* entities of the previous level, or some level/s lower than the one at which it is occurring. (cf. Lamontagne & Beausoleil 1984) As in fact *many* points must enter into the derivation of a *single* line, the process of line segment detection must involve a grouping process. However it should be noted that not all characterization processes are grouping (or, eventually, degrouping) processes. Although the above definition is very wide the operation of arithmetic addition for instance would *not* count as a grouping operation because it fails to satisfy the 'level' requirement of the definition, since the resultant object, that is, the sum of a number of numbers, although it could be construed as a characterization of the 'sumness' of the input entities, is still a number, and not, for instance, a function. However the detection of a pattern within a series of events, as 'showing' for example the imminence of a world conflict, *will* count as an instance of grouping. In this way, grouping operations can be argued to quite readily extend into 'pure' cognition as well. In chapter 2 it will be similarly argued how degrouping can equally be 'projected' into the higher cognitive realm.

Coming back now to the issue of grouping criteria, we will first recall, from the previous section and chapter, how the interface in fact affords a first grounds for grouping. More specifically we are referring here to the retina and to the fact that its intrinsic and fixed geometry provides a natural framework for the specification of relations between the entities it defines. Formally speaking, these entities are themselves the product of a 'primary' grouping process entrenched within the fine structure of the interface, of which the local geometry is responsible for giving rise to the primary, or 'point' feature of these entities (according to a grouping process occurring within the confines of the receptor cells), as well as the secondary feature of intensity; positional characterization for its part arises out of the global geometry of the interface. There is a third secondary feature, namely moment of occurrence, but we shall not be concerned with it here, since the model bears upon a 'frozen' type of analysis, i.e. one pertaining to a 'single' moment, or more precisely one that is not concerned with characterizations involving time. (cf. Lamontagne 1976, 1986) Furthermore since by default we are dealing exclusively with extant (non-null) retinal stimuli and since the intensity characterization has been reduced to an ON-OFF scheme, intensity will be seen to play only an implicit rôle in the process of leaping from point stimuli to line segments; consequently the main burden of analysis will be on the positional values.

Now it is rather straightforward to attach an explicit positional value to individual retinal stimuli by transforming the retinal grid into a cartesian one. However there is another way of doing this using the idea that positional values are already implicitly assigned through the fact that a retinal point occurs at a specific place. In actual fact what this means, and the way it will be seen to happen in the model, is that, in order to be exploitable, positional information need not be explicitly specified; it can just as well be used 'as is.' Of course this requires that in that form it has to be 'understandable' by whatever is intended to use it. For the helicoidal ring, this 'raw' information is first of all used to detect linear relations and then exploited to derive line orientation as well as position; however line length is calculated explicitly. Of course there is a limit to the advantage that can be gained through the geometry of the interface. But the fact that at post-retinal levels of analysis certain well-defined structures (such as the helicoidal ring) are, via the structure of the interface, involved with the structure of interfacial input opens up directly into a possibility of exploiting *their* geometry as well for carrying the analysis into the next level, etc. In fact reflection shows that as far as the central nervous system can be argued to be a 'structure' system, very little in the way of explicit computation need occur within it: to phrase an expression with a gestaltist flavour, the brain would not use equations, it would be one. There is not, at this stage of theorizing, a general rule or strategy that we could offer as to the relative 'dosage' of explicit and implicit computations that should be striven for in the construction of particular systems; however inasmuch as one would allow oneself to be guided by the above comment on the choice apparently

favoured by nervous systems, the direction in which to lean is quite clear.⁰

Grouping criteria will usually serve as primary characterizations, and these latter will consequently usually be single-valued, as the fact that they bear upon the existence of an object (e.g. a line) or a feature (e.g. squareness) already suggests. However the secondary characterizations, such as length, orientation, and position of line segments, or size, position, and orientation of the 'square,' will on the contrary usually be multi-valued. Thus line segment detection entails deciding, in the presence of all potential and possible variations within and among the secondaries, on those cases which satisfy the single primary criterion of linearity, and furthermore, once this is achieved, of sorting out their respective secondaries. (For it should be evident that it is not possible to work directly from the secondaries since there is as yet no object on which they could bear.) In the context of this problem the essential feature of the helicoidal ring is the way in which a *structural* embedding or interpenetration of the primary characterization within the pool of virtual secondary characterizations is achieved for the case of line segments. The structural peculiarity of the model derives from the fact that it is a template-matching device.

1.1 The helicoidal ring

In the remaining two short sub-sections of the present chapter we first present an intuitive overview of the helicoidal ring in terms of the basic geometrical facts attendant upon this particular theory of line detection; for a much more detailed treatment of the model, as well as the whole theoretical context from which it is extracted, the reader is referred to Lamontagne (1976) and Lamontagne & Beausoleil (1982). In the second place, a synopsis of the main purpose and outcomes of a computer simulation of the model is presented; again, for a much more complete discussion of the technical details of the implementation the reader is referred to Appendix I. Finally, this second sub-section should be considered as complementary to, rather than a recapitulation of, the material in Lamontagne & Beausoleil (1982).

⁰Incidentally it is worth mentioning how first Mathematics, and then Computer Science - with the exception of non-discrete or Analog Machines - has gotten us used to explicit computation. This is changing however with the appearance of a new type of computer, namely the Connection Machine (Hillis 1985), in which, as its name implies, the establishment of variable connections between processing nodes would serve as a basic computational element or construct. Thus a 'tree' would no longer be implemented by means of pointers but rather directly specified by the nodal connections that define it. This allows for 'easily' implementing parallel processing, a notion, at least in terms of visual processing (and the helicoidal ring is a parallel processing model), that can also be traced back to the Gestaltists' image of physical processes and interactions which, as is well known, do occur 'all at once.'

1.1.0 The model

If a line in a plane stands at an angle α° with respect to one of the (arbitrarily chosen) cartesian axes of that plane then, keeping the line fixed, a rotation by α° of that plane along with its axes will bring the line parallel to the reference axis and perpendicular to the other. Thus, supposing that one is very good at detecting parallelism, the orientation of a line with respect to the frame of reference defined by an arbitrary coordinate system can be measured by rotating it and the plane in which the line lies (see Figure 1).

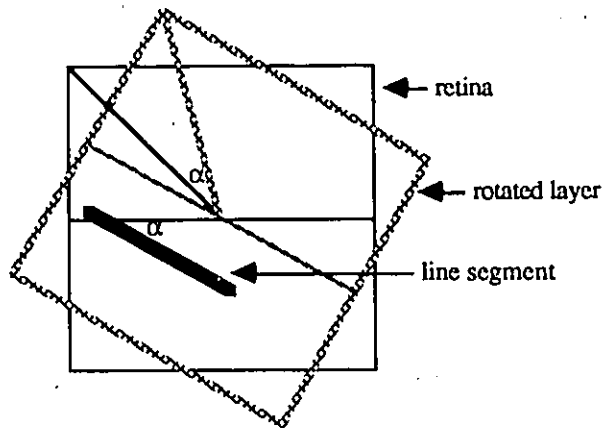


Figure 1. *Geometry of layer rotation with respect to a given line stimulus. Note how the line segment subtends the same angle with respect to the retinal frame's X axis as the rotated layer's does, rendering the same line segment parallel to the latter's X axis. (Not shown, retinal fields and layer tessellation.)*

Now suppose, furthermore, that in that planar universe, lines are *only* recognized with reference to *one* of those already defined by the coordinate axes. This would imply, for a line segment that originally stood perpendicular to that axis, that it would first only be 'seen' as a point, but then, as the plane system (plane plus axes) is rotated, the 'line' would grow in size until it reached a maximal value whenever was reached the parallelism stage.

Instead of actually rotating the plane system however, and having to keep track of the amount of rotation, one could start with a *stack* of already rotated planes on top of each other and then, for a given line segment, simply try out each one in turn until a largest 'line' is found. Again since for a given segment all planes will 'see' some 'line,' having to choose from among them the largest 'line,' or equivalently the best-fitting plane, can get one involved in lengthy comparisons (that in fact become impossible to carry out if more than one segment is projected at a time). Why not

instead simply let these 'lines' do it for us? How? By having *every* one of the 'lines' themselves (actually the weighted components of each 'line,' the weight being the total number of its components) 'travel' directly up and down through the stack of planes and *deleting* any of its 'parent lines' (actually again their components) which it encounters that are of smaller size than itself: since there is by definition none larger than the largest (and this is the one we are fundamentally interested in) it should be the only survivor of that elimination process.

If instead of a single line segment we now consider the case where there are two or more different ones, then it should be clear that the above procedure will work *individually* and *independently* for each one of them. That is, for each original line segment there should be, not necessarily on the same plane and not necessarily at corresponding places on these different planes, a largest or 'best line.' The actual plane in which a particular 'best line' occurs will provide by itself the measure of orientation, this measure of course being the same for any 'best line' that is found within it. Thus two 'best lines' found within a particular plane would have had to correspond to two original parallel stimuli, i.e. sharing the same orientation, but not position. The difference of position of course will be 'conserved' in a difference of position of the 'best lines;' this position can easily be generated by 'reducing' the 'best line' to a single signal, say the middle component of the 'line.'

This is the basic geometric theory behind the helicoidal ring. Yet from the abstract point of view of geometry the original stimuli already exist as lines. Not so in the 'concrete' terms offered by the visual periphery to which, it will by now have been guessed, the original plane where stimuli are found now corresponds. The visual periphery offers, in geometrical terms, inputs of just one kind, namely 'points.' It is then the task of the visual system (remember Köhler) to recognize among them, or derive from them, the existence of certain types of relations, namely, for the case at hand, linear relations. Now quite simply it is because each plane in the stack associated with the retina already knows something about linear relations and positions by being linearly 'biased', and because the whole stack of planes, by being twisted, knows something about orientation (which is why it is called a template-matching device) that the helicoidal ring works.

Thus *any* and *all* signals issuing from the helicoidal ring bear a message of linearity (just as any and all retinal signals bear the primary mark of 'pointedness'); *which* plane it emanates from signals an orientation; *where* on that particular plane signals a position; and finally *how large* it is (in explicit terms of a number) bears upon the size or length of the line. Thus is accomplished the basic task of recognizing line segments (through primary characterization) and qualifying them according to their secondary characterizations of orientation, position, and length.

We have stated that the retina provides the system with points. This is not exactly true, as a retinal object corresponds in fact to a small surface, which is only called a point because it is the

most primitive object in the system. Following upon the same 'logic' of discreteness, it should be clear that whereas geometrically speaking a given line has a specific, and potentially infinitely precise, orientation, such a level of precision cannot be reached in a finite system. That is, the helicoidal ring can have only a finite number N of layer planes, each one being twisted by an even fixed amount, say θ (such that, $\theta \cdot N = 180^\circ$), with respect to its immediate neighbours, and, like the retina, tessellated (not shown in Figure 1, but see Appendix I, Figures 2 and 3, as well as Lamontagne & Beausoleil (1982), Figure 2). It is the fact that in terms of the rotational geometry of the ring's basic structure the *first* layer plane has the *last* layer as immediate upward neighbour (and vice versa for the last layer plane) that the model takes on a (twisted) *donought shape*. The degree of torsion from one layer to the next corresponds to the orientational resolution of the system. The same reasoning applies to length: being of fixed and integral size, there is a limit to the length and precision of lines 'available' to the retina for projection onto the layer planes because of its, by definition, limited positional resolution.

What this amounts to for the workability of the helicoidal ring is a series of problems, known as projection problems. They arise as a consequence of the fact that in mapping or projecting a set of retinal points onto a given plane, one is in fact dealing with the 'interaction' of surfaces, meaning for instance that having to decide whether a given retinal point, on its way down into the 'ring,' does in fact 'meet' a given cell in one of the layers becomes a question of whether the two surfaces overlap significantly enough for a 'meeting' to have occurred. This problem and many other similar ones have largely been resolved by the introduction of refinements to the basic system, namely *retinal fields* and *tremor*, solutions in fact that turned out to be highly suggestive with respect to the workings of the human visual system, at both the neurophysiological and psychological levels (cf. Lamontagne & Beausoleil 1982).

Bringing all of these elements together more formally now, the helicoidal ring works according to four well-defined processing stages. In the first, the *projection* step, the whole of retinal point stimuli is individually mapped, and as many times as is necessitated by the number of tremor positions (all tremor positions plus projection occurring within a single retinal sampling by the system), onto every field layer or projection plane, adding a value of 1 to those field cells' content that are encountered by the retinal signal; now depending upon the size of retinal fields, a certain threshold must have been reached by these field cells in order for them to have been judged stimulated (or 'met' by a retinal 'point,' which now corresponds to a retinal field), a decision that is only taken after completion of the whole cycle of 'tremor' projections; upon taking this decision for a particular field cell its content is reduced to 1, after which it is again changed according to the number of adjacent field cells - adjacent, that is, with respect to one of the axes of the field layer - in

which the particular one lies. (Thus if a series of four adjacent stimulated field cells is found, each one's content will now become 4.) This in fact is the second phase, called *summation*. In the third stage, the *inhibition* stage, these weighted signals are made to 'visit' their vertical neighbours in a number of field layers determined by their current content; with force decreasing as a function of 'planar' distance, they annul any field cell content thus encountered whose content is strictly less than their own. And finally in the *reduction* stage each of the remaining strings of adjacent field cells in all layers is reduced to a single excited field cell, chosen as the one closest to the middle of the string. This constitutes the terminal form of the signal proceeding from the helicoidal ring.

Retinal outputs signal 'points,' while helicoidal ring outputs signal 'lines:' thus is achieved the leap from one to the other. What happens between these two events is a process of characterization (primarily of linearity, and secondarily of position, orientation, and length). The primary characterization is achieved through the grouping criterion of (linear) adjacency as embedded in the fine structure of the field layers, while the secondary ones are derived directly or indirectly from various other local and global structural aspects of the system: directly when it is computed implicitly (linearity, position, and orientation) and indirectly when it is computed explicitly (length). In whatever form it occurs however, this information can be used, within the *same* sampling period, to pursue the shape analysis of a retinal scene into the detection of curves for instance (now taking the form of patterns of helicoidal ring signals - see Appendix I for an example of a circle), or, by moving *across* the time dimension into the next sampling period, it can be used to compute various types of motion (rotatory, expansion, translatory - see Lamontagne (1976) for examples). Thus, even though the raw richness of retinal stimulation has been lost in the grouping process, that same process has given rise to new and more powerful information (in terms of the extent of the events which it signals as compared to those occurring at the retina). From being immediately periphery-bound, the analysis can now (and must in fact) become helicoidal ring-bound in order to accede, at least as concerns the 'frozen' type of analysis, to even more 'remote' levels of relations. The more subtle the kinds of relations to be extracted, the farther the system has to move away analytically from the periphery; but no doubt that is because even though these relations already (implicitly) obtain at the retinal level, achieving a visual description of them is *only* possible by a recursive mode whereby each new level of analysis added builds deeper and farther on the foundation laid by the previous one.

1.1.1 The simulation

Once the theory is clear, the next step consists of testing it. As described in the previous sub-section the helicoidal ring could not be tested yet however, since three values have to be established for the three parameters attaching to it. These are retinal size, retinal field size (related to

positional resolution), and orientational resolution, i.e. the number of layers in the simulated ring. A general (abstract) proof of the model, that is a demonstration of its effectiveness for all possible parameter values, is not only beyond our mathematical means, but would arguably be of little intrinsic value, given that the interest of the model lies within relatively well-defined bounds. For various reasons, some of which are explained in Appendix I, we chose the following: a retinal size of 96×96 cells, a field size of 3×3 receptors, and a 'twist' of 3° . This was the best, i.e. the 'largest,' helicoidal ring we could simulate at the time on the available system; nevertheless a closer study and familiarity with the model led us to establish the equation relating those three parameters' interdependencies, and we used these to calculate the dimensions of an 'ideal' ring (cf. Appendix I, section 3): it would require a retinal size of 228×228 cells, a field size of 5×5 receptors, and a 'twist' of 1° . Unfortunately such a ring, especially with the fast table lookup scheme used, would comprise what amounted at that time to prohibitive demands on computer memory (including disk storage). However one further interesting 'coincidence' (cf. Lamontagne & Beausoleil 1982), this time of the anatomical type, surfaced as a result of these calculations. Using Hubel & Wiesel's (1979) estimates of the number of "cortical units" in the primary visual area of the brain, as well as current estimates of the number of retinal cells in one eye, and dividing one into the other, one comes up with the approximate figure of 200. Given that we know, on empirical grounds, that human visual linear orientational discrimination is of the order of 1° (Jastrow 1892), which is the discrimination level used in the above theoretical calculation, such agreement is worth noting in the context of the neurophysiological 'extension' of the model. However the assumption, inherent to the division, of a homogeneous partition of the retina into equally homogeneous corresponding cortical fields is somewhat dubious, and should be kept in mind when considering this result; but about this, we shall say more lower down.

Now once some specific dimensions are decided and implemented, along with the four processing steps (in fact we limited this to the first three), the next one is to stimulate the retina. Again, just as it is practically impossible to simulate a 'general' helicoidal ring, it is impossible to offer a 'general' stimulus to the system. The specific ones to be analysed must therefore carefully be chosen in order to be computationally maximally demanding of the particular helicoidal ring one is working with. However, as in any empirical undertaking, these stimuli (or 'conditions') should also not be too complex, for fear of losing diagnostic power; the simultaneous interplay between the various structural-functional aspects of the ring, as 'applied' to a given stimulus, can, even in the case of relatively simple ones, rapidly lead one to loose track of events. On the other hand that is also a major reason for carrying out the simulation in the first place. In our case we chose two sets of stimuli, one being a pair of line segments (see Appendix I, Figure 8), while the other consisting of a circle (which, as far as the helicoidal ring is concerned, is a set of small line segments - the

circle's tangents, in fact -; see Appendix I, Figure 6). It should be noted that these inputs, in the discrete framework of the retina, were 'perfect' mathematical objects; there was no compulsion, or possibility, at this stage at least, to offer more 'natural' stimuli to the system.

As could have been predicted from the development of the ring (cf. *ibid.*), the largest source of difficulties in the simulation of the model was the projection stage. Here again an interesting (however obvious) fact stands out: the size of tremor is 'independent' of the size of retinal fields, for the very simple reason that cutting down the size of the room in which tremor should occur (from say two to one receptors) automatically cuts down the number of directions in which it can be exercised. As a corollary to this, our experience has shown that there is very little (except confusion for the ring) to be gained by increasing the tremor 'region' beyond a width of two receptors. Immediately connected to this is the size of retinal fields; as a consequence of the above argument it can now be argued more clearly that it should not be too large, that is, that little advantage can probably be gained from a field size surpassing a value of 5x5 receptors; to which, again, is directly connected the question of positional resolution within the ring, since by definition stimulus width is equal to field width. (Consult Appendix I, section 3 for a fuller discussion of these points.)

Now in the immediate vicinity of this question lies another consideration pertaining to the overall effect of tremor: from the point of view of the retina as a whole the translatory component of tremor exercises a homogeneous effect while the rotatory component does not, its effect being the more pronounced the farther away one moves from the center of rotation. For the retinal size with which we dealt this actually led to no acute problem; however it is worth mentioning that in the context not necessarily of a larger helicoidal ring but of an array of rings (which in fact amounts, from the standpoint of this problem, to the same thing) that the peripheral swings incident upon rotational tremor can become quite large, and that such a situation, if it were to arise in a 'normal-sized' ring, would be problematical indeed; it would give rise to what was called "double detections" (*ibid.*, p. 356). This is in fact the reason why the amplitude of rotational tremor was kept at a minimum; however that minimum is established with respect to a 'normal' ring, and if we now revert back to a consideration of the retina as a whole, the fact of limiting the *boundary* amplitude of rotational tremor to such a small size will immediately imply that a relatively large area in its center will *not* be tremored. Since this can entail difficulties for the ring, the only apparent way in which *both* constraints (that is, reducing the amplitude of the swing at the boundary and reducing as much as possible the area in the middle of the retina that is unaffected by rotational tremor) can functionally be satisfied is by *changing* the size of retinal fields, i.e., enlarging them as a function of radial distance. Although it would presumably be rather straightforward to calculate this increase, it would be an altogether different matter to implement it; we have not done so yet. This problem (along with its tentative answer) is the reason why the above

assumption on the homogeneity of retinal fields might have been too strong.

Mathematically speaking, the analysis of a perfect circle stimulus by the helicoidal ring should have produced a set of 60 pairs of tangents. But this is not what happened; only 35 pairs came out, the other 25, being of a lesser size (by just 1 in fact, but that was enough), getting eliminated during the inhibition phase. In Appendix I, the reason for these imperfect mappings (which follow a well-defined pattern - see Figure 5' there) are tentatively ascribed to the 'micro-geometry' of the ring, that is to the retinal, as well as layer, partition into *square* units. This problem is the problem of neighbourhood relations established by a given retinal tessellation, and 'carried' down into the layers through projection. In the simulation for instance we had to compute the field layer coordinates of projection of each retinal cell; now this was calculated *precisely* with the usual trigonometric function and then rounded out to the closest integer value (Appendix I, Addendum 1), again according to the usual rule, a rule, it should be added, that is particularly well suited to the case at hand, involving *square* units, but not necessarily particularly well suited for the problem at hand. To see this consider that in the square retinal tessellation case the decision that a retinal cell 'meets' a particular field cell, used throughout, might not lead to the same result overall as if, for instance a hexagonal tessellation were used, along with a suitably modified 'meeting' decision procedure; in fact there is reason to believe, since neighbourhood relations are much less ambiguous in the hexagonal case, that in that case overall mapping would be enhanced through a marked reduction in stimulus imperfections. Again, since this arose as a *consequence* of the reported simulation, we have not been able to test it, although in an initial study of the implications of such a fundamental readjustment of the micro-structure of the ring, an algorithm for quickly transforming cartesian into hexadecimal coordinates was developed (see Appendix II). It remains an interesting open question at this point whether our above ideal numbers would survive the shift; it is tempting to believe that they would, if only for the fact that the 'cartesian' ring, being far from ineffective (as the series of problems that are here successively raised does *not* intend to suggest, most of them having surfaced as a result of quite demanding tests of the model), *can* 'sustain' such a result, while the 'hexagonal' ring, by allowing for neighbourhood relations that would surely be even closer to those found to hold in the nervous system, would a priori have a greater chance of giving rise to values concurrent with it. This of course is only a tentative and in no way final step; indeed it is possible to imagine much more complicated projection designs, ones that would involve overlapping projections (i.e. *one* retinal cell mapping to *many* field cells), and so on.

In fact, it is our contention that this gradual perfection of the helicoidal ring could lead to an ever closer resemblance to, and hence understanding of, the visual cortex's wiring. As to the fundamental nature of the cortex's processing, the helicoidal ring already offers a number of explicit and precise avenues in which further empirical probing could be directed.

1.2 Concluding remarks

In this chapter we have covered extensive ground in the problems of visual processing. Starting in the very broad strategic terrain, where the problem of defining the 'purpose' of visual processing occupied a central position, the discussion then moved toward the problem of the tactical methods through which it could be achieved, which centered about the definition of the computational principles to be used in formulating a theory of processing; and finally, the operation of a specific working model was demonstrated.

Just as in moving into the tactical level of concern one is constrained by the decisions taken at the strategic level, one must similarly act in the context of those latter decisions while attempting to move into the area of specific and working models. Problems attaching to a particular type of concern can therefore develop as a direct function of those higher-level decisions; conversely they can, for the most part, only be understood in the context defined by them. There is of course no possible guaranty that solutions to these problems exist, that is, that the strategic or tactical decision/s that brought them about was the right one; indeed certain problems *might* only find their solution in a revision at the tactical or even strategic levels. However in the present chapter, although the discussion was pursued at all levels, the type of problem that stood out most prominently was of the computational kind (i.e. pertaining to the specific model); furthermore their solution was searched at that level only, as they did not seem to be 'acute' enough to render necessary an explicit reconsideration, and possible revision of, the adopted tactical, or even strategic, guidelines.

The importance of such apparently academic considerations should not however be minimized; in fact the flavour of the argument will be felt in chapter 2, especially in section 2.1, where it will be used to gauge the relevance of certain models found in the research literature with respect to the goal of constructing a *working* model of movement production.

The 'logic' of (computational) research we have just recapitulated has its direct counterpart in the pursuit and further development of the visual model. For as hinted at toward the end of sub-section 1.1.0, the task of 'expanding' the model into wider areas of shape analysis for instance would now have to occur within the primitive framework defined by the helicoidal ring; and although the helicoidal ring already suggests ways in which this could be undertaken, yet there is a possibility that some problem in shape analysis might arise which could render necessary a revision of that particular line detection scheme.

But that is not the direction of research we chose to pursue. While using the same conceptual tools as in the perceptual-visual case, we have rather attempted to analyse and understand the problems (as well as model some parts of their proposed solutions) associated with the 'other' periphery, namely the motor periphery. Presently, we shall turn to this.

Chapter 2

Periphery-bound Knowledge Processes II: Producing Movement

How is an apparently simple thing like reaching a door handle in mid-stride achieved? How is it that humans so precisely calculate movements of the arm that they very seldom miss reaching and grasping a great diversity of objects of so many different shapes and at so many different positions and orientations with respect to the body? Why can we so easily trace paths of various shapes, sizes, and directions in the air?

These are but a few examples that illustrate the amazing variety of performances generally subsumed under the heading of movement control, skills, and integration; yet, as for vision, we are rarely struck by what otherwise seems to be a natural feature of behaviour, namely the ability to carry out so many goals, actions, and intentions in a smooth continuum of movements. An important part of our world representation is visual and, being *in* that world, we seldom recognize its visual nature; in the same manner, a large segment of our reality is intentional, and revolves about our capacity to transform goals and purposes into complex series of movements: the ease with which this is usually achieved, the fluid interplay between intentions and realizations, the gracious way in which one realization melts into the next and one intention blends into another, and the extent to which the movement domain, and hence the intentional one, is meshed into a functional unit with the visual domain, all conspire to form an equally unquestioned whole, along with an apparently inescapable feeling of continuity as the background of everyday experience. The challenge to explain this interplay, and more specifically to conceive of a way in which movement, as a purely physical phenomenon, can be meaningful or correspond to a goal or an intention, and conversely to account for the transformation of an intention, as a purely psychological datum, into a physical manifestation, lies at the root of our endeavour.

We cannot of course claim to have achieved complete success yet. It turns out, at least according to our understanding and as hinted at in sub-section 1.0.0, that the task of imposing meaning onto movement is somewhat akin to the problem of extracting meaning from sensory inputs, particularly of the visual kind, a similarity that seems to depend (cf. section 0.1) on the presence in both cases of an *interface* whose purpose is to establish an informational connection with the physical-energetic domain. This interfacial transition, which we recognize as belonging to the most primitive (either first or last) level of visual and motor processing, appears to be quite definite: that is, any move away from the interface is a move into the informational; consequently this is where the problems have to be defined and solved. Following this path, and in accordance

with the adopted formalism, namely the theory of computation, we are first led in the present chapter to rephrase the basic problem of movement control in a way consistent with the computational analysis introduced in the previous chapter. Not unexpectedly our definition requires that a solution be expressed as a *productive system*, that is, one capable of realizing at least a part of the performance, and overall variation, observed in natural movements; this of course also entails the specification of a concomitant intentional variability.

The essential feature of our conception of movement control is that it involves a particular type of knowledge process, namely *degroupping*, much as, in the previous chapter, vision was conceived as resting largely on a process of *grouping*. The concept of degroupping is informally developed in the following section (2.0). Using a 'visual' metaphor, our discussion starts with a general view of 'interfacial' control from the point of view of the functional requirements that seem to be associated with it, which then unfolds into a preliminary discussion of the problems of movement control that effectively sets the stage for a critical review of the research literature associated with the problem of movement (2.1). In section 2.2 some elements for a working model of movement production are set out, along with a detailed description and discussion of the present state of the model.

2.0 The problem of control

Suppose one were set the task of writing a graphics programme that, much like a human, would accept as directives commands like "house" (or "cat," or "face") and then proceed to draw one. Even stipulated in this embryonic fashion, it immediately raises a whole set of issues relating to problems of description, identification, and representation. Let us look at some of these. The major thrust of the task seems to be to take a certain descriptor, for example "house," and to transform or use it somehow to eventually specify a certain subset of pixels (points) to be activated on a screen. But what subset exactly? Or, exactly what house? This is related to the 'fuzziness' of the concept of house, that has, so to speak, a lot of windows left open into which might be introduced a variety of missing pieces of information, that are themselves more or less precise: viewing angle, style, size, number of storeys, number of doors, windows, etc. Of course these specifications might be provided by the user, or they might be left up to the programme to decide: the latter case requires that it also have some store of default values for them, upon which it can call whenever it needs to in the intermediate stages of realization, these default values themselves possibly having to call upon further default values, etc. Right away it is evident that, even for the single case of "house," the programme would become quite involved; but if it is to work it must have this kind of generality in its representations, a feature that in fact is intimately related to the generality (or 'global imprecision') of the interface, i.e. the screen, that can equally well display houses, cats, or faces. That is,

the programme can be seen as exploiting or realizing some of the *potential* of interfacial representation (that, at a certain resolution, in fact covers the whole of static visual experience): it forces it, as it were, to 'behave,' at least visually for the onlooker, like a house, a cat, or a face; in a word, it *controls* the state of the interface.

It should be evident at this point where our hypothetical example is leading to: exploiting a reversal of the point of view of visual perception, where the task for a similar programme would be to *recognize* instances of houses, cats, and/or faces, to be achieved through a process of *grouping*, we have tried to draw a picture of the basic control problem as one of *producing* very local events as a function of very non-local or global ones, bearing on the interface as a whole (or at least a sizeable part of it), and to be achieved through a process of *degrouing*. Beyond the evident difference in the direction of 'effect' between perceptual and control tasks, we therefore discern a fundamental similarity in the supposed overall 'character' of perceptual and control processing brought about by the problem of functionally relating two types of descriptions or events: the central or highest and the interfacial or lowest.

But what makes one 'high' and the other one 'low'? One of the more intuitively appealing reasons for this distinction is that 'higher,' more global, or more central events are functionally subsumed under the 'lower,' or more local events, placing the former in a functionally dependent relationship with the latter: while in the 'motor' case lower events are *controlled* by the higher ones, in the perceptual case they are '*included*' within them. Another reason might be that, in general, interfacial events are simple whereas central ones are more complex: thus for example the event representing the symbol "house" is believed to be more involved, and judged 'higher,' than the event corresponding to the activation of a single, 'lower,' constitutive pixel. Closely related to such semantic scaling is the belief that to a *single* central event there corresponds a *multiplicity* of interfacial events that are critically related to it only as a *whole* and not *individually*: i.e. their 'contribution' or 'identity' derives from their interrelationships, which are, ultimately, defined or imposed by central events. (cf. chapter 1; Köhler 1947) The fact that interfacial events can be *related* to one another, and that these relationships appear to be more important than the *single* interfacial events themselves, naturally leads to the view that *any* set of peripheral events meeting the requirements from above can constitute what might be called a 'satisfactory realization' (while one would not in general suppose different central events to lead to identical realizations): accordingly, not only a particular set, but a set of sets, of peripheral events is possibly subsumed under a single central event. For this reason also central events might be judged more 'important' than interfacial ones. This is formally known in the literature as the phenomenon of *equivalence* (e.g. Lashley 1951), and it lies at the very heart of interface-related or periphery-bound knowledge processes.

Once this general framework is established and understood, it is possible to start seeing more

precisely some of its implications with respect to the construction of a particular working system. What we have so far indicated in terms of functional requirements for control could be classified under two general headings: [1] the specification of an object, or a goal-state, for the interface as a whole, corresponding in our example to "house" and the like, and [2] the means for achieving it, that is, functional ways of proceeding from this global description to a specific set of corresponding peripheral events; however, so far these have been only vaguely and indirectly hinted at in the idea of a satisfactory realization.

One should be careful here not to confuse the interface as an object of control, which is by default *always* specified, with the particular object, or more precisely the particular objective, that is intended by the system, and which is *not* specified by default for that unique object. Of course our experience of control in general is such that we do specify many different objects, abstract and concrete, for each of which we of course also specify many different objectives; yet, at bottom, at least as concerns the movement domain, there is always but a single object, namely the motor body, that just happens to be used quite adeptly as a very effective *means* of expanding and transcending the immediate interfacial bounds. These higher levels of control however are beyond our immediate concern, although we might point out that we believe that the more modest course we have chosen can eventually be indicative of, and certainly remains consistent with, an equally rigorous study of their mechanics. This is mainly because these higher levels of control are believed to be built, and therefore dependent, upon the lower-level ones.

The fundamental problem, then, is to functionally connect, or associate, specified goal states with actual peripheral states, or modifications thereof. What we need is a way to specify a variety of goal states (just as the retina allows for the specification of a variety of initial states) supported by a *dégrouper* mechanism that will effectively carry them out. Now as we saw earlier a *single* central or initial event, namely a specified goal state, will 'provoke' a *multiplicity* of interfacial events: hence the degrouper or effecting mechanism (or "exécuteur" as it was called in Lamontagne & Beausoleil (1984)) will have to somehow transform or generate from the given goal a satisfactory set of much simpler events. But which ones exactly? What is a satisfactory realization? How could the mechanism *know* or *decide* what it should be?

Let us consider some extreme solutions. Is it possible, in the first instance, that an altogether different piece of this mechanism is involved in the realization of every single instance of goal state? If this were the case, then in our drawing programme the execution of every "house" command would be computationally *independent* from every other since the programme would be made up of as many different pieces of code as there are potential commands, one for each in fact. (The complete command set, and hence the programme, would be *very* large, but not infinitely so.) Clearly this is not what is wanted, although such a programme, taken as a whole, could be argued

to be a degrouping device. In the second instance, and at the other extreme, could the degrouping mechanism be reduced to a single simple piece of code, something like a read-map-read loop that would accept as commands members from the very large sets of individual pixel coordinates making up each particular house, cat, or face? Again, this is presumably not what is sought.

These extreme cases bring out an important point which could be stated with reference to the multiplicity of interfacial states that is supposed to be generated by the controller or degrouping machinery: in the first case the weight of this variability is transferred onto the programme while in the second case it is transferred onto the command set. (From a formal point of view of course the command sets in both cases are of identical size, and so very well comparable as regards this aspect of variability; but the point here is that if one looks at individual commands, one kind is much simpler than the other; given a certain degree of variation which is to be 'circumscribed' by a programme, a preference for simplicity in commands is thus 'paid for' by an increase in the complexity of the controlling mechanism, and conversely a preference for simplicity of execution is 'paid for' by an increase in the complexity of commands.) However one should not be led to believe that "the weight of variability" or the "degree of variation" associated either with the commands or the periphery can somehow *independently* be brought down to their respective minima. In fact the "degree of variation," irrespective of whether it is defined with respect to the interfacial states or the diversity of commands, constitutes a sort of fixed absolute value of the problem of periphery-bound processes. Reducing that value merely amounts to reducing the *size* of the problem; it does not affect the *nature* of the problem.

Why are these extreme cases uninteresting as potential candidates to do the job? Why would one *not* seriously set out to develop such programmes? Naturally enough, the answer is that much better can be (and is) achieved. But what is a "much better" programme? It is one where an *optimal balance* is achieved between the complexity of the commands and the complexity of the operations that are set in motion by them. Generally speaking, for a given task one will attempt to keep *both* the command set and the programme as simple (or the least complex) as possible. Thus, for instance, the 'weight' of production of a subset of interfacial variability (such as that associated with "houses") should be spread between the programme and the command set. How this should be done constitutes the core problem of control as we have defined it: *it is the problem of degrouping*.

Let us now define it more precisely. The periphery is made up of a *set* of components (e.g. pixels) to each of which is associated a group of characterization dimensions (e.g. position and intensity, which each component can independently assume) describing their *state*: this in turn gives rise to a new set, called the *variation* or *diversity*, whose elements now correspond to a particular peripheral *configuration*.

In the present case the interface is defined as a screen consisting of an array of components,

namely the individual pixels. Each of these components' states can individually be described according to a position and an intensity. An interfacial configuration is the set of all the components' current states, while the complete set of configurations is called the variation or diversity.

Now the interface is the object of control, and the purpose of control can be stated in terms of the realization (selection and production) of one configuration among the set of possible ones. (This however is a static view of control, where each configuration is for simplicity taken as self-contained; the context of movement control extends into the further problem of defining 'paths' among the set of interfacial states. To this we shall return in sub-section 2.2.) Now a particular configuration (or eventually a particular modification of configuration) can be specified and achieved in many ways, the extreme cases, as we tried to show above, leading to a clustering of variation either within the command set or within the programme itself. This leads to an unsatisfactory state of affairs, for while on the one hand the programme becomes impossibly large, on the other hand the commands become unwieldy. The tradeoff is evident here: as the programme grows in complexity, the commands can be simpler, and conversely as the commands become more intricate the corresponding programme can become much simpler, in fact trivial.

But variation is inescapable; large parts of it might be ignored of course, but any non-trivial domain of variation, such as "house," could be argued to involve an 'infinite' diversity. If it were only a matter of *describing* the diversity of "house," then there would be no problem. But the situation is more complicated than that: we want a system or programme that will *produce* instances of "house," and not simply a formal description of them, although that is also requisite. In fact, as von Neumann (1951) pointed out, the only possible formal description of concepts like "house" *might be* in terms of a programme or effective decision procedure (he was referring more explicitly to nerve-networks, but the idea is the same). This is what we have done here: the *command* "house" has gradually shifted into the *concept* of "house" and become related much more closely to the programme than to any other aspect of the situation: now neither the trivial nor the complex programmes could seriously be argued to embody the concept of "house;" at best, they could be argued to allow for it. But if they do not embody it, on what basis could they be said to *know* about it? From this vantage point also one can see how both of the extreme programmes fail, for in neither is embodiment achieved in a sense that would be judged satisfactory: in the first case because the concept is reduced to an 'infinite' number of specific and different instances, which clashes with our intuition of the uniqueness of the concept, and in the second case because it presupposes completely what is to be realized, which clashes with our intuition of the fuzziness (i.e. the potential precision) of concepts: something essential seems to have been missed.

But what exactly is missing? We seem to be faced here with the problem of relating different instantiations of something *unique*, where the difference is related, ultimately, to the interface,

while the sameness seems to rest upon an essential feature of the generating mechanism, i.e. its capacity to somehow maintain or conserve the *identity* of these instances while allowing for a large (but not too large) measure of variation within them. This means, from a functional point of view, that however that essential feature is implemented within a system, it must be done in a way that leaves room for variation. This is why our two versions of the programme are not satisfactory, for neither allows for 'looseness' or for identity: they both display a very tight relationship between the command set and the set of resulting interfacial states. Now at this point it could be objected that a loose relationship, which seems to be what we are pointing toward, is sure to be non-functional because for example the command "house" does not contain sufficient information to account for the generation of the very detailed set of interfacial events that are supposed to be 'drawn' from it. This is correct, except that the command itself should not *account* for the house drawing, it should merely *initiate* it: what does account for it is the programme; it is supposed to know enough to successfully produce the drawing, it must understand.

But aren't we going in circles? If our stated objective is to explain the embodiment of effective knowledge, such as that needed by a programme to produce instances of "house," how can we also state that in order for it to know about "houses" it must know about "houses"?! The answer is quite simple in fact: knowing about "houses" involves something *unique*, which we call the concept of house, and of which the programme itself is the embodiment; but knowing about "houses" also involves much more: beyond the essential knowledge of how doors, windows, walls, etc., must be critically related to form an instance of "house," the programme must know something of the geometry of windows for example, which again implies that it knows something about lines, and eventually the constitutive points ... Note how, as we move down closer and closer to the interface, to the level of final local events, we are also moving away from the *critical* or essential features of "houses," from the identity of "houses," and how furthermore there is less and less room open for variation: while at the top level of the concept of "house," a (relatively) large measure of variation in execution is potentially allowed, by the time the lowest level of individual pixels or points is reached not only is identity lost completely, but there is no room left for variation.

An interesting pattern is starting to emerge here, one that should be reminiscent of the general discussion about knowledge processes carried out in the context of the computational formalism in section 0.1: relying heavily on an example of programming, we have arrived at a recursive formulation of the computational requirements for the task of control or production, which simply states that 'higher' knowledge depends upon 'lower' knowledge. This might probably not be considered a very surprising conclusion, but its statement in light of the formalism not only makes the principle clearer, it also suggests a way to proceed. It is starting to look as if knowledge in general is not, and cannot, be a one-shot affair, as our extreme programmes would indicate, but that it has some

recursive depth to it.

Now the pattern that seems to be followed as one moves into those depths can be described along a few essential features, two of which have in fact just been suggested.

The first feature is the apparent *tendency toward a loss of identity* resultant upon the transformation of the specified goal-state into an actual one as defined by the set of entities making up the interface: as one moved through the steps of realization, the object of control would progressively change from a representation of the whole interface down to a representation of its constituent elements. Thus, starting at the top level, where identity concerns the interface as a whole, the path toward complete realization would lead through successive reductions of the object of control right down to the final level of the constituent interfacial elements, whose identity conserves but a very tenuous relationship with that of the specified goal-state. However this tendency toward a loss of identity of the successive objects of control must be counterbalanced by a tendency to conserve it over the increasing *set* of objects, since if control is to be achieved, the degrouping mechanism must somehow ensure, through this necessary process of object reduction and multiplication, that *overall* the goal is still satisfied. *This process of transformation, which guarantees the conservation of global identity while having to localize or disperse it for realization, is the essential constraint which degrouping must satisfy.* There is no clash between the overall goal and the more local goals it gives rise to since they are derived from it. Consequently, this tendency to conserve global identity concerns the type of relationship that must obtain in the transition from different objects of control, since this is where it has to be assured.

The second tendency could be called *resolving variation* until it is exhausted. We had previously defined variation with respect to the interface only, namely with respect to the set of its possible configurations; now in parallel to this one there exists a higher-level variation which can be described as the variation inherent to the diversity of possible commands. However this nearly goes without saying. The tendency toward resolution of variation which we are talking about here does *not* pertain to the set of commands itself but to the relationship that exists between specific commands and the degrouping mechanism. This is directly related to the character of 'looseness' of commands to which we referred above. It should carefully be noted here that we do not mean that commands can only be imprecise, or of a general character, but rather (given the fact, which we noted at the beginning of this section, that, at least at the peripheral level, realization can occur in a variety of satisfactory ways) that the kind of relationship obtaining between any given command and the degrouping mechanism cannot in general be one-to-one; rather it must be of the type one-to-many. Consequently, if realization is to be achieved, the degrouping mechanism must take on the task of deciding *which one* among the many is to be selected, a feature of degrouping we are calling *resolving variation*. This opens the door to the selection of default values, in cases where there is no

other concurrent constraint acting upon selection, as well as to the possibility of computing the specific values using some intrinsic constraint arising during realization. We are supposing of course that the mechanism is built in such a way that exhaustion is reached, i.e. that all non-pecified (and in a sense non-critical) aspects of realization are computed at or before the terminal peripheral level.

For instance suppose one intends to grasp an object that is within reach: the intention itself, or the initial command, is formulated very precisely at its own level, but it presumably does not 'contain' the specifications pertaining to the extent to which *each part* of the arm is to move, in what direction, at what speed, etc. - this is left up to the controlling mechanism, which can achieve the desired action in a variety of ways depending for example on the initial position of the arm, on the exact distance it must move, on the shape of the object, etc. Again, and irrespective of this potential difference in the details of execution, one thing does remain unaltered throughout, from the first intention to the final realization: the overall goal of reaching the object. So, with respect to the controlling mechanism, at the level of the initial command there is potential variation, and realization consists of resolving it completely. Thus the relationship that exists between a command and the degrouping mechanism is one in which the precision of the command is *not* directly linked to the precision achieved in realization: thus a command is 'loose' essentially because of the variety of possible, and by definition 'particularized,' peripheral instantiations associated with it.

However this is only part of the story. We have so far only dealt with an indefinite command, that is, our attention has been focused on the fate that is reserved by the degrouping mechanism to *one* goal specification. In actual fact, as noted rapidly above, degrouping equally implies a diversity of commands. Since we are not supposing, as the first extreme programme we considered suggested, that there exists a diversity of mechanisms corresponding to this diversity of commands, but rather are taking the view that there is but a single one (or at least one with a far fewer number of computational components than the number of potential commands), something more must now be said concerning the nature of degrouping given the diversity of commands.

We have, in the above paragraphs, dwelt on the subject of the variety inherent in any command with respect to its execution. We have also used an example, that of reaching, which could have been misleading in that respect in the sense that "reaching for ..." is much like the previous "house" command, that is, it defines a *class* of commands, and hence a variety *within* commands which in fact is different than that which we meant when speaking of the relation it holds with degrouping. Although we did not insist upon it, it should be clear that when we spoke of the precision within commands what was meant was the specification within the command itself of the details of the particular member of the class to which it belongs. Thus there is a kind of match between the details that are left up to degrouping *once* the command is given and the details that *must* be specified

within the command itself, making it exact, even though with respect to the degrouping mechanism this precise command still holds much that is indefinite. So, returning to our initial example, if one wanted to draw, say, a *victorian* house, one would need to specify not only its size, the angle of viewing, the number of windows, etc., but above all its 'victorianness'; anything touching upon the general structure of houses, or of windows, etc.; things that are common to *all* houses need not be specified, since they can be 'deduced' by the programme.

We seem to apparently have strayed away somewhat from our initial question about the diversity of commands, but in fact this is not the case. What we have been discussing in the last few paragraphs is crucial not only to the possibility of a variety of command classes but critically concerns the nature of degrouping as well. Following the above suggestion, the possibility of a variety of commands depends upon the 'support' provided by degrouping. By this is simply meant that some aspects of goals need *not* be explicitly specified in the command but can be taken care of automatically, or 'supported' by the mechanisms. So for example in the second extreme drawing programme no support at all is provided by the control mechanism for the diversity of commands: every detail of final goal state must be specified in the command itself; while in the first drawing programme there is no intrinsic means of organizing commands into classes, that is, it does not provide criteria for identification. (Of course we know that the set of commands constitutes a class of "house" commands, but the point here is that it is not defined in a way, supposing it for the moment to be incomplete, that would lead for example to the generation of *new* elements, or indeed of new classes: this presumably requires that the set of commands available be such that it is possible to 'combine' them into new commands, a feature which can be seen to depend critically on the extent to which the programme, i.e. degrouping, can support this set of more primitive commands; but this is clearly not the case there.)

What we are suggesting here is a sort of language for commands, one that would allow the specification not only of a diversity of commands but of a diversity of *classes* of commands. The degrouping mechanism then would play a rôle similar to that of a compiler, accepting high-level formulations and transforming them into series of low-level operations. One is easily reminded here of our earlier but much more general discussion of universal Turing machines, where the defining set of primitive operations (a list of input symbols associated with changes of state) was such that *any* procedure could be described in it; the principle involved is the same one. Everything we advanced concerning the nature of degrouping and the relationship it holds with the command set could be rephrased into the 'language' of programming. Little would be gained from this however. But, following this supposed comparison, and bringing it back to the domain of control and degrouping, we do wish to remark upon the fact that beyond allowing for and supporting the diversification of commands, the language would also allow for a high degree of precision within its

own realm, that is, at the global level of goal specification; thus execution neither creates nor does it impede command precision; rather it conserves and transforms it in the process of degrouping.

A final caveat concerning the use of the computer and programming metaphors, that is widespread in the literature on movement control, is also worth mentioning at this point. While we believe that the principle involved is identical, we do not believe that, for example, a model of movement control requires a language as general or as abstract as that of a universal Turing machine; the generality that is sought does not concern the computable *per se* but rather a particular instantiation or subset of the computable that concerns and applies to the domain of movement control or production. It so happens that the idea of degrouping is similar and related to the idea of programming, but that is to be understood more as an outgrowth of the adoption of the computational formalism as general conceptual framework than as a statement about the sufficiency of the programming metaphor as a model for the task of control, and in particular movement control. This is simply because, taken as a model, the programming metaphor is as general as the computational metaphor, and so, as concerns the problem of movement control, nothing can be gained by exchanging one more concrete instantiation for the other more abstract one. (Recall, from chapter 0, that computers are implementations of universal Turing machines.)

It is the leap from metaphor to working model that is lacking, as indeed the formalism calls for and allows. Our whole effort can be viewed as an attempt in that direction. We have tried to define the basic problem of control as involving a certain kind of knowledge process (degrouping), focusing on the task of using a description, namely the command, bearing on the object of control as a whole, to produce a corresponding configuration (set of interfacial component states). Each domain, the command and the interfacial one, attaches to intrinsic variation that must be reconciled in the productive system. While interfacial variation is easily defined, the axes of higher level command variation (as supported by the degrouping machinery), or its *characterization dimensions*, are more difficult to pinpoint. For instance, following chapter 1, we would need to define the degrouping mechanism embodied in the drawing programme in such a way that it implements the *primary* characterization of "house" commands, as well as some *secondary* characterizations such as its style (e.g. victorian) or size, etc. By adopting a recursive approach to this problem, that is, as stated above, by establishing a primitive level of control upon which others can hopefully be built, we believe it should become more tractable in that one level of control, once entrenched, would in principle guide the elaboration of the next one.

This completes our initial attempt at rephrasing the problem of control, and, though still very indirectly and partially, the problem of movement control. On the one hand, the central elements of this definition concern the interface, the multiplicity, homogeneity, and independence of its components, along with the locality, borne out of their primitivity, of the events to which each

constituent entity gives rise, and finally the associated diversity of goal configurations, borne out of their combinatorics. On the other hand, they concern the features of the controlling or degrouping mechanism with respect to the commands: the fact that the commands bear upon, i.e. identify, configurations, or modifications thereof, and the fact that in view of their singleness, relative to the multiplicity of interfacial events triggered by each one of them, they are global in nature, and finally the fact that the mechanism allows for and ensures this by transforming a specified goal into a satisfactory realization, i.e. one effectively corresponding to the stated goal.

This idea of a satisfactory realization can now be more fully appreciated. By qualifying degrouping with the requirement of 'satisfactoriness' we are simply stressing the fact that it is not in itself sufficient to define the problem as one of functionally relating two different domains, the local and the global, or of associating the single with the multiple, but that a particular *kind* of association is called for, one which we expressed intuitively as conserving the identity of the goal throughout a process in which it is gradually 'lost.' As the levels of control move closer and closer to the interface, the object of control presumably shrinks while getting 'subdivided' into the interfacial multiplicity of entities. It is at this point that the idea of a satisfactory realization enters the picture: the system as a whole must be built in such a way that the simultaneous splitting up of the object of control, along with their respective management, occurs in such a way that *overall* the goal's identity is respected or conserved. This is a direct consequence of the fact that goals concern the interface as a whole while levels of control must become more and more local.

With respect to the constructivist or recursive approach to model building just discussed, it should now be clear where the major problems lie: [1] one wants to build up a (global) level of control that will eventually concern the whole interface by first setting up a most primitive level where, by definition, it is not yet the case, but upon which will be built the next one, etc., until is achieved control of the complete interface; [2] however this functional connection of the global with the local levels of control should also be defined so that the system has the largest possible (and eventually complete) 'hold' upon variation. Since variation is most clearly described with respect to the interface, the development of a control system aiming toward the production of 'total' variation has every reason to *start* with the peripheral level of control. In section 2.2.0.1 this issue will be developed more fully.

Consider finally that, since in principle each new level of control has to deal with more and more elaborate objects defined by the previous level of control, building up a new level of control on top of an existing one not only means dealing with a group of its objects (each of which presumably encompasses a group of objects of a *lower* level of control, etc.), it also means dealing collectively with the way these objects are dealt with; literally, each new level will be called upon to control the lower-level controller/s. In this way each new level can open onto ever expanding

realms of more and more powerful possibilities.

We believe that these elements can already provide a framework within which can be elaborated computational models of movement control. However before we proceed with this task we will first attempt to place the present understanding within the larger context of past research efforts.

2.1 Approaching movement control

The central feature of our analysis of the problem of movement control is that it is based upon an interpretation of control as a problem of *production*. The computational approach to the problem of movement control therefore requires the specification of an intended level of performance along with the details of the functional system that is supposed to achieve it.

Now we would argue that for the problem at hand a rigorous or precise specification of the level of performance is not needed before one can proceed. There are two main reasons for this: [1] a computational argument: if, for example, one intended to develop the ultimate human robot (i.e. aimed for a human level of movement performance) then we hold, following von Neumann (1951), that the task of formally describing the full range of human movement is contingent upon and equivalent to the creation of a system whose performance reaches that level; it would therefore not precede but follow that enterprise; in fact it would only be accomplished as the end product of the whole endeavour; [2] an analytical argument: there is a deep and pervasive conviction upon the nature of most complex biological and psychological phenomena: that they are built from and dependent upon simpler ones of a similar nature; the same of course holds for studies of movement control, so that for example a human level of performance is thought to rest upon lesser or more primitive levels of control, etc. Now we have called this last argument analytical, but within a productive approach to movement it should be clear that the point of departure for analysis is not the most complex level, although from the point of view of a working system that highest level (now of prescription or command of course) would presumably be found at the beginning of any one realization. The point of departure for analysis in fact is the lowest level of control possible, which should be such as to allow for the implementation of successively more complex or elaborate control structures. It follows that the nature of this most primitive level of control, which we might call *interfacial*, is critical with respect to the possibility of achieving the final goal. In a word even though it should be constraining, the interfacial level of control should not be restrictive: it should be constraining in that it indicates, directs, or opens up onto further developments of the control system; but it should not be restrictive in the sense of blocking off or disallowing certain possibilities or types of known movements. For example the control system for a hand should permit movement of the individual fingers, or that of an arm allow for independent movement of its components; movement of the effector (hand or arm) as a whole - a more complex control task -

presumably depends (but not necessarily directly) upon such a capacity and so either a control procedure for the effector as a whole which does not allow the locus to be brought down or to be shifted to the level of its components, or a control procedure for the components which cannot be brought up or included within the control mechanism for the whole effector would both be unacceptable.

However even though an explicit or formal description of intended level of performance is not needed, one must still be able to show or argue that it could *potentially* be reached by a given system. This is the critical rôle of the definition of the motor interface. Following the same reasoning as above, it is clear that interfacial restrictions that are too great will defeat the purpose even before work is started on the control system. Thus, one would define an effector system or interface that allowed for a 'good' description of any *actual* movement; in that way one would be satisfied that any such realization would in principle never lie beyond the reach of an eventual control system tied to that interface. This might be called the argument for qualitative interfacial completeness. Of course in most contemporary research settings the *sizes* of the interfaces used fall quite short of their actual human counterparts: in vision for example one does not work with a retina of one hundred million receptors, or in movement studies one does not (usually) directly deal with a 23 degrees-of-freedom arm. Yet the argument for completeness still holds because beyond a certain critical size the difference is quantitative, not qualitative. For example enlarging the size of a retina beyond what is needed to discriminate 1° differences is of little consequence, even for shape analysis, while in movement production a minimum of 6 degrees of freedom are needed in order for an arm to fully access 3-D space. Geometrical arguments of the sort adduced in the previous chapter about interfacial representational capacity needed for a certain level of visual discrimination of oriented line segments readily provide a measure of these critical thresholds.

The above view on the problem of movement control, that is, that it consists mainly of a problem of production and that intended level of performance is defined implicitly with respect to the interface, leads to some important consequences for its study.

It implies that the basic task of a movement control system is that of specifying and realizing, out of a very large pool of *potential* movements defined through the interface's properties and capacities, an *actual* one. Now as argued in the previous section since these products are defined with reference to the periphery through which they materialize, or more accurately the last and hence most primitive level of control, it follows that in order to specify the desired entities or sets of events, the overall process of control must embody, or call upon, some form of knowledge. From a formal-descriptive point of view, it is certainly true that the problem of control is solved *in principle* once it is shown that muscular action is necessary and sufficient for movement to occur; however from a functional point of view this is only the beginning. A (formalized) solution to the

problem of *production* further requires a group of functions that map (or select) from the *set* of potential movements to an actual one. It is because of this property and because these functions actually encompass within themselves the whole domain of movement that they are said to constitute an instance of knowledge. A major effort of research must therefore bear upon the nature and mechanism of such knowledge.

Again, as argued in the previous section, for interface- or periphery-bound knowledge processes, and therefore in particular for movement, there seems to be a critical interplay between what is known on the one hand and how it is used or functionally embodied on the other hand. In our view this goes some way in explaining how the psychology of movement control and the neurophysiology of movement production are related. Now as we shall soon see in some detail the same quest for the nature of the knowledge involved in control or command can be seen to animate both approaches to movement production. However in that common search one set of efforts can be seen to more clearly incline toward understanding the system, while the other favours studies of performance. We shall call the latter type of research PERFORMANCE-DRIVEN and the former type SYSTEM-DRIVEN since not all efforts can be clearly identified as belonging to either Psychology or Neurophysiology. In that respect, Artificial Intelligence (AI), and in particular Robotics, provides a third area of research of potential interest to the problem of movement control. As opposed to both performance- and system-driven research, these endeavours are much more explicitly concerned with model-building, especially of the computational kind. We have labelled it TECHNOLOGY-DRIVEN research. But though Robotics research may be viewed as occupying a strong intermediate position on the system-performance axis, closer study reveals that it has generated very few hypotheses of interest to the psychological or biological aspects of the problem of control, a feature that seems to result not from a lack of conjectural effort but from possibly placing too large initial restrictions on the intended range or type of performance. This will be discussed at greater length further down.

So far in this chapter we have endeavoured to define a general framework for movement production based upon an epistemological interpretation of the theory of computation developed in the introductory chapter. The main feature of this view is that it appears consistent with both psychological and neurophysiological approaches to the problem of movement control, a main consequence of the fact that from the computational point of view, problems of knowledge naturally unfold into concerns for both psychological (global, behavioural, or performance-related) and neurophysiological (functional, local, or system-related) issues.

It is worthwhile therefore to reexamine this understanding in the wake of past efforts and relate it to some of the main issues that have become central foci of attention in the field of movement production. In so doing our main emphasis will alternate between the search for explicit statements of principles of movement control, as they pertain to a general conception of the problem of con-

trol, and the discussion of more specific issues connected indirectly but critically with these ideas.

2.1.0 Performance-driven research

In a large measure the conceptual and paradigmatic prototypes for the performance-driven approach to movement production were set down by Woodworth (1899). Attentive reading of the original paper reveals that the idea of movement programming (the "initial adjustment"), the centralist-peripheralist issue (arising from a consideration of the "sensory basis of motor control"), the fundamental observations on the speed-precision tradeoff, and finally even the problem of equivalence, can all be traced back to it.

All this is set within the background of an overall conception of movement centered upon performance. Now the aspect of performance that he considered essential to the study of movement was its *effect* or *result*, which provided an empirical basis for the measurement of dimensions of performance (or control) such as accuracy. He states (*ibid.*, p. 18) that "the emphasis is laid on the *production* of movement" - as opposed to, or complementary to, its perception, which had been the object of most earlier studies - and that "in [this] connection the movement must always be studied with respect to some result aimed at." This is a crucial pronouncement. He had stated earlier (*ibid.*, p. 4) that "the movement must have a particular direction, a definite extent or goal, a definite force, a definite duration, a definite relation to other movements, contemporaneous, preceding and following." From the point of view of the production of movement these characteristics of performance could be understood as addressing or at least pointing toward the issue of 'system variables,' that is, the problem of the nature of the control (or production) process itself. But there is in fact little ground for believing that this is what was intended or planned in general or in that particular study; rather it seems that direction, duration, extent, etc., were to be associated with dimensions of particular *tasks*, and not taken as indicators of the nature of the movement to be produced, with the view of arriving at a model of production. This also holds for questions surrounding "the sensory basis of motor control," where one would naturally expect production 'events' to play a central role in recognition 'events.' Yet there again the discussion is very often carried out in terms of performance or task variables, presumably because production itself is largely conceived in those same terms.

We are of course aware that it is equally possible to argue that stressing empirical or performance facts acts as a safeguard against erroneous generalizations. However it is more difficult to argue for what is sometimes considered an equivalent proposition: that it is *only* by a strict adherence to the facts that we will come to correct generalizations, that a constant reference to them constitutes the *only* possible means of steering research toward the truth. Be that as it may, the contributions for which Woodworth is remembered in the field of movement control have much

more to do with the ideas he put forth than either the paradigms he developed or the data that he accumulated.

Of central importance in that respect is a distinction Woodworth made between the "initial adjustment" and "current control" of a movement (ibid., p. 41). He defined the second as "reactions to stimuli set up by the movement" (ibid., p. 58), be they of visual or other kinæsthetic origin, but without providing indications, even of a general nature, as to the possible mode of sensori-motor integration which it calls for and upon which it supposedly relies. That may have been because such immense empirical possibilities were open for investigation that even preliminary considerations upon mechanisms seemed premature. However this was apparently not the case for the first 'mechanism,' namely "initial adjustment," which he postulated and characterized in somewhat greater detail.

If lines of considerable length, say a foot or two, are made at a rapid rate, the changes in direction will probably be found to be about the same in them all. [...] The initial impulse takes the hand along a curve. We aim around a corner; not according to geometrical straight lines, but according to the make-up of our arm. From the geometrical point of view, the simplest movement that we can make - the movement as determined solely by its first impulse, and not complicated with later adjustments - is still a complex affair. The initial adjustment is itself complex. It includes the innervation of different muscles one after another. The coordination adapted to produce a straight line is probably more complex than that to produce certain curves. The first impulse includes also a command to stop after a certain distance. These later effects of the first impulse are probably in some degree reflex. [...] Yet the first impulse of a movement contains, in some way, the entire movement. The *intention* certainly applies to the movement as a whole. And the reflex mechanism acts differently according to the difference in intention. We must suppose that the initial adjustment is an adjustment of the movement as a whole. (ibid., pp. 54-55)

This passage is uncharacteristic of the whole paper in that it is one of the few places where explicit reference to the system is made, and it provides, along with those on "current control," the report's most detailed theoretical statement on the process of movement production. It is argued that intentions somehow mesh with "the make-up of our arm," that they involve ordered "innervation of different muscles," and that "reflex mechanisms" must enter into play at some point. While it could equally be argued that such statements are trivially true, since by definition it is the arm that moves, and that the arm is moved through systematic activation of different muscles, into whose workings it is known that various reflex mechanisms are operative, it could much less trivially be shown *how* an initial adjustment can use or take them all into account 'simultaneously,' that is, how it could "contain the entire movement," or comprise "an adjustment of the movement as a whole." In that respect the implicit distinction between control events of a peripheral and of a central nature is worth noting, based apparently on the fact that even though every arm *movement*

is a function of arm geometry, of specific sequences of muscle activation, and of reflex adjustments, the process of movement *production* must involve more, a sort of generalized capacity for correctly selecting or specifying, for *every* intention, *all* of the interfacial or control elements required for its realization. This interpretation is reinforced by a later statement on the basis for the judgement of the extent of voluntary movement, which must be "a sense which is not reducible to a sense either of its force or of its duration, or of its initial and terminal positions," a sense, furthermore, that is "the real basis on which the control is based." (ibid., p. 80) The implication here seems to be that "the real basis" for control has no directly specifiable empirical counterpart, that is, it does not appear to correspond to any single 'familiar' physical dimension of movement, i.e. of the final outcome of the process. But then where is it to be found?

Even though Woodworth had stated that "the question raised [here] is not as to the possibility of any movement at all" (ibid., p. 3), where we interpret "possibility" as referring to production, and that his purpose therefore differed from our own, it seems remarkable that he should nevertheless have allowed himself to speculate on the nature of movement production, and that such theoretical sidelights turned out to have provided a very important conceptual framework for nearly the whole of subsequent research efforts. Indeed the idea of "initial adjustment," along with the notion that it should encompass the whole movement, etc., were among the earliest formulations of the idea of motor programming, an idea which has pervaded the field since then, both in Psychology and Neurophysiology; and his initial discussion of the problem of "the sensory basis of motor control" continues to ring a familiar note. However since he had declared his work to be a study in "the motor side of consciousness," (ibid., p. 1) the appearance of such theoretical concerns should be less surprising, if only in view of the desire that its conclusions should have some bearing upon the conscious - as opposed to strictly behavioural - correlates of movement production; as it turns out however, given a basic allegiance to the then growing empirical tradition in Psychology, it seems that to successfully 'pull' (behavioural) findings into the domain of consciousness is an undertaking requiring first an answer to what Hebb (1949) called the organization of behaviour. In the specific case of movement control this requirement translates into a need for a model of production. Now Woodworth had recognized that "the *intention* [a conscious correlate] certainly applies to the movement as a whole," (loc. cit., p. 55) but unfortunately in the subsequent discussion very little room is devoted to what seems to us to be the basic problem of the relationship between the wholistic character of the intention on the one hand and the systematic relations that can be established within the set of physical characterizations of, and the conditions associated with, the resulting behaviour on the other hand.

But this question does not seem to have been allotted the central position which we assigned to it, for even though one could have expected subsequent research to pick up on the basic theoretical

issue that Woodworth, however implicitly, had raised, it is not easy to follow its thread in what otherwise grew to become a relatively rich area of study. Having recognized so early on that, for example, the "sensory basis" of control seemed to require the *integration* of many aspects of movement, and having advanced furthermore a similar idea concerning the "real basis" of control, should have prompted the search for *possible ways* in which this could be achieved or how it could be the case, just as, from the general point of view, Woodworth had already endeavoured to conceptualize movement control or production as a two-stage business, resulting in a 'model' that provided a starting point from which could then be empirically debated and tested questions and possible answers about the accuracy of movement. But the initial emphasis on performance and the related eagerness to maintain strong empirical claims and roots, seem to have resulted in a drawing away from such large theoretical issues, and a shying away from conceptual efforts that are supposedly in danger of losing sight of, or straying too far from, performance data. Indeed many years later, and having recognized Woodworth's seminal conceptual contribution to the field of movement studies, Fitts (1964) stated that "one of the most important advances made [...] is clarification of what is meant by a program or plan governing a sequence of operations." (p. 261) The "adjustment" metaphor had then been replaced by the more elaborate, but still closely related, (computer) "programming" metaphor, an advance that unfortunately did not seem to carry over into the more specific question of the nature of a *motor* programme. And still later, Kelso & Stelmach (1976), commenting upon "the so-called switch in emphasis in motor skill research from a "product-" to a "process"-oriented approach wherein the underlying processes and mechanisms are of primary concern," concluded that "The major problem, however, is that psychologists are still measuring product most of the time, rather than process." (p. 35)

The performance approach to movement production which Woodworth initiated did not however come to methodological fruition until many decades later, although his ideas played some instigating rôle in the growing mass of physiological, and eventually neurophysiological, studies of human performance. The most often cited reference in that respect is by Lashley (1917) who, with some measure of success, put some of these ideas to the 'neurological' test. Although he observed and deplored the "almost total neglect [by psychologists] of the neurological problems presented by skilled movements," (p. 169) he had little else to offer in the way of "the nervous mechanism of initiation, continuation and cessation of adaptive movements" (ibid.) than criticism of earlier 'implausible' ideas, especially those pertaining to the cessation of movement. But most of these ideas, it should be noted, were derived from the psychological domain, Physiology not yet having attempted to model motor control functions beyond the level of reflex integration. Yet as we shall see later on, Lashley himself is responsible for a very clear formulation of the essential problem of control, in which neurophysiological answers were sought to questions arising from the analysis of

performance data.

However, greater attention was eventually paid by psychologists to the physiological correlates of "adapted" movement (e.g. Dodge & Bott 1927; Wilson 1933). Not surprisingly these initial studies were carried out in the wake of Sherrington's no less pioneering work (Sherrington 1906), in particular in the area of the co-contraction of agonist-antagonist pairs in reflex and non-reflex movements. But by concerning themselves with such matters, they effectively started shifting the emphasis toward the system and its working, although one could equally argue that such attempts to reach physiological rigour and precision were rooted in and inspired by earlier behaviourist theories regarding the critical rôle of efference in consciousness (cf. Max 1934). There appeared suddenly a new pathway for the characterization of performance, one derived not from the *effect* of a movement but from an *intrinsic* part of it. Again this inevitably led to the problem of linking "certain temporal relations which characterize muscular contraction and relaxation" (Dodge & Bott 1927, p. 245) with "voluntary regulations of behavior." (ibid., p. 263) Yet in trying to organize their results they revert back to Woodworth and propose that "it seems reasonable to suppose, if a complicated reaction is to be effected evenly and continuously, that there must be no real hiatus between the preliminary adjustments and their subsequent correction." (ibid., p. 263) This new index of performance had apparently not simplified matters, but only shown them to be at least as complicated from the physiological as the behavioural point of view. The search for the physiological key to performance however did not stop there. We shall see shortly that it was pushed up higher within the system, toward even more intimate neurophysiological indices.

It would be wrong to think that these studies, and a handful of similar ones, were heralding a major reorientation of the field, for even though they did foreshadow a tendency and an outlook that eventually came to typify (if not dominate) the area, the road from there to the present state of affairs was not a straight one: two sharp conceptual bends, imparted from Cybernetics and Information Theory, were yet to significantly mark its path.

It is well known that the immense efforts expended toward the solution of the numerous and varied technological challenges that issued from the problems and necessities of the Second World War led to a series of conceptual breakthroughs that have exercised a profound and lasting influence in many areas of scientific research since then. In particular, certain problems of control (especially related to A.A. gunnery and piloting), which is usually recognized as providing the original impetus toward the development and use of cybernetic ideas, were of vital importance. A closely related problem concerned the operators of these new machines. Its study seems to have led some psychologists, among which Craik (1947, 1948) figures prominently, to observe some striking similarities between the performance of human operators and the functioning of the basic cybernetic control model. It followed from this view that the senses and perception could be treated as the

input component of the model, while the effector system acted as the output component; the brain, along with the motor system, was associated with the comparator or computing center and the amplifying system, respectively. (cf. Craik 1948) Thus the cybernetic model provided at once an integrating and an analytical framework for viewing and studying man's behaviour.

Even though Craik's most important contribution in that respect as concerns movement control pertained only indirectly to movement production, it did so in a highly suggestive manner. This took the form of a hypothesis about the necessarily intermittent nature of correction. Basically he postulates the existence of commands for (more or less) "complicated patterns of movement" (Craik 1947, p. 58) that are "individually triggered off ballistically, and the sensory feed-back must take the form of a delayed modification of the amplitude of subsequent movements." (ibid.) Those commands, he further suggests, must be such as to allow for the transformation of the control situation into one where the task consists of detecting and/or maintaining the *constants* characterizing that situation. (ibid., p. 59) Thus in a pursuit task the operator will detect a constant speed (if that is the case) and issue a "constant speed" command, while, again if speed changes, and in a constant manner, he may in turn adjust his movement with respect to that new invariant. As we shall soon see this idea of invariants resurfaced three decades later, although under a different guise.

It is of course with the (more or less) "complicated patterns of movement" that we are mainly concerned. With respect to these, Craik suggests that they may have to be learned (Craik 1948, p. 147). Just as for a visual pursuit task involving a complex stimulus shape it is very difficult to conceive of a system that would exclusively rely on local speeds (cf. Lamontagne & Howe 1980), so for a motor pursuit task involving the whole arm it is not at all clear how a system could depend or work on speeds associated with single joints (i.e. local speeds). Presumably if a motor system is to rely on speed information per se then a way must be found of assigning a 'global' speed to the whole arm. This single value could then be used in tracking, or for computing accelerations, etc. That *may* be a product of learning (see further down), but it is certainly not by simply suggesting that it is learned that the problem is solved. The typical complex movements that are used as examples rarely refer to this intrinsic complexity of even a relatively simple effector system. However, to these problems we shall return in section 2.2.

Most of the research on movement production and performance carried out during the first few decades following the publication of Woodworth's ideas and results was pursued in the spirit that he had infused it with. In the main (cf. Keele (1968) for a review of some of this material), the same basic performance variables (speed, duration, accuracy, distance, direction, effect of the presence or absence of feedback, etc.) continued to be studied in the context of very similar tasks (reproduction, positioning, tapping, etc.). These studies however resulted in little addition to, or modification of, Woodworth's original formulation; his model retained its dominating position as overall framework

within which were made, or rather were set, most attempts to identify lawful relationships.

Remarkably, a sizeable portion of that research was eventually brought within the scope of Fitts' information-theoretic equation (Keele 1968). Fitts' Law, as it came to be known (Fitts 1954), constitutes the clearest and most synthetic statement involving performance variables: it relates movement time (MT) with precision or tolerance (W) and movement amplitude or displacement (A) in the following manner:

$$MT = a + b \cdot \log_2(2A/W)$$

where a and b are constants. The famous speed-precision tradeoff, observed and documented since the time of Woodworth, had at last achieved a formal expression. Yet soon afterward arose the problem of explaining Fitts' Law, that is, of relating it to aspects of the movement *system*. We view in that single question the origin and source of much of the developments in the field in the past two decades and a half. For, strangely enough, even though Fitts' Law was *not* derived from, or based upon, a direct study of movement systems, its greatest contribution might well have been to provoke a major renewal of theorizing about it. Keele himself, in stating that "The first step in understanding movement control is to discover what variables influence speed and accuracy," (Keele 1968, p. 388) was still framing the problem in woodworthian terms of performance, but the gist of the paper in fact consisted in exploring "three possible ways in which movements may be controlled." (ibid., p. 393) Of those three possibilities the first two, namely the rôle of visual and kinæsthetic feedback, could appear to belong to a different category than the third, namely the possibility that "movements may be preprogrammed in the sense that the amount and timing of innervation to the muscles may be determined before the movement begins, and once it has begun, peripheral feedback would exert no further control," (ibid.) for while the former are based on the supposition that feedback plays some rôle in the control of a *given* movement, the latter is rather a suggestion concerning the nature of the *production* process, though admittedly it is not a very precise one. It could well be argued, on the other hand, that this merely reflects a difference in emphasis, and that the ultimate goal, from whatever direction the approach might be made, remains the same in all cases: understanding the production process. However we have not encountered such discussions and arguments. Our main point here is that the problem of control seems to be understood in two different ways, one referring to the identification of factors acting upon the movement command process, the other referring to the command process itself. This dichotomy is probably focal to what became known as the centralist-peripheralist debate (cf. Jones 1978), in which one can be left with the impression that a convincing demonstration of the critical rôle of peripheral input (as a major

control variable) within the command process would in fact settle the problem of the command (control) process. We view the situation in the opposite, that is, that it is only through efforts at settling the command or, equivalently, the production problem that the necessity and importance of peripheral input or 'control' can be decided.

Nevertheless the tide seemed to have turned and opened the way to renewed attempts at expanding the original woodworthian model. This, as we shall see, was marked by an explicit recognition of the cognitive nature of the problem of control or production. Even though the woodworthian variables and general outlook survived, growing attention turned to theory. Historically, motor learning was the first area to benefit from this, but the fundamental problems of control attached to motor learning, and left in a more or less satisfactory state by theory, eventually sprang back into the fore. But before picking up on this trail, and as a prelude to it, we will digress for a moment into a contribution to the larger field of cognitive theory from which we shall draw one of the clearest statements we know of on a particularly important, but so far largely neglected, aspect of movement production, namely equivalence.

Coming from an apparently very different direction, and addressing an apparently wholly different set of issues, Hebb's proposal of the theory of cell-assemblies (Hebb 1949) nevertheless proved to be the source of what we consider to be essential observations on the nature of control in general, and by extension, of movement control in particular. Viewing movement as the constituent and basis of all behaviour, and conceiving behaviour as essentially being 'driven' by environmental conditions, had by then become a long-standing psychological truism; consequently, the main goal of Psychology was to discover the laws of association between stimuli and responses. On that view can be understood the importance of the study of the 'foundation' of behaviour, since movement is thought to be the most primitive level of response, to which presumably could be linked similarly primitive levels of stimulation or perception. The very primitivity of that response level however steered its study into the physiological arena; but by doing so it automatically engaged criticism. Thus the theory of reflex-chaining, purporting to resolve complex behaviours (such as language, and by extension thought) into discrete series of (learned) stimulus-response pairs, and according to which the sensory consequences of a particular response acted as stimulus for the next response, and so on, was not only put under heavy psychological fire but was also found to be neurophysiologically wanting. (cf. Lashley 1951) From the start, the study of performance found a niche in neither of these two disciplines, and although Woodworth had "not thought it worthwhile to attempt to decide on which side of the border [his] studies lie[d]," (Woodworth 1899, p. 27), the persisting 'distance' between them has never been resolved clearly either in terms of methodological or conceptual differences. We believe that Woodworth was right in principle, and, following Popper, that problems, not methods, are at the heart of research endeavours. We furthermore tend to view Hebb's

attempt at a general *neuropsychological* theory in this same light, namely as an attempt to simultaneously satisfy what had otherwise grown to become the entrenched canons of psychological as well as those of neurophysiological theory: by first stipulating the existence of a "process that is not fully controlled by environmental stimulation and yet cooperates closely with that stimulation," (Hebb 1949, p. xvi), and second by requiring that the hypothesized nature of that process be neurophysiologically plausible, he was led to the well-known idea of cell-assemblies, whose consequences were laid out in great detail and with respect to a great many problems of behaviour.

In particular he devoted a few pages to the phenomenon of motor equivalence, which he defined as "a variability of specific muscular responses, with circumstance, in such a way as to produce a single result." (ibid., p. 153) He also observed that "the variability characteristically appears when there is a variation of the animal's initial posture, of his locus with respect to the goal, or of intervening barriers." (ibid., p. 154) His account of the phenomenon however is somewhat surprising in that equivalence is seen as based critically upon earlier and simpler learning, in which for instance is established "an ability to move either hand or foot directly from any point to the line of vision." (ibid., p. 155) This might indeed be the case; and neither would we quibble over the idea that a human being "does *not* learn to make a certain series of muscular contractions," (ibid., p. 157) but learns rather, for example, "that the depression of a certain lever starts the motor on his car." (ibid.) Indeed if "learning" were replaced with "intends," little of the essence would be lost. But this brings up our major and central point: even though one does not intend certain series of muscular contractions but rather to depress a certain lever, it is still true that a certain pattern of muscular activation *must* result, the one in fact that corresponds or amounts, from a global point of view, to starting the motor. Hebb insists, correctly we think, on the fact that there is quite a distance between "starting the motor car" and the corresponding "series of muscular contractions." But note that this depends on the existence of another sort of even more primitive equivalence phenomena, based on the system that translates the one into the other (and eventually many *equivalent* others). This other type of association could presumably be the one he understood to occur earlier, and upon which could then be built up higher levels of equivalence. But approaching this from the point of view of learning is paramount to supposing that such a mechanism already functions, even if only in an embryonic fashion. For example, the child who acquires the general 'concept' of reaching or forms the generalized schema of reaching, would not be said to be acquiring a brand new concept but rather to be refining, or generalizing upon, an already existing approximation, one that depends on the prior functioning of a mechanism that allows him to first 'conceive' and realize (i.e. translate into muscular patterns) something like and amounting to "extension of the whole arm." For if this were not the case then on what basis could the supposed generalization ever occur?

We have not, for our part, adopted a learning approach to movement production, our main

concern being centered upon problems related to this more primitive type of equivalence. But, following Hebb, we believe that the higher levels of equivalence displayed by organisms are in direct functional relation to it and that it must be resolved before a strong solution, whether or not it is afterward based on learning principles, can be found to this "ubiquitous, fundamental psychological problem." (ibid., p. 157)

Even though Hebb had hoped "that such conceptions [would] be a good basis for the study of motor learning," recognized "that the study [was] still needed," and "that nothing more ha[d] been provided than a rational approach which may or may not be found adequate in the long run," we find very little reference to them in the subsequently vastly expanding field of research into motor learning, the context in which nearly all theoretical efforts into performance henceforth occurred. None of them having picked up on Hebb's suggestion however, it is not surprising to find that the problem of the more primitive level of equivalence was not tackled either. We will try to show that the same remarks that were made above are still applicable to these theories. While this lack pervades much of the ideas put forth in the context of performance-driven research, we would argue that this problem of 'primitive' equivalence lies at the very heart of system-driven research. But further discussion of this point belongs to the next sub-section.

While Lashley (1951) had criticized the idea of reflex-chaining, which was then the only available explanation for temporally ordered behaviours such as motor skills, and argued that it was unsatisfactory as a control and as a learning theory, it was only twenty years later that an alternative motor *learning* theory was formulated. Adams' theory (Adams 1971) is expressly cognitive in nature. Departing from the classical thordikean model of the acquisition of response patterns, he proposed a new model in which KR (knowledge of results) provides *error information*, not simply reinforcement. His theory postulates and relies heavily upon an ability to *detect* and *correct* errors. He states that "motor learning is best viewed as a problem to be solved, where S tries a movement, is given KR, tries again on the next trial, and so on, and the KR is the information used to solve the problem." (ibid., p. 122) The basic assumption behind the model is a conception of control whereby "the choice of direction and the extent of movement are the two main properties of a simple, self-paced movement." (ibid., p. 123) To each of these properties he associates a construct, namely a "perceptual trace" with the extent, and a "memory trace" with the selection of a direction of movement, the latter's purpose being "to select and initiate the response, preceding the use of the perceptual trace," (ibid., p. 125) adding furthermore that "in effect, the memory trace is an open-loop motor program," albeit "a modest motor program that only chooses and initiates the response rather than controlling a longer sequence, as advocates of motor programs usually imply." (ibid., p. 126) The main loop of the model depends upon two functions, one of *recall*, which is responsible for selecting and initiating a movement, and the other of *recognition*, responsible for

"knowing whether the movement is proceeding correctly or not," and governed by "the perceptual trace along with ongoing feedback." (ibid.)

As long as learning theory remained empirically oriented, there was never any need to formulate models and hypotheses in terms of internal mechanisms. When, for a variety of reasons, such models were found either too limiting or otherwise unsatisfactory, and attention turned 'inward,' there appeared the need not only to *refer* to internal constructs but also to specify their functioning within specific mechanisms. By removing the problem away from a strictly empirical domain and into the speculative one, there appeared also the extra burden on theory of proceeding from it back into the empirical field where all scientific theories should be laid to the test. Thus Adams' theory of learning is erected upon a basic concept of how control works. His main interest was of course to show how this skeleton model allowed for the superimposition of a learning dynamics: first one imagines how it works and then one shows how this allows for betterment or improvement of performance. The two are very closely linked. Now at the heart of Adams' theory is the notion of error. Within a working system the treatment of error implies firstly its recognition, and secondly its correction. This is the crux of the matter: learning occurs as better and better approximations to the desired result are selected with the help of error information. But what is a better approximation? That is, how is it selected? Within an empirical framework it is sufficient to note that improvements occur, and according to what particular responding profiles they do. But in the cognitive mode this is not enough: one must describe the selection process itself, a process that, although recognized as crucial, could on first view be thought only accessory or of secondary importance in the basic learning dynamics. It is not sufficient because in the absence of a detailed picture of the selection process one is incapable of describing the *variations* in this selection on which the whole matter of learning depends. Unfortunately this was the case with Adams' theory and he admitted that "this difficult problem remain[ed] untouched." (ibid., p. 143) Even reduced to "modest" dimensions, the issue of the nature of motor programmes, movement representation, commands, or whatever it is called, reappears as a central problem of movement production, whether it be from the point of view of the systematic description of performance or from the point of view of its gradual enhancement.

From its inception the line of thinking initiated by Adams (1971) has had a strong influence on the orientation of research in motor learning, as well as in performance studies, to the extent in which the model proposed by Adams is erected upon a theory of performance. That theory drew heavily from cybernetic ideas, and in the case of low levels of control achieved its most polished and powerful formulation in the guise of various correction servos diagrams. (Pew 1974) Adams' theory can be seen as a tentative extension of these mechanisms to the level of voluntary control and learning, a leap that is essentially made possible through the complexification of the input and



output servo parameters. Input now takes the form of a memory trace, while output, at least partly, is to be considered or compared in the light of a perceptual trace. These structures or traces cannot easily be quantified simply, although they could be quantified, and will have to be eventually if the servo scheme is to give way to a true model of performance and not merely remain at the level of integrative analogy.

Schmidt (1975) revised, criticized, and extended Adams' original model, following closely on some suggestions made by Pew (1974). His theory enlarged considerably upon the idea of the recall-recognition function cycle while maintaining the same basic framework. The recall and recognition traces now grew into recall and recognition *schemata*, a notion derived from Bartlett's (1932) theory of memory. The whole theory then attempts to specify the mechanisms by which schemata are modified, how they evolve and become more precise (adapted) and powerful (general); it does *not* however specify the mechanism by which a given schema originates. But given that "the individual makes a movement that attempts to satisfy some goal," (Schmidt 1975, p. 235) the implication seems to be that some independent selection process handles the original command. Now the commands themselves are said to be of the nature of "generalized motor programs for a given class of movement." (ibid., p. 232) This is postulated. We must also assume that the selection of a particular programme is also postulated. For only by supposing that a particular programme has already been selected is it possible to understand that "the performer's problem in choosing a movement is the determination of the response specifications that will modify the existing stored motor programs." (ibid.) Here "a movement" must mean the instantiation of the programme into a particular member of the set of equivalent movements which it presumably defines. The performer's first problem, viz. that of the selection of a member of the class of programmes, as well as the theorist's problem of the nature of such elements, are both left untouched. And although the admitted central concern of the paper was upon "the development of motor programs and their response specifications," (ibid., p. 233) the most precise statement concerning a possible mechanism is the following:

... four sources of information - initial conditions, response specifications, sensory consequences, and response outcome - are stored together after the movement is produced. When a number of such movements have been made, the subject begins to abstract the information about the relationship among these four sources of information in a way suggested by the dot-pattern discrimination experiments. (ibid., p. 235)

Now it seems here that the original command, which in our view should be of primary importance in the abstraction process, has been left out. Of course if this process is supposed to give rise to *new* schemata, one would not expect them to play a central rôle in their own generation. However that leaves open the question of the grounds for the original similarity of the produced series of

movements. That is, we cannot now suppose that the similarity depends upon the prior existence of a "stored motor program." On the other hand if the suggested dynamics concerned explicitly the modification of an *existing* schema, then one would not expect it not to be mentioned. One should not, furthermore, be left with the impression that the reference to "dot-pattern discrimination experiments," studies that seemed to demonstrate the action of a schema-forming process in a perceptual task, was a reference to an accompanying schema-forming mechanism; to our knowledge the development problem has not yet been solved in perception either. In that respect it is interesting to note the fact that the idea of schema was originally proposed in the area of *receptive*, as opposed to *productive*, functions. This of course meant that it had to be adapted to the latter kind of task. Even in Pew's (1974) original reformulation however this perceptual bias is still quite evident when, speaking of skill acquisition, he asks "What properties of a *movement sequence* are *encoded*?" (ibid., p. 28, our emphasis) If one must suppose a class of movements to be extant when the learning process starts, then in a sense one is also postulating the existence of that which is to be learned (for how could they empirically be similar otherwise?). This will not do.

But it could be argued that such a statement can be understood to mean that the encoding (i.e. the learning process) is carried out on *command*, as opposed to *movement*, sets. On that account Pew does make some suggestions, but they do not seem in conformity with this supposition. He states that "whatever it is that is stored, it is not simply a specific set of motor [i.e. muscle] commands," (ibid., p. 27) but rather that it could be a "representation of a path in space through which the members of the body will move," (ibid., p. 31) entailing of course the necessary existence of a stage where "translation of the stored program into a temporal string of motor commands" (ibid.) occurs. Introducing this intermediate phase in the *production* of movement however probably further removes the basis of abstraction by one step, rendering even more important the problem of knowing what the nature of "motor commands" is. But then it is no longer clear where exactly learning is supposed to take place. Pew's main purpose however was not the formulation of a full-blown theory of skill learning; he was in the initial stages of exploring an idea, and trying to "organize our ignorance on the basis of logic, speculation, and some limited evidence." (ibid., p. 33)

In a later paper (Schmidt et al. 1979), moving away from learning issues and reverting back to the classical woodworthian problem of the accuracy of movement, we find the subject of programming still occupying a requisite position. However a definition is now provided that lifts some of the concept's ambiguity as encountered in the above discussion. They state that

the motor program should be considered as an abstract memory structure containing codes capable of being transformed into *patterns* of movement. The patterns produced from a given program have certain invariant properties, even though two responses from the same program might have large differences in other respects. Under this view, the program is *generalized*, so that *parameters* are required to

specify the particular way in which the program is to be executed. (ibid., p. 417)

Here "patterns" must certainly refer to the empirical counterpart, as evidenced in performance, of the schemata, that is, the "abstract memory structure." Now it is presently suggested that these patterns have variant (i.e. parametric or non-essential but determinant) *and* invariant (essential) characteristics. While it is further suggested that the secondary or parametric values are used within the system to specify the *command*, we are left with the impression that the primary or essential features of movement command are to be deduced from a study of performance. This is in fact what they then proceed to do, as regards in particular the time and force dimensions. Such an approach seems to rely on the basic assumption that "the invariant features of motor *responses*" (ibid., our emphasis) will be found in *all* movements. They would in fact correspond to a sort of command language through which, by the proper specification of parameters, a system might finalize a program. This is indeed the central problem of movement production. Yet again it seems that a residual ambiguity persists in their treatment of the "generality" of motor programmes. For on the one hand it is possible, as we have just done, to interpret it in terms of the whole domain of movement, while on the other hand we find the following statement in which we must now understand the generality of the programme to be associated with a *single* schema: "motor programs contain certain invariant features of the movement, and the movements can be produced differently on different occasions by application of different parameters." (ibid., p. 418) In the earlier quote it was suggested that programmes are built up from codes, that codes give rise to patterns, and that these patterns have invariant properties. Now the whole discussion boils down to the question of the relation between these codes and the properties of the movements to which their combination gives rise. By framing the statement in relation to *performance* they effectively bypassed the question; and indeed there is little effort to bring their experimental findings within its purview. In that context, it is interesting to note that the authors conclude that "Unfortunately, the principles that have emerged are slightly different for the different kinds of movements we have used, and we must deal with each of them in turn." (ibid., p. 445) They did however hope to be able eventually to "combine these relations with principles of error detection and error correction to describe a composite model of movement control." (ibid., p. 448)

In a later paper (Schmidt 1980) some distance appears to have been covered toward its achievement, although problems of error detection and correction are no longer a main concern. We do find however a further discussion on the nature of both variant and invariant properties of movement. It is proposed (ibid., p. 149) that

these invariant features are (a) the phasing (or relative timing) inherent in the sequence, (b) the sequencing of elements, and (c) the relative force, or the relation

among the various forces produced by the different contractions in the sequence. If we assume that these aspects of a response are represented in an abstract program, the addition of certain *parameters* should allow this basic pattern to be produced with slight variations.

It becomes quite clear now that invariant properties are bound to *single* movements, and that they should be understood as referring to those properties of *all* movements that define their identity. Thus although the model of control proposed by Schmidt emphasizes the variable nature of the parameters (movement time, force, and muscle selection), it must also be assumed that some mechanism exists for determining the specific values that must attach to the supposed invariants in order to produce, along with the selected parameter values, the programme whose execution unfolds into the observed movement: i.e. for a *given* movement the phasing, the sequencing, and the relative forces involved are *invariant*. Evidence contrary to the supposition that phasing occurs for the simple case of unidirectional movements leads to a tentative revision of its rôle, which initially was conceived of at the level of muscles, into a higher level of control whereby it now determines "when a new equilibrium point is to be specified." (ibid., p. 162) This revision appeared necessary in view of the apparently strong empirical support for the "mass-spring" hypothesis according to which "the limb moves to a new position via the specification of a new equilibrium point, achieved by changing the level of excitation to one or both of the opposing muscle groups." (ibid., p. 156) In the more general multi-joint case, this leads to a view of programming, at least as concerns phasing, where "each of the finger moves is defined by an equilibrium point, but [where] the distribution of these equilibrium points across time is defined by the phasing in the motor program." (ibid., p. 162) But this revision seems to have drawn a new invariant (or parameter?) into the picture: the equilibrium points, for *each* joint. On this question, which is apparently not recognized explicitly, and on the further and more subtle question of the overall 'feasibility' of the scheme, little is said. We would simply note at this point that the requirements of the phasing and equilibrium specifications inherent in the production of even a medium-sized class of movements could apparently become prohibitive: if this were the case it would defeat the "storage problem" which was at the root of the schema solution proposed earlier (Schmidt 1975). On the other hand it might well be that in fact only a handful of different phasing and equilibrium specifications are necessary in order to achieve a very large variety of (empirically) different movements. Yet neither the theory, nor for that matter the experimental research, which typically restrains the range of responses to two or three 'archetypal' movements, offer avenues for such crucial problems to be approached and such central questions to be decided.

This completes our survey of the main contributions to what we called the performance approach to the problem of movement control. We wish to conclude the present sub-section by

drawing attention to some of the salient features of this overview. This will also be useful for the next sub-section covering some contributions to what we called the system approach to the problem of movement production, as we believe that very similar remarks will equally hold true for them.

The development of ideas on movement control from Woodworth to the present does not follow a linear progression. Rather, its path is strewn with various twists and turns. A unifying thread or general conceptual framework can however be perceived throughout. Basically it seems to be organized around the idea that movement control obeys a dynamics which involves both planning and correction functions, where planning refers to a form or intention or meaning for the movement that is not directly (in the typical case at least) related to the immediate motor circumstances, while correction refers to those aspects of specific instantiations that are highly related (if not completely dependent) upon them, and comprising the whole domain of perceptual involvement in movement.

An important aspect of the reflex-chaining hypothesis for movement control that, to our knowledge, has rarely been underlined is the idea that it met, however unsatisfactorily, a fundamental preoccupation with the *production* of movement. This concern however has apparently failed to remain central within the domain. Most movement studies in Psychology since the time of Woodworth have *not* dealt with this problem, possibly because it might have been thought to belong more properly to Physiology. But this view is too simplistic. It would nevertheless be untrue to say that performance studies led to no suggestions at all concerning production, or that those ideas were all borrowed directly from Physiology, which, on the other hand, borrowed nothing from Psychology (cf. next sub-section). Even in studies oriented purely toward performance, some assumptions still need to be made on the nature of the production process or mechanism. For a long time however simply supposing that movement was preprogrammed seemed to suffice, and indeed many studies in the recent decades have been built around the question of whether or not, and to what extent exactly, movements could be preordained (cf. Jones 1978); but even though such suppositions say something about production they are far from constituting a theory of production. Starting sometime in the fifties, a tendency to refer to and use neurophysiological evidence in movement studies became more pronounced, but again, contrary to the simplistic suggestion above, this enlargement of preoccupation did not seem to lead to a renewal of interest in the problem of production. Rather the deepening of theoretical concern manifested itself as a move toward 'cognitization' rather than as a move toward 'physiologicalization.'

This at last led to more detailed pronouncements on the possible nature of a movement control system, whose most precise formulation is to be found in Pew (1974) and Schmidt (1980). But in our view, the schema model, as a theory of production, does not go very far beyond suggesting that certain aspects of movements are essential while others are variable. It is a start, certainly, but one

would expect more from a theory of production, something like a demonstration, if not direct proof, of its plausibility, or capacity to handle the diversity of movement which it wants to account for *before* it is put to the empirical test. However it seems there is a widespread belief that such proof can only be empirically forthcoming, and that model building should (i.e. can only) be very closely guided, or guarded, by experimental work. Of course it does not follow that models will spring from such research; this is contrary to the very process which is unfolding in the area. However, given this, it is still remarkable that nearly the only criterion for the interest of a model is that it should be amenable to empirical testing within a short time; while the idea that it should be testable, and preferably *critically* testable, is a basic tenet of the scientific process, the desire that it should be so (nearly) *immediately* is not. Thus the specification of a model seems to be carried only as far as it can lead to empirical prediction; in particular this leads to a situation where it is not necessary for a model of production to effectively 'produce' before it will allow inference of some empirical consequences. This was the case for the hypothesis that "movement can be preprogrammed" for example.

Movements are intrinsically highly variable, and while a lot of that variability seems, to use statistical expressions, to be of the "between" (across movements) type, it has been known for a long time that a good measure of "within" (single movements) variability also exists. Furthermore while we "explain" the first type with reference to intentional or purposeful variation, the second type, which is much more visible, and in which we control for intention or result, is much less amenable to factors external to the product. "Within" variation of course refers to the problem of equivalence. This phenomenon, although it has not to our knowledge received the kind of experimental attention which its importance would imply, is nevertheless not completely ignored within models; in fact the idea of a schema as originally proposed by Bartlett (1932) was issued with reference to the existence of equivalence classes in perception and cognition. However, again, the observation of the fact, for example that we will perceive something called "linearity" across a relatively large variation of differing instances (cf. chapter 1), and the supposition that there must exist some mechanism which allows for the perception or extraction of an identical characteristic, namely linearity, within the set of variable instances, does *not* constitute a model or an explanation of the phenomenon; it merely constitutes the formulation of a problem. Similarly the observation that many different movements lead to the same result or end, and the supposition that a singularity of purpose (or whatever it is called) can account for, or must lie behind, this multiplicity of equivalent responses constitutes the formulation of the problem of motor equivalence. As discussed previously, Hebb was the only one to approach the problem with any degree of explicitness. The adoption of the schema idea by Pew (1974) for example did implicate some attention into equivalence, but the resulting model only provides a demonstration in principle of the achievement of a single

purpose in equivalent ways. It is basically argued (e.g. by Schmidt 1980) that changing some parameter values, for example relative force, will do the trick. This is paramount to saying for example that in the perception of linearity, orientation is ignored. That might be the case, but the question remains: *how* can it be ignored in some cases and not in others. This issue is in fact related to the difficult general question of the relation between the two types of variability and how they are to be treated within a system. It has not been addressed in the schema theory.

Some further central points that appear to be left in a state of fuzziness in current performance models have to do with the characteristics of commands and their relation to resulting movements. The question of the extent of a programme (or command), that is, its 'size,' or the question of how 'large' a movement it can sustain, are of particular interest. Nearly all possibilities have been held at some point, from the idea that no programming occurs at all to the opposite suggestion that all movements are programmed. The fact of the matter seems to be that the existing performance models do not allow for the identification of a programme, or at least the description of what a command or command sequence might be, along with a description of the generated output. This is probably because they are not elaborated as *working* models, which boils down to the simple fact that one cannot use the current performance models as blueprints for the construction or simulation of a control system. We have no doubt that the process of theorizing would be markedly enhanced and would greatly benefit from the adoption of such an outlook (cf. chapter 0). This is in fact what we have attempted to do.

Pew (1974) had noted that while "Process-oriented views of perception, memory, and decision are already well-advanced, [] Motor performance has been the laggard in this respect." (Pew 1974, p. 1) Interestingly, movement studies nevertheless seem to have caught up quite rapidly since then, which might leave one thoughtful about the true depth of these views in cognition at large. Be that as it may, he then defined "the unique subject-matter of motor performance" to consist of "movement control, utilization of response-produced feedback, and the organization and patterning of behavior in time." (ibid., p. 2) Too often however such departmentalizations seem to derive from the specific problem-context arising out of the use of particularly 'rich' paradigms (such as pursuit) rather than from an analysis of the basic problem. (Thus the area of motor equivalence for example is not explicitly recognized as contributing to the subject-matter, although we would argue that it would be if an equivalence paradigm existed. It could probably not, however, be fitted neatly into one level.) In our view the unique subject-matter of motor performance is the problem of movement production.

Pew suggests that "The reader should think of motor control in terms of a hierarchically organized system [an idea which also dates back to turn of the century performance research; cf. ibid. p. 32] in which the distinction among levels is diffuse and in which there is a rich interplay among the

various [presumably feedback] processes." (ibid., p. 3) He concludes that "The question of what is a motor program and what do we mean by "automating" a movement are to me inseparable from a more global analysis of temporal and spatial organization. There is no level in the motor system at which I am willing to say, "Here is where a motor program comes into play."" (ibid., p. 33) Indeed we also believe that most likely every part of the system is intertwined within all the others, but we also believe that understanding should aim for their gradual disentanglement. Given this formulation of the ultimate goal, one could wonder how a "global analysis" should be carried out, and indeed be surprised to discover that a solution to the problem of "the organization and patterning of behavior in time" is said to depend upon an "analysis of temporal and spatial organization."

The idea of embedded layers of control, or equivalently of a hierarchical organization within a control system, is one which has pervaded in turn what we call system-driven research in movement studies, especially in Neurophysiology. To this we shall presently turn.

2.1.1 System-driven research

Whereas, using Woodworth's terms, performance-driven research into the problem of movement production and control essentially focuses on the "result" of motor action, and endeavours to uncover what might be called the laws of significant or purposeful bodily motion (at least insofar as the achievement of specific motor results can be said to touch upon the matter of meaning), system-driven research for its part appears to be overwhelmingly preoccupied with the "possibility" of motor action. Now it is generally held that this possibility depends upon and is a direct function of certain parts of the nervous system along with their associated effecting system, i.e. the musculo-skeletal system. Simply put, "Movement is the result of the motor activity of the muscles acting upon the levers of the skeleton." (Paillard 1960, p. 1682) Since "Every movement results from the cooperative action of a number of muscles" (ibid., p. 1683) and since "the motor activity of the muscles" is a function of their innervation, that is, the effect of the activity of the neurone groups associated with them, the problem of "the neurophysiological basis of the patterning of skilled movements must be defined in terms of 'coordination'." (ibid., p. 1684) Hence the problem of movement is immediately stated in terms of the (nervous) system that is believed directly responsible for it. This contrasts with the movement-as-performance view where, even though the product's existence is recognized as dependent upon a system, the variables and the problem are primarily defined with respect to the product itself and not with respect to its 'physical' generation. The emphasis on generation, or production, is thus at the very core of system-driven research on movement. Paillard (ibid.) states the problem very succinctly:

We must seek for a mechanism which can select a privileged combination of mus-

cies from among a considerable variety of possible combinations.

In this connection, two fundamental questions arise: what are the factors which determine the choice of the muscles to be engaged in a given movement and the variable spatiotemporal ordering of their activation, and how does this selective operation operate in the nervous system.

It turns out, unfortunately, that muscles are not simple elements within the motor chain, in the sense that they would constitute the functional units from which are built up all movement commands, and that the level at which they apparently can become functional units (the motor cortex) is found quite high up in the nervous machinery. Muscles are in fact structurally made up of a variable amount of variable-sized components (bundles of muscle fibers) that, along with the single motor neurone that binds them together functionally, make up what Sherrington (1925) defined as the "motor unit." (cf. Buchthal & Schmalbruch 1980) It constitutes the functional unit upon which everything motor, from postural adjustments (Roberts 1978) to the most sophisticated human voluntary productions, rests. It follows that from the neurophysiological point of view the selection of particular muscle groups has to be effected through the selection of particular groups of muscle nerves, called " α " (alpha) motoneurons. Since these nerves subserve the whole variety of motor 'purposes,' it is not surprising to find them impinged upon by very large numbers of other neurones, some of which stem from the muscle tissue itself (the sensory nerves, of type "Ia," of the primary or spindle muscle receptors, also called intrafusal receptors), others (the interneurons, among which are found the Renshaw cells) belonging to the grey matter of the immediate or distant spinal environment, and finally still others emanating from a variety of regions in the heights of the central nervous system. There is more: a second type of spinal motoneurons, the " γ " (gamma) fibers, also attaches to muscle control but in an indirect manner, since these neurones stimulate the intrafusal fibers lodging in the spindle receptors. These muscle fibers in turn are equipped with a second type of sensor, the secondary receptors, from which proceed the group II afferents. (Prochazka 1981) Two other systems complete the sensory apparatus of the musculo-skeletal system, namely the Golgi tendon organ (group Ib afferents) and the joint receptors (Proske 1981; Skoglund 1973; Matthews 1972) In effect, muscles are not simply the purely motor or output end of an overall organismic input-output loop: rather they already constitute, along with the associated spinal circuitry, an elaborate, intricate, and semi-autonomous, sensori-motor system, one which must be taken into account by the higher centers. Thus, the follow-up servo hypothesis of movement production (Merton 1953; Hammond, Merton, Sutton 1956) attempted to show how some of the spinal reflex circuits could be exploited by a central command. However this hypothesis was later invalidated. (Matthews 1972; cf. Burke 1981)

The spinal machinery appears formidable. Indeed it apparently led one author to nearly admit that it is beyond comprehension: "One can speculate [on the rôle of the fusimotor system] but, as

in the past, any speculation is likely to be seriously flawed." (Burke 1981, p. 121) Yet he goes on to suggest that

The role of the muscle spindle in the control of movement is probably not related solely to its reflex effects onto the motoneuron pool. Equally and possibly more important are its projections to the higher centers involved in elaborating the *movement program*. The muscle spindle's contribution is probably subtle, its *feedback integrated* with other sources of information. (our emphasis)

Paillard (1960) had for his part already suggested that the cortical "keyboard" would play a theme to the accompaniment of the spinal "keyboard," while itself being subject to the whim of a (possibly sub-cortical) director. Thus he advanced the idea of an elaborate communication system in which each node would respond in a specific fashion, and in which, for example, outside of their own intrinsic 'logic,' the spinal circuits would be subject to external (cortical and subcortical) driving impacts, mediated through the pyramidal and extra-pyramidal systems of fibers. It is as if the spinal neurones could be driven, but only in certain directions; the superimposed cortical commands then proceeding to 'push' them in certain of these directions, and in such a manner as to produce, unknowingly to the local spinal systems, more global, or 'significant,' patterns. This would explain not only that for example local signals are fed back into the command centers but also, as Burke (above) suggests, that they must be *integrated* at that higher level. However we are already trespassing the bounds of neurophysiological pronouncements.

Nevertheless we think it is worthwhile to consider the contents of the above quotation. In it we find reference to two major concepts that seem to have become "monnaie courante" in a large part of neurophysiological research into movement. Since these form, to our understanding, the basic components of what might by extension be called the current model of movement control and production in system-driven research, they deserve closer attention.

The idea of a movement programme, which seems to have arisen independently in the context of electromyographic research during the twenties (cf. Paillard 1960), has apparently survived well into the 'neurographic' era that soon afterward flourished. Reflex control seems difficult to extend functionally outside of the spinal domain, where lie the motoneurone pools and where its mechanics is apparently largely confined, although there have been more or less successful attempts to do so. (Phillips 1969; cf. Asanuma 1981) The relatively local or piecemeal nature of reflexes does not appear to be easily reconcilable in neurological terms with the otherwise temporally and spatially much more extended features of voluntary productions. In that respect the parochial character of cortical control is to be noted (Asanuma 1981), as it does not seem either to correspond to the site where intuition would have the necessary prior organization of movement commands occur, these commands of course still being expressed as sets and series of local events. Although little

reservation is manifest within the neurophysiology of movement production as concerns the existence of some sort of mechanism responsible for the observed overall ordered organization of local activity within the pyramidal tract for example, there is a much more definite restraint as concerns analysis of this problem from the more general viewpoint that seems required in order that the underlying process eventually be apprehended with microelectrodes. (cf. Paillard 1978) Some effort however has gone in that direction, the results of which occupy an important place in the latter part of the present sub-section.

The second concept to which Burke alludes is feedback. Although he had speculated that efforts at theory were "likely to be seriously flawed," it must nevertheless be remembered that the 'application' of the idea of feedback to nervous function is closely related with some of these attempts at theory (von Holst & Mittelstaedt 1950; Merton 1953), and that although the specific servo system proposed by Merton for example did not survive, the general idea of feedback control and the overall neurophysiological plausibility of the concept continued to be invested. Indeed already in Paillard (1960) not only is the idea of reflex itself tentatively revised in light of the principles of self-adjusting mechanisms, but an overall case is made for viewing the whole sensori-motor apparatus in that manner. This view was developed more systematically in Paillard (1980). But rather than constituting the basis, as in the case of Merton (1953), for the elaboration of specific models of control, it seems that the cybernetic framework has remained both general enough to encompass a large number of diverse studies touching upon nearly every known motor structure and yet specific enough to provide ideas for their functional as well as anatomical investigation.

In that respect let us recall here a statement made by Cajal over ninety years ago concerning the study of nervous microanatomy. Speaking of the devoted patience required on the part of anatomists "pour débrouiller la trame délicate de l'axe encéphalo-spinal, [] la machine la plus subtilement compliquée que la vie puisse nous offrir," he added that "ils étaient guidés, cela ne fait point de doute, par l'espoir que la découverte de la clef structurale des centres nerveux jetterait une vive lumière sur les importantes activités de ces organes." (Cajal 1894, p. 444) In other words this meant, approximately, that in order to understand how the neuronal machine works, we must first understand how it is built. We believe that the appeal and influence of cybernetic ideas in system-driven research in movement stems largely from the fact that they allow for the sort of inference from structure to function which Cajal talked about. Unfortunately this seems to work only in a very approximate manner, in particular for what we might call the general or overall layout of the system, which is indeed organized according to an input-compare-output scheme; however that had been proposed before the advent of Cybernetics. That it also turned out to shed some light on local circuitry, especially at the level of reflexes (in particular the stretch reflex), constituted a major demonstration of how what until then had largely been defined in behavioural terms could be shown

to be coherent, or rely functionally upon, local circuitry; however this was not a case of going from a knowledge of structure to an understanding of function but rather the reverse.

At one point it seemed as if the micro-structure of some of the circuits, in particular the spinal relays, were in fact simple servo and servo-assisted loops. But even there it is no longer clear that such simple schemas tell the whole story. When it comes to central structures whose micro-anatomy is well-known, in particular to the cerebellar cortex, then it seems clear that the circuits involved are no longer recognizable at all in cybernetic terms, even though from a global point of view it is difficult to abandon the notion that somehow we are dealing with a comparator or computer of some sort. The complexity of the nervous system seems nowhere to be more apparent than in the fact that events at a local or micro level do not recognizably generalize to global function *and vice-versa*; while it is generally believed, following Cajal, that the key to the nervous system is to be found at the level of local events, we find very little effort expended toward the problem of relating them with global functions, or indeed with other closely related local events. Thus the micro-structure of the cerebellar cortex is crystalline, which seems to suggest that locally it is oriented toward a single task; yet even though we have good information on the nature of its global integrative expression, the problem of proceeding from one to the other, of showing how the global operation ultimately depends upon the local network, or conversely of showing how a particular local operation, when 'multiplied,' can lead to a functionally new whole, has not received within Neurophysiology the kind of attention which its importance suggests. Following Paillard (1960) this might have been due to the fact that at the time such problems entailed insurmountable experimental difficulties. However as he later added, "It is my firm conviction that, in the last resort, every significant advance in neurophysiology must always refer, explicitly or implicitly, to models of how we imagine the system works, as well as to models of what we imagine the system is doing." (Paillard 1978, p. 156) For our part, we believe that solving the above problem would indeed constitute "a significant advance in neurophysiology."

We also believe that finding a solution to this problem, or more precisely that part of it which consists of relating local nervous events with global events or functions, is the overriding purpose behind some major experimental efforts into the involvement of the higher centers in movement production and control that have been expended over the last twenty years. Following upon the experimental success of such attempts in vision, particularly in the work of Hubel & Wiesel (1962), single-cell recording techniques have been adapted and applied to the movement domain. Although the idea that neurones constitute the primitive or local 'physical' constructs of the nervous system dates back to Cajal (1894), the idea that single neuronal activity corresponds to local events, or reflects local function, and that it is relatable to global function is a much more modern hypothesis that apparently also has roots in "models of how we imagine [a particular]

system works" like the one conceived by Selfridge (1959). That such a system was 'artificial,' and grew out of a preoccupation with the creation of machines (in the sense of chapter 0) displaying human-like capabilities, is secondary to the fact that it proved to have some relevance to the problems of studying natural systems like the CNS; no doubt that is because the problems to be faced in *constructing* human-like machines are fundamentally not different from those encountered in trying to *understand* their natural counterpart.

However, even though the problems are similar, the methods markedly differ, probably because the single-cell neurophysiologist is in the presence of a working system while the 'system' theorist is not. Yet having to select the most promising or revealing site in which to implant his micro-electrodes the neurophysiologist must refer to some understanding of the system. Now the 'best' kind of knowledge in that sense is anatomical. Thus for example Evarts (1980, p. 224) states that "In initiating studies on the activity of the brain in relation to volitional movement, it seemed reasonable to start by examining the output of the cerebral cortex - investigating those neurons which send their axons to the spinal cord." But there is more to this choice. Over a century ago Fritsch & Hitzig (1870) had demonstrated that part of the cerebral cortex, until then believed to be exclusively engaged in higher mental processes, was involved in the 'lowly' task of muscular control. It seems nevertheless that the tendency to ascribe higher functions to the cortex, even though they were no longer restricted to the 'purely' cognitive domain, remained strong: the motor cortex and pyramidal tract neurones were now thought to be involved in the *programming* of movement, and in particular *voluntary* movement, while the sub-cortical structures, including the basal ganglia, the lateral cerebellum, and the ventrolateral thalamus, and constituting the extrapyramidal system, were concerned with *reflex* or *automatic* movements. (cf. Evarts 1980; Paillard 1978) These two systems were furthermore believed to be quasi-independent. But these views and assumptions have had to be revised drastically in the light of "anatomical and physiological studies [that] now make it clear that pyramidal and extrapyramidal systems are *not* separate." (Evarts 1980, p. 242) The revised picture now states that "Whereas the traditional view held that the cerebral motor cortex was at the highest level of motor integration and that the subcortical structures were at a lower level, that is, closer to the muscle, it now appears that the situation is quite the reverse." (ibid., p. 244)

Like others before him, Evarts had, on comparative anatomical grounds, argued for the critical importance of the pyramidal system in "those aspects of motor behavior that are most especially human." (ibid., p. 224) in particular manipulation, speech, and fine adjustments. It is certainly an easy step to conclude from this that the pyramidal system's motor contribution is of the highest nature, an assumption that seemed to have biased the interpretation of his "finding that discharge of some PTNs [pyramidal tract neurones] varied with load in tasks requiring either fixed-velocity

movements [] or maintained posture" (ibid., p. 231) in that sense. He explains that "My own early formulations concerning the significance of this finding on the relation of PTN discharge frequency to load involved the assumption that the PTN was providing a central program or command specifying some peripheral event. I had given little attention to the altered feedback that might have been reaching the PTN with different loads until Phillips [(1969)] proposed that changes in corticomotoneuronal activity with changed loads might be the result of a mismatch between actual and intended movement." (ibid.)

This interpretation is largely responsible for a most remarkable revival of servomechanistic models and ideas in the neurophysiology of movement control during the seventies. (cf. Desmedt 1978; Miles & Evarts 1979) Interestingly, even though the supposed programming tasks associated with the motor cortex have now been relegated to the subcortical structures and centers, which are now seen as 'driving' it, Phillips' hypothesis remains consistent with the cortex's supposed critical contribution; for it is easy to imagine for example that the corticomotoneuronal loops might be 'preset' in a manner similar to the way in which Merton had earlier imagined the spinal loops to be preset by the γ system. But all these views were foreshadowed in Paillard (1960) who also referred to Penfield's (1954, p. 5) "central co-ordinating integrating system which we may call centrencephalic" from which would issue streams of "volitional impulses." Paradoxically, this growing understanding of the 'proper' rôle of the cerebral cortex and the pyramidal system in movement production has removed even further, and into still deeper recesses of the brain, those structures and processes whose 'tapping' would finally unveil the mystery of movement coordination, or initiation, or as we now say, programming. In that context it is noteworthy that just as Paillard (1960), in discussing the neurophysiology of will, had reverted to clinical data, so twenty years later still does Evarts (Evarts 1980), adding however, speculatively, that since "The entire cerebral cortex sends fibers to both the basal ganglia and the cerebellum, and these two structures in turn send massive connections back to the motor cortex by way of the thalamus," that "The inputs going into the cerebellum and into the basal ganglia may be coded in a more abstract and complex manner than the outputs leaving the motor cortex." (ibid., pp. 243-244) Efforts to establish and understand the functional relation of local nervous events within the cerebellum (Thach 1978; Harvey 1980) and the basal ganglia (DeLong 1974; Marsden 1980) with selected global movement characteristics has indeed proved to be difficult and shown them to be abstract and complex; that is, they still remain largely unknown.

We would suggest that this is because apart from the servo loop there are little if any concepts that pertain directly to possible functional modes of neuronal interaction, of the type that would be useful in interpreting these data and capable of explaining how such a homogeneous structure as the cerebellar cortex and its nuclei can have and process "primary representations [] of signals relating

to start, possibly stop, direction, joint position (or muscle length), pattern of muscular activity, and force; and [] a velocity signal." (Thach 1978, p. 674) Certainly one can argue that such observations are consistent with a hierarchical model of control (ibid.) in which an ascending order of variables, from the 'concrete' muscle-related signals to the more 'abstract' non-muscle-related ones, could in principle be established; however serious doubt can still attach to such interpretations in view of the fact that "none of the neurons whose discharge is correlated with some aspect of movement have been shown to actually play the causal role the correlation implies." (ibid., p. 675) Indeed following Paillard (1978) we would further submit that showing for instance that muscle length signals do enter critically into a cerebello-cortical loop requires first an idea on how a length signal is *functionally* represented in the structures involved, not in terms of single neuronal signals, which we already believe it to be, but rather in terms of the control process itself; this in turn implies for example that the problem of the *integration* of the (local) neuronal signals, probably connected with the spindle receptors' central projections, from which are presumably built up the muscle length signals, be solved. But we have not seen these issues, which we consider of central importance, addressed explicitly or even recognized as such. Of course we fully recognize, on the other hand, that the intricacies of central nervous anatomy, and the imposing mass of neuronal tissue more or less intimately involved in the production of movement, renders the task of unravelling the specific functional contribution of each to the whole business of control a very difficult enterprise. We are merely suggesting that it might be facilitated by a deeper preliminary investment in some theoretical problems such as the one described above.

As a further example, consider that as Thach (1978) seems to suggest, and as we are also inclined to believe on anatomical grounds, the fact that the cerebellar cortex 'acts' according to length, direction, etc., does not allow us to conclude that these are actual control variables since, except for the "pattern of muscular activity" signal, they are also all *performance* variables. Furthermore, as suggested earlier, it seems difficult to reconcile the cerebellar cortex's homogeneous micro-structure with this apparent multiplicity or diversity of processing. Now as we tried to argue in the previous sub-section, a performance variable can only be transformed into a control variable through its integration within a model of production. In the case at hand this would require that a cerebellar control or processing mechanism be postulated that can account for the observed diversity.

In their review on "Concepts of Motor Organization," Miles & Evarts (1979) reserve the largest space to a discussion of models derived from Control Theory pertaining mostly to the oculomotor system, in which "two types of control - open-loop and closed-loop - operate [] with a machine-like quality [making it] particularly amenable to an engineering approach." (ibid., p. 337) This is prefaced by a survey of the critical studies purported to have given rise to the idea of central motor programmes; according to them, "Sherrington's work on the scratch reflex underlies our

concepts of the *triggered movement* based on a *central program* involving a *spinal rhythm generator*." (ibid., p. 329) Thus the concept of motor programme was originally based on automatisms such as the stretch reflex; it was later extended to cover locomotor pattern generators. Interestingly, it was equally extended into the realm of voluntary control, the apparent converse of automatic or reflex control. Our understanding of this generalization is based on the necessity of assuming in *both* instances the existence of a mechanism responsible for the *patterns* which characterize these very elaborate sequences of activation, be they of the muscular or neuronal type. In the case of automatisms, the margin of variability within these patterns appears to be much more restricted than in the case of voluntary productions, although as often pointed out and as shown through many experiments, it is not inexistent. The case of the oculo-motor system is attractive in that its domain of variability is very well circumscribed; and indeed this is where the models are most developed. But the domain of variability associated with an effector system such as the arm is much less tractable within "Control Systems Models;" consequently the models are much less developed in that area. For this broader case the input signal, which in the case of the oculo-motor system would be a signal representing "Visual Target" for example (ibid., Figure 8, p. 349), is simply referred to as a "Central Program Control Signal." (ibid., Figure 2, p. 336) It is not specified according to what principle, or what type of variable, it would drive the PTNs that are thereafter engaged in an elaborate correction loop.

It seems then that the idea of programming largely remains at the level of descriptive metaphor, referring either to the character of highly ordered sequences of movement and muscle activation or the supposed highly ordered sequences of local neuronal events that are deemed responsible for them (of which the strict and fixed ordering of spinal α -motoneurons, possibly γ -motoneurons, and Ia receptor recruitment within movement provides the most recent and spectacular indication - cf. Henneman 1981). With respect to the models encountered in the oculo-motor domain, where the nature of the input signal is relatively well known (e.g. retinal slip, etc.), the discovery of the nature of the command input/s to the servo loops presumably involved in the voluntary control of other effector systems such as the arm, along with its (their) mode of production would certainly transform them into a comparably powerful model of control. An identical remark holds for tentative applications to voluntary control of the principle of efference copy or reafference (von Holst & Mittelstaedt 1950; von Holst 1954) - e.g. to explain or make sense of α - γ co-activation (ibid., p. 358) - and corollary discharge (Sperry 1950), both of which stemmed from work in oculo-motricity, and both of which were not addressed to the issue of patterning as it is understood in the larger field of movement production. Thus even though it is rather easy to suppose that these same principles are at work in *any* control system, it is rather more difficult to demonstrate how, by changing the nature of commands and subsequent afferents, it is possible to

transform an oculo-motor control system into an arm-movement controller. Such difficulties are apparently directly linked to differences between the interfaces' structures, which require different commands, and hence to consequent differences in the meaning of resulting afferent signals. While in the case of the eye it has been thought for a long time that retinal slip can be precisely associated with oculo-motor signals, in the case of the arm the equivalent 'meaningful' input-output pairing or pairings has remained much more elusive.

In effect this leaves largely untouched Paillard's two fundamental questions, especially the second one concerning the operation of selective muscle activations, which we understood to be referring to the processes *preceding*, or at least concomitant with, but still certainly different from, the engagement of auxiliary or subservient servo loops. Now it is generally agreed that some requisite nervous activity is essential for movement, activity whose purpose is now understood to be concerned with the generation of signals which will thence be fed into cortical servo loops and possibly fed back into their place of origin. Again, referring to the oculo-motor domain, what is needed here is the specification of an overall *directive*, such as pursuit, which is then taken up by servo loops explicitly contrived to carry out, among other things, such a purpose. However it is apparently not possible to encompass as large a part of the activity of an effector such as the arm within the confines of a directive such as pursuit, even though humans of course do engage nearly as successfully in tracking tasks involving their arm as their eye; that is, the domain of variability associated with an arm is much larger than that associated with the eye. This is because it is a 'richer' effector; it has a wider movement potential associated with it. On that view the temptation to generalize from the case of the eye to that of the arm, given the measure of success attained with the former, can be great; but it is equally possible that the ultimate integration of both control structures will show oculo-motor control to be a special or simplified case within a control scheme involving processes and principles that are not immediately apparent in it. Since the eyeball is also muscularly driven, these principles will *not* be found active at the lowest or peripheral level but rather in what we called above the requisite preparation.

Contrary to Miles & Evarts (1979), Paillard (1960) had suggested that the idea of programming, or requisite preparation, had evolved out of a belief that what appeared to be incomprehensible order, namely the electromyographic activity concomitant to a variety of movements, must in fact obey a higher order. (cf. Hallet et al. 1975) He writes: "To patterning of spatial distribution which determines the participation and the degree of intervention of the diverse muscles implicated in the action is added an extremely complex patterning of temporal succession." (loc. cit., p. 1683) Ten years earlier Lashley (1951) had discussed this problem at much greater length. Viewing the spatiotemporal and energetic relations among the muscles involved in movement as but one instance of a much more general feature of behaviour, a feature which he referred to as serial

ordering, he argued that "Not only speech, but all skilled acts seem to involve the same problems of serial ordering, even down to the temporal coordination of muscular contractions in such a movement as reaching and grasping," (p. 121) adding furthermore that "Analysis of the nervous mechanisms underlying order in the more primitive acts may contribute ultimately to the solution even of the physiology of logic." (ibid., p. 122) According to him the till then prevailing (and only) theory to account for "temporal integration" was "associative chain theory," according to which "the performance of each element of the series provides excitation of the next," (p. 114) a theory that turned out to be untenable in view of performance and neurological facts. Deeper analysis of the problem of the (temporal and spatial) organization or integration of discrete bits of behaviour into meaningful wholes led him to the idea of a "syntax of movement" (p. 122) and a restatement of the overall problem as concerning "the existence of generalized schemata of action which determine the sequence of specific acts, acts which in themselves or in their associations seem to have no temporal valence." (ibid.) He then went on to discuss the implications of this new outlook in regard to known features of the neuronal substratum. But little was offered beyond this essential reformulation except, for us at least, the very suggestive statement that "Its [the visual cortex's network of short axon cells'] integrative functions are an expression of the properties of such a network." (p. 132)

It had been one of McCulloch's hopes, as far as we understand it, to produce an analytical tool that would enable, given the structure of small neuronal networks, the deduction of their function, in a manner not dissimilar to but much more extended than, as we saw above, the application of the servo concept or understanding to the large-scale connective architecture of the CNS. Unfortunately it has turned out that each case must be studied individually and become the subject of specific hypotheses. Work on the cerebellar cortex, whose micro-structure is the best known of any part of the CNS, shows how the inference of function from structure, even for a relatively simple case like this one, is not at all an evident or straightforward step. (cf. Pellionisz 1980) Now if this holds for the simple wiring of cerebellar cortex, how much more so would it be true for the much more extensive neuronal arrangements of the cerebral cortex's columns. (cf. Szentágothai 1979, Figure 6) On that view McCulloch's dream appears to lie on a conceptual continuum in Neurophysiology dating back at least to Cajal's hope that unveiling micro-structure would bring us some way towards the discovery of function. For us however the importance of Lashley's discussion in that context lies in an apparent break with that tradition, as it is instigated by a fundamental *problem*, one whose understanding (and eventually solution), it is argued, "holds some promise of a better understanding of cerebral integration." (Lashley 1951, p. 135) Unfortunately the attraction to the (nervous) system seems to have remained stronger than the desire to approach and analyze the abstract problem of serial ordering and to propose possible solutions to it as prerequisites to experi-

mental probing. We had suggested earlier that the cybernetic concept of feedback along with its various attendant realizations had provided for Neurophysiology a ready-made formal context in the study of certain tasks where the rôle of sensory input seemed to be as important as efference in the establishment of dynamic stability. The problem of serial ordering appears either complementary to it, from the point of view of the establishment of *successive* equilibrium points - if is conserved the notion, which might be plausible at least in the domain of movement production, that *every* act is realized through the attainment of a pre-fixed reference state -, or in opposition to it, from the point of view of the apparent impossibility of encompassing within the *single* principle of feedback - appropriately tailored to every given situation - the infinitely variable possible and actual sets of spatially and temporally independent elements that define the domain of these acts.

But whichever way it is conceived, 'typical' solutions or models, or what Lashley called "conceptions of the fundamental organization of the nervous system," (ibid., p. 130) - that incidently referred more to structural than functional organization - are lacking. Following Lashley (1951) and Paillard (1960), insofar as the problem of movement production is a special case of serial or selective ordering - at whichever level, muscular or neuronal, it is stated - and insofar as it must involve principles other than the interplay of various servo processes such as is suggested by the idea of programming, or more specifically by a reference to the processes involved in the generation of most of the extra-peripheral signals impinging upon the pyramidal neurones (cf. Wise & Evarts 1981), this absence of general solution is reflected in an absence of models of movement production within Neurophysiology.

This completes our overview of system-driven research into the problem of movement control. We would like now to conclude it with a brief discussion of the relationship, from the point of view of the system-polarized approach, between the concept of degrouping and the idea of control-as-production. We have not thought it worthwhile to extend our review of work in Neurophysiology beyond the minimum that was needed to reach a clear formulation of what we think is the essential problem of control within the domain, and one moreover that would allow us to place it in the perspective of our own understanding of the matter.

As we see it the neurophysiological problem of movement control can be brought down to three main components: [1] a multiplicity of spatio-temporally quasi-independent physical elements, elements that are also structured and fixed and possess an intrinsic dynamics connected with physical motion - the 'effectors' (segments, muscles, motor units); [2] a very extensive and elaborate nervous apparatus that attaches to it; [3] the question of the relationship between events, however defined and however characterized, or their effects, in the effectors and in the associated 'motor' nervous networks. While a very great deal is known about interfacial physiology [1], and more and more is discovered about the intricacies of nervous motor relays [2], the question raised in

[3] appears much more to define the overall programme of research into the neurophysiology of motor control than the statement of a particular, albeit very great, problem to solve. However some formulations, which are not yet precise enough to constitute testable hypotheses but that are clear enough to provoke theoretical interest, have been advanced in order to maximize the search for an answer. Here we are alluding of course to Lashley's definition of the problem of order and to Paillard's question concerning selective processes. While one spoke of elements having no temporal valence and the other of muscles, it can readily be seen that both fit into the three-point framework outlined above. Following on these lines, the problem of control then becomes the problem of *defining* and *imposing* a variety of specific, or 'meaningful,' orders (or spatio-temporal structures) upon a set (or subset) of these more or less loosely connected elements. This being approximately where our introductory discussion of degrouping started, we now refer the reader back to it. For the purposes of the present discussion however a few more remarks are in order.

Cybernetic theory consists of a set of principles that are, as any theory should be, more general than any specific instantiation it allows for. We mentioned earlier Merton's (1953) application of the theory into a possible scheme (the follow-up servo) of production. This was one instantiation that proved erroneous. Next in line came Phillips' (1969) idea of the trans-cortical loop, which has been very useful for setting up experimental paradigms for the study of motor cortical neurophysiology. Its status and value are still debated. Afterward Paillard (1978) set out to attempt a complete integration not only of movement, but of the whole behaviour of organisms, within a set of interdependent embedded control loops. (Admittedly, he was not proposing a theory, but defining a setting within which could be elaborated and organized problems and programmes of research.)

In the early sixties, according to Miles & Evarts (1979, p. 328), "The use of formal *Control Systems Models* in studies of motor organization was pioneered by Stark and Young in studies of eye movement." Now considering this case of control from the point of view that we are adopting here, it can certainly be argued that the cybernetic model of the visual pursuit mechanism does solve a problem of order, and does answer to Paillard's request for a selection mechanism. However a closer look at these models reveals that they actually *presuppose* the existence of a mechanism that will for instance *translate* a (single) visual target - situated at some angle θ_t (ibid., Figure 8) in terms of eye position - into the (multiple) eye muscle commands necessary to achieve the desired goal. What does this mean? It means first of all that it is possible to arrive at a very good model of the *performance* of an 'effector,' namely the eye, in the absence of any explicit recourse within the model to the eye's true effectors, i.e. the eye muscles. Indeed these models are apparently of little help when the sensory neurophysiologist enquires as to what rôle the very rich spindle output associated with eye musculature plays in such a scheme. They might however be very attractive to performance-related researchers who believe that the same can be done for the whole body as has

been achieved for the eye, even though we have encountered very few instances of transposition from one problem into the other. This might simply be because the eye is untypical among effectors; for us the essential difference lies not with its anatomy, even though the reason for the segregation is closely related to it. That essential difference lies in the variety of purposes which an arm can answer to as compared to an eye, and most importantly to the fact that as concerns the eye the critical signals, that is the sources of information upon which operates the system, are more limited and fewer in number than in the more general case of the arm for example, with its much larger number of degrees-of-freedom, muscles, and associated sensors. Thus it is *conceivable* (and even computable) that activity levels in the six eye muscles might be translated into a single position value with respect to the head, but the equivalent calculation for an arm with twenty muscles and three articulations is a different story, the first question to be answered, supposing one wanted to arrive at a *single* positional value for the arm, being the nature of that parameter.

Now the fact that such signals are readily definable for the eye along with the fact that a large portion of the eye's behaviour is describable in those positional terms, has apparently overshadowed the issue of the variety of purposes mentioned above. If a way could be found [1] to describe all possible positions of the arm with a single (or at least a small set of) parameter/s, and [2] to describe all possible behaviours of the arm as a change in that value, then we would suggest that the problem of control for the arm has been largely solved. This would amount in effect to equating "the variety of purposes" with "the variety of (changes of) positions" of the arm. Our main point here is not to argue against the relevance of servo loops per se, but to emphasize the requisite need for the proper signal, which constitutes a problem in itself, that on first view at least, has apparently little to do directly with the cybernetic mechanism even though it is directly defined with respect to it here.

A performance approach to this problem might attempt an initial simplification, something like choosing the end-point of the arm and characterizing arm position as position of the arm's end-point in euclidean space. However such a procedure is computationally largely intractable, and quickly resolves into a system question; in other words it is only an apparent simplification, the kind of spatial knowledge and computational power presumed to exist for managing the task successfully and within natural time constraints being at least as great, if not exceeding, that required for alternative, but still unknown, *modus operandi*.

A very similar remark could be made concerning possible attempts to expand the eye movement control scheme into an arm control scheme by the superimposition or adjunction, in the manner of Paillard (1980), of larger and larger control loops, and whereby each level would only deal with one single parameter that cascades down into the next one, which in turn gives rise to its own particular signal and so on down the line, a design which certainly would not lack of variety of

intention and purposes: however again, the main problem is not with the design per se but with the resolution of the product, namely again the variety of movements associated with an arm for example, into a handful of workable 'integrated' parameters.

Because of system-driven research's interest in the periphery, or the motor interface, it is much easier to see in that case how the basic problem is in fact the problem of degrouping than it is in the performance-driven approach. Lashley's use of a syntactic metaphor for movement, which is certainly attractive if only for the fact that language is defined primarily as a problem of production and that it is thought largely to depend upon a body of syntactical rules (ideas that incidently were still embryonic within Linguistics when Lashley wrote), and Paillard's idea of a selective operation, relying implicitly on the idea that the components to be selected and ordered were 'independent' since otherwise the need for such operations is difficult to realize, both call for the solution of the problems of degrouping as outlined in section 2.0. Unfortunately neither of them pursued the analysis beyond that level. Thus even though, as discussed above for the idea of programming, one observes a certain readiness to ascribe to the nervous system a capacity to carry out processes in response to an understanding of certain problems, one also observes a marked reticence to pursue the implications of such processing (in terms for example of the demands it would make in order to occur) outside of strict experimental guidance. This probably explains why from the twenties up till today the concept of motor programming has neither noticeably evolved nor become more precise. And it has been used indiscriminantly in a great variety of contexts, ranging from the study of automatisms such as insect flight to the problem of skills learning. Certainly however, just as for the basic cybernetic ideas, the notion of programming (or selection, or ordering) is more general than any of these supposed specific instances. But there is not to our knowledge even one workable realization or demonstration of the principle to which could be referred, and which could possibly enhance, the study of any of those specific cases. In our view the insistence upon the problems of control as problems of production and the assumption of efforts to resolve them even for 'artificial' cases would go some way in helping to resolve the natural ones, or even possibly in demonstrating the inadequacy of the concept. But this calls for a greater investment in pure theoretical concerns, or in speculation, than has occurred so far.

The problem of degrouping is a theoretical issue that arose out of, and was strengthened by, a study of the problem of movement production within both system-polarized and performance-driven research, and it is one which we believe to be at the heart of *both* approaches. The fact that in the present essay it has in large measure overtaken and overshadowed concern for apparently more well-defined and/or restricted, empirical aspects of the phenomena of control, especially movement control, is the expression and consequence of that view within the framework of the above argument.

It is, nevertheless, a problem which can be clearly identified within yet a third approach to movement control. We have called it technology-driven for reasons that will become clear in the upcoming and final sub-section devoted to a review of past research.

2.1.2 Technology-driven research

As mentioned in the introduction to the present section, there exists a third area of research in the problems of movement control and production which belongs to neither of our two previous classes of performance- or system-driven research. This third series of endeavours, being intimately related to, and dependent upon, recent (i.e. post World War II) developments in information theory and technology, it has therefore had the shortest history. In fact it was only in the early sixties that the idea of automating the operator of various manipulators was born and took shape; initially this occurred for example in the context of the demands associated with planetary probes, but quite rapidly the "fantastic potential" (Mason 1979, p. 6) of computer-controlled manipulators became evident, as it appeared that the diversity of possible applications of such tools was to be compared to nothing less than the variety of uses to which is put the most versatile tool of all: the human arm and hand. Thus while Craik (1948) had dreamed of *modelling* the human operator as an element in a control system, the new goal took the form of the apparently more ambitious prospect of *building* a mechanical or artificial operator with human-like capabilities, at least in the domain of manipulation (in its widest sense). This development paralleled and expanded upon the crystallization of the technological challenges, formulated in the middle fifties, that came to be known as Artificial Intelligence (AI). The intrinsic open-endedness of the formal background (cf. chapter 0) to this new approach to the problems of intelligence and knowledge at large (including therefore those knowledge problems specifically associated with the control of manipulators), carried over and realized into its technological instantiation, particularly in computers, probably contributed to a noticeable degree to the establishment of the field's most outstanding characteristic: the polarization of attention on information *technology*. Of course the peculiarity of Robotics research in movement, i.e. AI research bearing on the control of manipulators, as opposed for example to research in reasoning or problem-solving, stems from the fact that it involves a further concern with the physics of the manipulator itself. This added dimension, moving away from the purely symbolic concerns characteristic of most of the work in AI and closer to the 'concrete' preoccupations of the engineer, has in fact largely determined the orientation of movement research in AI in which the specific knowledge problems associated with the control of manipulators seem to be dictated and indeed defined much more forcefully by the physics of the object of control than by any other consideration.

It might be worthwhile, before proceeding further, to mention how this third, *technology-*

driven, approach in fact shares in and integrates the characteristic or prominent preoccupations of the previous two. Thus while it is evident that Robotics research does not start out with a 'given' such as a performance level in some (e.g. human) organism, it does in fact, just as performance-driven research, strive toward a theory of performance; in the same manner, while it is evident that Robotics research does not start out with a 'given' system such as a (motor) nervous system, it does in fact, just as system-driven research, strive toward a theory of the controlling system: the main point here being that in order to build a *system* capable of achieving or demonstrating a certain level of *performance*, neglect of either aspect is unwarranted by definition. On the other hand, it is clear that a solution to the manipulator control problem would at once provide a model for both system and performance since the theory of the system and the theory of performance would be indissociable, and probably indistinguishable. As to whether or not such a solution would 'apply' to biological systems and levels of performance, or indeed as to whether it is feasible at all, even though it is generally agreed that in principle it is, remains of course to be seen. Such questions would require, to be treated rigorously, that the field be far more advanced than it is now. (cf. Hildreth & Hollerbach 1985) Conversely it is certainly true that, beyond a shared grounding in the ideas of Cybernetics, system research as well as performance research has apparently provided little that is of use to the stated purposes of Robotics.

Nevertheless we believe that certain trends are already detectable within technology-driven research in movement control which allow and call for review and comment in the above sense, if only for the fact that AI research must tackle both aspects of movement simultaneously, a feature which demarcates it from 'pure' performance research or 'pure' system research.

Technology is all about making things that work and that do some work. Information technology, and in particular the developments associated with AI and Robotics, participates by definition in that endeavour; however here the kind of work (or tasks) that is aimed at and the kind of principles that are at work within the products of that technology are centered upon information. In the theory of computation (cf. chapter 0) there is a term used for characterizing that particular kind of work: effectivity (referring to effective decision procedures). This explains the overriding concern within AI and Robotics for effectivity: i.e. producing results that work. In that context, the relevance or non-relevance of Robotics research to the biological problems of control (either from a system or performance point of view) is a non-issue by definition. However we have tried to argue in the previous sub-sections that both system- and performance-driven research in movement would benefit from a greater concern for problems of production of movement, that is, with questions of effectivity. The fact that the whole theory of computation, and therefore in particular the practice of information processing which is derived directly from it, is founded upon the notion of effectivity (or decidability or computability) justifies this overriding concern; it furthermore implies that a

serious use of the information-processing metaphor, which pervades both system and performance research on movement, can only become powerful if it is carried to effective fulfillment.

The fact that a reconciliation of this requirement of effectivity with the empirical domains of both neurophysiologists and psychologists has not so far been achieved successfully (cf. Hollerbach 1982; Hildreth & Hollerbach 1985) does not indicate that it never can be; indeed the widespread adoption of the programming metaphor within these diverse research domains probably corresponds to some intimation of a possible 'shared' solution, and is probably incomprehensible outside of the assumption that it can be found. However it should not be understood here that the current physical instantiation of the theory of computation (mainly computers) offers the *only* possible avenue to the achievement of effectivity. If the brain is a computer, it *must* be an effective one (i.e. its functioning should be describable by the theory of computation); however it would, and apparently does, constitute a quite different instantiation of effectivity. Thus it should be added that whereas AI work is by definition committed to effectivity, it is also largely dedicated to a particular set of tools which happen to currently provide for its achievement. For the neurophysiologist or the psychologist this simply means that for example a *motor* programme need not (and probably should not) take the initial form of a *computer* programme, although, given that computers are *universal* computing machines, *any* effective decision procedure can be simulated by a computer programme. This should be kept in mind by neurophysiologists and psychologists in view of the above remark on the nature of the brain as computer: the 'easiest,' or at least the most 'natural,' terms of effectivity for system and performance problems do not necessarily correspond to those that are available in current information technology, nor should they be forced into their mould.

This should in fact be quite apparent from the previous chapter, where a description of an effective visual perception procedure was done along with a reference to its exact implementation and testing within existing information technology. As will be remembered, the peculiar problems of simulation associated with the helicoidal ring had little to do with the model itself; rather they all derived from the helicoidal ring's partial 'incompatibility' with respect to the above-mentioned *instantiation* of effectivity (cf also Appendix I).

So where does this leave Robotics and AI? Essentially, we are not primarily concerned with the details surrounding the specifics of their 'paradigms,' although some attention will have to be spent on them in order to be able to seize clearly where technology-driven research stands for example with respect to the typical problems of movement such as variability (see section 2.0, as well as the two preceding sub-sections). Our search here is much less constrained than it was in the previous sub-sections by the adherence to a fixed given; typically, although it might at first sight appear to be contrary to what was said earlier, AI, and to a lesser extent Robotics, expends considerable energies on theoretical issues, apparently much more so than is the case within the other ap-

proaches. Although the absence of a fixed object of enquiry (performance or system) might at first sight be considered a drawback, we believe, and will try to show, that these problems and issues are nearly exactly those we identified as being in need of theoretical attention within system- and performance-polarized efforts. Technology-driven research brings solutions, but it also underlines basic problems. Generally the solutions appear to fall quite short of the problems; however that is not unexpected, nor is it a feature particular to that field.

If within Neurophysiology one of the most natural ways to think about and explore the system is in terms of structure (macro- and micro-anatomy), and if within Psychology one of the most 'natural' ways to think about and explore performance is in terms of results, technology-driven research for its part draws heavily from the physics of the effector, a science unto itself, and one that appeared to naturally lend itself to, and provide the means with which could be achieved, the desired goals of AI in terms of both effectivity and realizability. This physics constitutes a set of concepts derived from Physics, and more specifically from Mechanics. It defines the theoretical basis upon which is erected and within which is carried out the vast majority of research in Robotics. Thus "Motor control is [still] first and foremost a mechanical problem." (Hildreth & Hollerbach 1985, p. 43)

According to Lozano-Perez (1982) the five main areas of concern in Robotics are: programming and planning, trajectory planning, compliance and force control, dynamics and feedback control, and finally mechanical design and sensors.

Considered as physical systems, natural and artificial effectors can be characterized by three types of interrelated properties: the dynamic, the static and the kinematic.

By *kinematics* is meant the position, velocity, and acceleration relationships among the links of the manipulator. *Statics* is the relationship between the force and torque that the manipulator is exerting on the environment and the forces and torques at the links. *Dynamics* is the relationship between the kinematics and the statics. (Brady et al. 1982, p. 6)

In general the equations expressing these various relationships are quite complex. This arises as a consequence of the fact that a given manipulator's local structure (joint and link types) is completely determinant of its 'global' physical properties: a given link's position (and velocity and acceleration) and force - as defined with respect to the spatial frame of reference within which the arm is situated - is dependent upon all *other* links' positions (etc.) and forces. Thus the analysis is fundamentally directed towards part-'whole' relationships. Now for the purposes of control, which is always ultimately exerted locally, i.e. at the joints, one is usually interested in the forces developed and the position (etc.) attained at the *last* link; accordingly 'global' effects can be defined as those effects resulting from the cumulative or integrated physical effects of all other components upon the

terminal link.

Much like the purposes which biological effectors serve, the tasks that robot manipulators are intended to perform are usually described in global terms, that is, in terms of what the end effector has to do. Grasping, placing, or moving an object for example are all understood with respect to the 'hand,' even though in general these actions imply the active 'support' and participation of all the other components. Since control ultimately occurs at a local level (the joints), the determination of the level of 'collaboration' of these 'secondary' components necessary for the completion of a given (global) task requires the local-to-global relationships described above to be inverted into global-to-local ones. But whereas, given local parameter values, it is quite straightforward to calculate the unique resulting set of global parameter values, the converse is unfortunately not true: starting with a set of global, or task, values, the problem of solving the equations of motion to extract the corresponding local parameters is in general very difficult. In fact it is the outstanding problem of Robotics.

One way of describing a task is by specifying a *trajectory*; a trajectory consists of "a sequence of desired positions, velocities, and accelerations of some point such as the end effector or manipulator tip." (Brady et al. 1982, p. 1) This is the first stage of the control process.

Once a trajectory has been determined, the arm has to be guided along the trajectory. At each time interval, it is easy to determine, from the position and velocity of the joints, what accelerations should be applied to the joints so that the arm will be on the trajectory at the next time interval. However, the motors which move the joints produce forces not accelerations. The controller must determine what force F_i should be applied to joint i in order to produce the required acceleration [for joint i]. (Waters 1979, p. 11)

Accelerations and forces are lawfully connected through effector dynamics. The dynamics equations of an effector system describe the relationship holding between the geometry (position, velocity, and acceleration) of the system's components and the forces developed within them. For a robot controller, as suggested above by Waters, these equations are used to determine the local force values (F_i) corresponding to a given or desired kinematic state; they supply the instantaneous commands obtaining for the link motors. The value " L ," called the Lagrangian, defined as the difference between the kinetic and the potential energies of a physical system, is used in the following form of the dynamics equations due to Lagrange:

$$\frac{d}{dt} \frac{\partial L}{\partial \dot{q}_i} - \frac{\partial L}{\partial q_i} = F_i \quad i = 1, 2, \dots, n$$

where " n " is the number of degrees of freedom (d.f.s), or equivalently the number of joints (if as

usual each joint is associated with just one d.f.), " t " is time, and " q_i " stands for the controlled variable (either a rotation or a displacement, depending on the type of joint). There exists a different but equally useful form of the dynamics equations called the Newton-Euler equations (cf. Brady et al. 1982, p. 15 ff.). Since a trajectory is specified in workspace coordinates (i.e. the frame of reference in which the arm is situated, rather than the joint coordinates to which the q_i 's are related) a further transformation must be applied to the above equations to achieve the final local values. The arm controller must compute these equations at every time interval, extracting the required force values to be fed to the link motors. Since in general this interval must be quite short (in the order of 10ms), real-time control of manipulators using the dynamics equations had for a long time been impossible due to the fact that the computation could not be carried out within this limit. However it has recently become possible (Hollerbach 1980), by exploiting recurrence relations in the equations, to simplify them and hence drastically reduce their computation time well below the critical level.

Given that robots existed before either the Lagrangian or the Newton-Euler dynamics equations could be used by controllers, the question arises as to the mode of control that obtained for them. In that respect the following is illuminating:

The most common forms of control in use today require solution of inverse kinematics in real time in order to determine commanded joint angles from trajectories described in workspace-oriented coordinates. Independent joint position servos [] are then applied to each joint. (Johnson 1982, p. 145)

That is, it is *position*-based at *all* levels of control: the global level of the manipulator kinematics and command trajectory, and the local level of execution at the joints through independent position servos. (A position servo is one in which the control variable is a position.)

As is well known, ignoring velocities and accelerations produces movements of a lesser quality, where for example more or less abrupt changes of position, and of velocity, are not smoothed out. It also of course presents major difficulties in applications that naturally require precise speed control. Furthermore, contrary to a controller based on effector dynamics, independent joint servoing considers the effect of other joints as a source of error rather than as something to be possibly exploited, or at least recognized as a potential source of refinement of control. Of course even in a 'dynamics' controller, feedback is still necessary to the extent in which the mathematical model, even though vastly more sophisticated, is also an idealization of the situation; it still falls short of the reality and hence is open to error. For our purposes however, following upon the discussion at the end of the previous sub-section, we merely wish to point out again how a performance derived from the use of the principle of feedback control is, so to speak, only as good

as the controlled signals, and that consequently the main issues bear upon their nature and derivation rather than on the adequacy of the servo principle itself.

But what are these main issues? In describing the overall theoretical background to robotics research we have so far hinted only at a few. Chief among them, as noted above, is the achievement of computational efficiency, which sometimes seems to override all other considerations. In fact some of the major developments that occurred in the field in its expanding phase in the seventies (Whitney 1972; Albus 1975; Raibert 1978) were directly motivated by the search for ways to deal with the intractable and explosive complexity of computing in real time the control variables and the inversion of the kinematics and dynamics equations. Thus the problem was not one of conceptualizing control itself but of conceptualizing efficient means for carrying out the calculations implied by the existing physical model of the control situation; this problem of computational efficiency, it should be noted, is only indirectly tied with the fact that the model is a physical one. It derives from the formulation of the model rather than the model itself. On the other hand the existence of very efficient means to 'compute the model' is likely to play only a minor rôle in the more general issue of the adequacy of the model itself in providing a framework for building skill and dexterity into robots.

Given that technology-driven research relies nearly exclusively on the physical model, most of the issues in robotics and AI are also derived from that model. Thus problems of trajectory planning and compliant motion (i.e. movement which must comply to physical environmental constraints such as a table top or a wall), and even more so questions pertaining to design, are all directly related to, or 'weighed' against, this fundamental understanding of control. However it should not be concluded from this that physical theory provides an integrated view of the control situation. Although it is evident, for example, that the dynamics equations involve both geometric (position, velocity, acceleration) and physical (force) variables or dimensions, it does not provide a means for 'extracting' trajectories. This is because dynamics covers all possible trajectories, in a manner which does not allow one to derive a 'natural' method for trajectory planning, not because it has nothing to do with the pertinent variables. We shall return explicitly to this question in section 2.2. Briefly, as things now stand trajectory planning is done independently, and then combined with the dynamics equations; the converse, as just noted, namely finding ways of incorporating dynamics into trajectory planning, is still "an interesting open research area." (Brady 1982, p. 223)

This should be kept clearly in mind, especially when considering apparently critical comments like the following:

One of the distinguishing features of a robot system is that it be capable of physically processing the external real world in addition to its information-processing abilities. It is, therefore, more than just a realisation of a Turing Machine, because

we must take account of the range of *physical* possibilities as well as the range of logical possibilities in its design and operation. In a robot system we are more interested in what it is possible to *do*, and in what it is possible to *make*, rather than just in what it is possible to *compute*. (Rooney 1984, p. 193)

It is certainly true that the physical 'presence' of a robot or of a biological effector is impressive and inalienable. However it should not, we believe, be concluded from this that the physical and the information-processing properties of a movement system have to be segregated into independent analyses. Indeed, the whole of technology-driven research efforts up till now rests upon attempts to transform what Rooney calls physical processing into information processing. It might indeed, *qua* knowledge, bear much more closely upon the physical than, say, the visual world, but it is still computed. One can certainly be impressed by the *global* energetic manifestations attendant upon movement production, but it should be remembered that the energetic exchange consequent upon the activation of a single motor unit is probably not very much larger than the energetic exchange occurring at the level of a retinal field. It is far from clear or apparent that the nervous system proceeds like the researcher in AI, that is, that it computes commands according to a given physical (in the sense of Physics) knowledge base; it is just as likely that such a knowledge base, in whatever form it exists, is derived and consolidated in large measure, if not exclusively, from motor information-processing. Thus, attempting to infer the nature of motor information-processing from a study of the properties of its related physical results might not be as powerful a strategy as can appear at first sight. In any case, at least as concerns the relationship between an information task such as trajectory planning and the energetic interactions implied by the arm dynamics, it has not so far been achieved even for a single case. Roughly speaking this means that although we can make use of the geometric-physical properties of manipulators in trajectory planning, we have not yet found ways of exploiting their energetic-physical properties toward the same end. Moreover it should be noted that there does not yet exist a general theory of kinematics (Rooney 1984, p. 222), upon which, of course, trajectory planning relies directly.

It is certainly true that theoretical allowance must be made for the "range of physical possibilities" inherent to specific robot arm designs, just as attention should be paid to the physical design of retinas for example (cf. section 2.0); however the major problems do not appear to reside in that domain, since it is relatively easy to satisfy the minimal structural requirements needed in the attainment of, say, a human level of performance. But on that criterion, control of the ideal arm is completely beyond any extant computational tools (*ibid.*, p. 221). Robot arm design consists mainly in the art of getting the most (in terms of physical possibilities) out of the least (in terms of the information-processing needed to achieve control - at least with respect to the physico-geometrical model of control).

However even though in robotics research, as in vision research, it is rather straightforward to state the basic structural requirements necessary for the achievement of a human level of performance, it is not at all evident what kind of attendant information-processing it calls for. Its discovery admittedly constitutes the basic goal of AI; in large part it defines the field. It therefore seemed logical that the first step in "the creation of machine dexterity comparable to the dexterity exhibited by the human arm" (Flatau 1973, p. 1) should be a study of the question of the extent to which arm design can allow for, facilitate, or hinder these prospects. However this is a loaded question, in that a definite answer to it could only be stated in the light of a highly dextrous machine, that is, it presumes that the problem which it is supposed to help solve has already been solved, or at least *can* be solved through the geometric-physical model. By providing at least a context within which control could be conceptualized, the geometric-physical model has, however, partly satisfied this apparently circular requirement. Yet in the end the theory of dexterity has remained largely "undefined" (ibid.). Turning to system, or even performance, studies in search of a theory of dexterity that could be useful in that context proved, as is well known, to be a deceptive enterprise; as for the lost dream of bionics, it seems to have survived only as a weak echo in the form of introductory remarks about the fact that "It seems clear that the present mathematical formalisms are in deep trouble when addressing the type of mechanical control problems which are obviously trivial for the brain of the tiniest bird or rodent." (Albus 1975, p. 221) What is not as evident however is why, but unfortunately AI researchers have apparently chosen to expend little energy in finding out.

Having described and commented in a general manner upon what we believe to be the essential problems and methods encountered in technology-driven research on movement, we will presently examine their relevance in the context of our own understanding of the problems of movement production. Insofar as we share with these endeavours a common formal background, and to the extent in which the context of our own efforts (described in detail in the next section) is closer to those found within AI than within either of those associated with performance or system research, it will come as no surprise that a larger degree of concordance can be found with them. However these 'meeting points' will be seen to occur nearly exclusively at the level of the fundamental problem, and much less so at the level of methods or solutions. Thus, while in the previous sub-sections we tried to show that motor control was basically a problem of production, it should be clear that this view can be argued to constitute the starting point of technology-driven research. It should also be clear that, as outlined above, the basic task of control defined in AI involves degrouping.

A few pages ago we talked about how in Robotics "the analysis is fundamentally directed toward part-'whole' relationships," and how the latter naturally arose as a consequence of the cumulative and interrelated physical influences operating among the various links and joints of the manipulator. It was further suggested that in this analysis, following upon the requirements for the

vast majority of applications, attention was usually focused on the fate of the *last* link in the chain, i.e. the 'hand.' The notion of trajectory was then introduced as a way of describing specific tasks. Since a trajectory usually corresponds to the path to be followed by the 'hand' engaged in a certain task, and given that the 'hand' like all other links is repository to 'whole' effects, it is natural to speak of tasks (or trajectories) as having the character of 'wholeness,' or globality. So while the physical nature of the manipulator gives rise to an analysis which must take into account the whole (i.e. the sum of the effects of the joints), and while the task is also naturally defined and expressed as a 'global' procedure, the problem of 'forcing' a manipulator to do a specific task must involve the transformation of these global descriptions into more local ones, where, as mentioned previously, control is ultimately exercised. This problem should be recognized as an instance of the problem of degrouping.

In a sense it was automatically realized and recognized as the major theoretical problem in the initial stages of research in the fabrication of self-sufficient auto-controlling manipulators. The kinematics and physical models to which AI turned in its search for computational means of achieving these goals could even be argued to be directly based upon a grouping-degrouping dynamics analysis. Nevertheless, at this point some difficult and subtle questions arise. They concern not the recognition of the problem's key position but rather the mode in which it was approached. The fact that physical analyses were, to some arguable extent, also tied up with these issues, coupled to the fact that they were the *only* type available, could be adduced in favour of the strong initial attraction of this type of approach. Now, following Albus' comment above, it was nearly equally soon realized that "the present mathematical formalisms are in deep trouble" when it came down to comparing the products of artificial controllers with those of biological ones. Yet the belief that the problems were not *different* in the two cases, and the conviction that even though present technological performances fell quite short of natural ones, they were not *beyond* their formalisms' grasp, both remained strong enough for AI researchers to persist in them.

It will be remembered, from section 2.0 that we postulated the existence of global entities whose nature was such that not only did they not bear directly upon the *local* state of the periphery, but could not effectively do so in view of the purpose of control, which purpose was defined initially as involving the specification and realization of peripheral states (and eventually change of state) as *wholes*. Now the geometric-physical model readily allows for one description of global entities, a description that is, however, directly derived from a 'grouping' or integrative procedure, whereby each local element, specifically the links and their connections, are synthesized into a 'whole' description. Once this is achieved, as suggested earlier, control proceeds through a 'degrouping' process whereby the present global entity is 'solved for,' or brought down to the level of, or transformed into, the proper series of local entities (commands for the individual joints). This is

where the majority of computational difficulties arises, since the process usually calls for demanding computations such as matrix inversion to be carried out inside very strict time limits. However that is not to our understanding the source of the major difficulty. Rather we want to suggest that one cause of this difficulty, which, it may be noted, is already great even at the present low levels of performance, might be due to this apparently unquestioned means of deriving the initial global entity or description. It is certainly not our intention to question the fact that global entities are at the root of the control process, but rather to propose that their nature need not necessarily be based upon or derived from *physical* models of globality. It has of course never been argued, within AI or elsewhere, that there could be no other possible grounds for deriving them than through a 'physical' grouping process, but the fact that it has been adopted as a method so readily leaves one few alternatives but to think that the matter was considered evident and in no need of particular attention or demonstration.

Be that as it may for now, and coming back to the problem of the nature of the global entity, we will conclude for now on a brief statement of a possible alternative to the above-described method. The present suggestion holds that, in a sense, one cannot and should not start out with a very definite idea on the nature of the global entity; rather one should strive toward its achievement following a process of gradual construction, a process which attempts to build more complex levels of control upon simpler ones rather than in a wholesale manner, and furthermore is much more strongly guided by the highest extant level of control in the development of the next one than by a previously or initially defined 'final' level. Interestingly, this type of approach blurs somewhat the distinction between perceptual and movement knowledge processes, at least with respect to this kind of study. It is furthermore, if the parallel be allowed, reminiscent of the open-endedness of natural evolution.

In section 2.2 we shall pick up again on this trail in the more explicit context of our own overtures in the stated direction, and our partial attempts at resolution. It should nevertheless be clear at this point where we conceive the boundaries of technology-driven research to lie, and why we have not endeavoured to pursue our search within that framework. Indeed AI's dedication to one particular type of solution to the problem of degrouping (solutions which happen to be, at this stage at least, quite removed from those which are supposedly sought both within performance- and system-driven research), along with its disregard for the particular challenges - and more specifically of the quest for biological relevance - associated with them (even though reference is sometimes made to neurophysiological fact or to some products of introspection), appears, especially in view of its much more 'explicit' recognition of what we consider to be the fundamental theoretical problem of control than is evidenced in either of the two previously studied groups of efforts, to be somewhat paradoxical.

2.1.3 Conclusion

Although the Gestaltists, early in this century, were the first to insist and argue, mainly in the context of visual perception, for the existence and critical functional importance of global entities in visual processing, it was not until the advent of the theory of computation and the development of associated information technology that the arguments could be turned into a much more 'concrete' form (cf. Lamontagne & Howe 1980), where, being now convinced of these arguments, the matter became one of explicitly formulating *working* models of the processes that give rise to these global entities. It is in the context of these new challenges, namely those of actually deriving percepts and producing movement, that the functional concepts of grouping and degrouping took shape. Although the fundamental filiation with the epistemological domain, that is, the assumption that periphery-bound processes *are* knowledge processes, has remained unquestioned, it is nevertheless the case that concern with the problems associated with meeting these theoretical challenges has rather quickly driven concern with the purely epistemological issues from the forefront.

These problems have, indeed, turned out to be weighty. We have seen that while neither performance- nor system-driven research efforts have, with but very few exceptions, even attempted to formulate, let alone solve them, technology-driven research for its part has opted to concentrate on and pursue a single line of approach. Beyond the necessity of recognizing the problem of degrouping as central to any effort either at unveiling nature's solutions to it or at building a machine capable of a similar level of performance, one of course can be sensitive to the peculiar constraints and difficulties associated with these research efforts, whose demands render the absence of explicit recognition and the apparent reticence to embark upon a theoretical quest quite understandable. For our part, we would tentatively, and speculatively, formulate these 'reasons' in reference to the still apparently very great gap separating the empirical and theoretical domains. By this is simply meant that for some reason it has seemed most difficult to 'derive' a general (but precise and functional) theoretical understanding of movement phenomena, with respect to a given system or to a given level of performance, and, conversely, to 'produce' relevant and powerful empirical consequences from extant theory. The essence of the matter seems to reside in the question of formulating the problems in a way that satisfies both aspects of the research endeavour simultaneously. This can be seen most clearly with respect to technology-driven's theory, which, to our knowledge, has not generated a single hypothesis regarding the possible nature of natural systems' mode of production or even provided indications of promising directions in which to seek. Yet it is presumably the most precise and powerful extant theory of production. Admittedly, neither have we seen the latter's search steered in any but an anecdotal fashion by outcomes stemming of performance- or system-driven research.

Even though we have not as yet been able to bring our own research efforts to a level where a reconciliation between these three approaches or a resolution of their diverse demands could be afforded, the fact that three apparently quite different avenues to movement production can equally be argued to revolve about a unique problem, namely that of degrouping, lends some credibility to the notion that in its study and eventual solution might well lie the way to important if not decisive progress.

Even though it is fundamentally theoretical in nature, the problem of degrouping cannot of course be studied in a vacuum. Having first had to define a context for its study, we arrived at a 'paradigm' in which are brought together some central elements from the three domains of research. From performance-driven research we have chosen to study a humanoid arm configuration and have been particularly interested in the problem of 'shape;' technology-driven research's impetus is clearly underscored in our pursuit of effectivity and in what amounted to be a fundamental concern with kinematics; as to system-driven research, whose 'input' is less well-defined, let it merely be said that the belief that the brain is effective (in the computational sense), and hence the (corollary) search for the key to some of its powerful yet simple 'formulae,' has coloured and restrained in a very profound way the extent to which, on computational grounds, we were willing to go in order to preserve a semblance, as well as a hope, of biological relevance.

2.2 Elements for a working model

The focus of attention now shifts towards the computational realizability of degrouping, that is, the construction of a movement production system.

This third and last section of chapter 2 is devoted [1] to an extended discussion of degrouping in the context of the more specific problem of controlling an arm, and [2] to an account of some proposed degrouping structures for controlling this effector. It expands and further develops many of the issues discussed in the preceding chapters, bringing them to bear on our specific instantiation of the problem of control. Following closely the approach to visual perception presented in chapter 1, it is based on the definition of a virtual interface, namely an arm, onto which are connected the control structures. Although the interface as a whole, that is the arm, is the ultimate and single object of control, it will be seen that three distinct levels of control arise and that to each is associated an 'object' of control: this follows from the nature of the interface, which is made up of 'independent' parts. It is defined in a way that is argued to capture the essential features of a particular control context belonging to the biological class of effectors; more specifically it is patterned after the human arm. From the point of view of its structural components (upper arm, forearm, and 'hand') it is quite similar to the actual human arm, although from the functional point of view the conceptual effector is simplified, the loss occurring mainly as a result of reductions at

the shoulder and hand; but this is not really critical to the main issue at stake, for it can be argued that if the problem of controlling such a simplified effector were solved then it would carry us some way towards a solution of the more complex case, just as in vision there is no a priori necessity for dealing with a retina of human size in order to proceed with the study of perceptual problems at a human level of performance. However there is a lower limit to the degree of reduction allowed, which in the case of effector systems, as mentioned in sub-section 2.1.2, is derived from a geometrical argument, and which in the case of visual systems stems for example from consideration of the desired resolving power (see Appendix I, section 3).

Concerning the nature of the control model, it will also be apparent, as for the visual line detection model discussed in the previous chapter, that the main line of attack was *primarily* directed towards the achievement of computational efficiency, a necessary, although not sufficient, prerequisite for a theoretical effort aspiring to psychological relevance and neurophysiological plausibility (cf. section 2.1; Lamontagne & Beausoleil 1982). Given a stricter computational grasp of the neuronal substratum, or computationally more precise indications on the psychological nature of movement production, one might expect computational models of control aiming for these goals to be much more attentive to whatever Psychology and Neurophysiology have to offer, and in fact to proceed from their respective contributions in the task of pushing out even further the computational frontiers of movement control. But in general, as we saw, given the epistemological and formal assumptions underlying the computational approach, neither, singly or collectively, presently affords sufficient basis for such a path to have been followed by us, these assumptions being at best only partially met within most of these efforts. This implies that one has a nearly completely free hand in the quest for computational efficiency, the impositions stemming from performance and system research, to which we shall nevertheless sometimes refer as we go along, having more to do with the definition of the problem than with the nature of its solution. This is particularly clear with regard to the interface as well as for the level of performance to be aimed at by the system.

If, on the one hand, our most general conception of psychological phenomena puts them in the category of knowledge processes, our most general conception of neurophysiological phenomena, on the other hand, ascribes to them a supporting rôle, the conjunction of these two types of processes being thought to 'occur' or be expressible in computational terms. Now one of the leading themes within modern Neurophysiology, starting with Cajal's discoveries and persisting throughout the century right up to Hubel & Wiesel's (1962) findings, is that the *micro*-structure or *micro*-anatomy of certain, if not all, of the brain's sectors is as important (because it is well-defined) as its *macro*-structure or *macro*-anatomy. We of course tend to interpret this importance, at *both* levels, as indicative of or related to as yet undetermined computational tasks. The difficulty, from the point of view of an empirical investigation of many of these tasks that go beyond the evident

physical linkup with a perceptual or motor periphery, and even with these in fact, is that most of them supposedly depend, for their input, upon the output of previous structures whose function is unknown, and generate outputs to other structures whose purpose is equally unknown. Even though our foremost goal was not to produce a neurophysiologically 'tight' theory of movement control, in which case for example we would have worked with an interface much closer to the 'real' human arm instead of with our simplified effector, we retained as fundamental the assumption, derived directly from the above-mentioned interpretation (and that in fact could also be argued to hold within the framework of more classical neurophysiological theory) that the neuronal *structures* themselves, alongside their constitutive *neuronal elements*, partake of the computational tasks that characterize them. The implication for theory is that computational models need not necessarily involve only calculations but can also exploit wiring as valid and legitimate computational tool. Following in that respect the style of the helicoidal ring model, the elements of solution proposed in the following sections all partly revolve about the *structural expression of computational principles* deemed pertinent to the task of degrouping. There is of course no a priori *necessity* for theory in general to be bound by this, as is demonstrated by many efforts in mainstream AI, and particularly Robotics (see sub-section 2.1.2), that are of course quite deeply committed to computational efficiency; however we have argued that the fundamental problem faced within Robotics and Neurophysiology is the same one (degrouping), and on that view the question of whether or not the nervous system *does* carry out the task, and how, becomes subsidiary to an even more pressing problem, that of how it *could* be done.

The overall idea is simply to arrange for the possibility of producing at one level a *variety* of events that have repercussions at the next one down, where again the same can happen, and so on until the level of the interface is finally reached, at which point the effect dies out into an ultimate energetic manifestation. The general idea, following the argument in section 2.0, is that for *one* movement, successive levels give rise to *increasing* numbers of events while the number of objects controlled by them *decreases*. To see this, the nature of and the relationship between events, objects, and levels must be clearly delineated. Briefly, events correspond to signals in the system while objects correspond to whatever part or parts of the effector system are directly and ultimately affected by a given event or set of events; levels can be essentially described with respect to the objects, potential and actual, to which are tied the events that occur within them. Given a hierarchy of levels, it therefore follows that, for example, the highest level will equally have as potential object, through the intervening levels, the object of the lowest level; but as a rule we usually refer to a level's object as the most immediate one on which the control structure exercises its effect. The usefulness of such a wide definition stems from the fact that it stresses the embedded nature of successive levels of control, a feature, as argued in section 2.0, which appears inescapable in view

of the nature of the interface.

The above statement in fact concerns the whole control system's functioning, its basic thrust; it expresses the relation holding between successive levels, the very essence of degrouping lying in the fact that they are organized in the manner described. Levels are where things are happening, computationally speaking. We have, naturally enough, analytically subdivided the effector system into components to which are ascribed certain precise control subsystems. Now such a functional partitioning of the effector, in which the components targeted by the control system correspond to identifiable parts of the effector, can be unreceivable in some instances. This was the case for the idea that the muscle is a basic control object for the nervous system. But in our view the debunking of that hypothesis cannot be interpreted as a refutation of the fundamental nature of the problem: there *is* a multiplicity of objects at the periphery, and some indication of the existence of degrouping *has* been identified at the level of muscle units, the basic neuromuscular constituents (cf. Henneman 1981). These are not, strictly speaking, effector components; rather they are *parts* of the effector components.

For the case of the arm we will therefore be seeking to construct a system where a *single* highest-level event, pertaining to the effector as a whole and hence potentially addressing *all* subsequently identifiable objects of control, impacts on the next level by provoking a *larger number* of events (equal to or less than the number of components of the effector, which in most cases is larger than the one effector) that *each* in turn address a lesser group of objects, defined at the final level, where events now address but a *single* object.

Each level therefore will work according to the same principle: it will be built to accept or understand a certain type of event, bearing upon a specifiable immediate object, and to respond to it by generating a larger number of new events, each bearing upon reduced or lesser objects. It should be noted carefully that there is only one set of 'real' objects in the system, namely those defined by the interface: all other higher-level ones are *imposed* on them, and only thus identified: every level *creates* its object, except at the periphery, where it is given. This is another way of explaining the necessity of levels of control: achieving control of an effector system that is only really defined by a set of quasi-independent parts logically seems possible only through a prior association of some of these parts into an object (something to which a control structure is attached) that in turn is constitutive of further objects, etc. But these intermediaries are virtual only and are to be understood within the framework of the computational task at hand: therefore the objects, immediate and potential, which we talked about above, could just as well be understood as indicative of the overall purport of computations within the control model: the level of the 'sub-effector' (see next subsection) for example can be seen from the point of view of the structure of the effector itself, as its name suggest, but more importantly it can be seen as a conjectured efficient functional byway along

the road from intention to ultimate realization, which only really occurs at the last level, namely that of degree-of-freedom in our model.

What has been said so far concerning the system's overall task and functional mode does not in fact correspond to the way it was actually approached. Rather, the search for successive levels of control followed a bottom-up direction, starting with a control structure for the most primitive one, the basic constituents of the interface (the degrees-of-freedom), and proceeded from there towards the most global level, involving control of the whole effector. This was forced as a result of our choice of an embedded structure for the successive levels of control, a choice that in turn had a very important consequence: it caused *the problems attached to a particular level to be defined, and their solutions sought, in terms of the possibilities opened up by the previous one, and within the confines of the restrictions imposed by it*. For instance the second level of control, that of the sub-effector, is *defined* by and *built* upon the degree-of-freedom control structure. Accordingly, and following what was mentioned earlier, the original definition of the interface *already* comprised some of these general elements of global control design, and so we *knew* from the start that control structures for 'sub-effectors' would be highly dependent upon those for degrees-of-freedom. This implied, literally, that work had to be carried out at many levels simultaneously, results at one depth of control having more or less clear implications for the nature of processing at all others. In our presentation of the control models we will roughly be following this pattern of development, starting with a description of the structural and functional components of the effector, and then proceeding to expose the related control structures.

We would like to stress again that this internalization of the problem of control, the fact that it is essentially redefined at each new level as a function of the *existing* levels of control (which at bottom is simply the interfacial components) rather than as a function of conditions more external to the system, appears to us to lie at the heart of periphery-bound knowledge processes and seems capable of explaining, at least in principle, much of its power and diversity. It is of little consequence in that respect that the 'direction' of control is the inverse of the 'direction' of perception, for from the point of view of the beholder and the controller, in both instances it seems true that the objects of experience lie at the top of a large construction whose productions are very far removed indeed from either periphery's objects, that nonetheless are directly participant in their genesis and therefore profoundly determinant for the nature of experience.

2.2.0 Degrouping

In section 2.0 we presented our fundamental thesis according to which the problem of movement control, viewed as a *production* process tied to an *interface*, resolves into the problem of *degrouping*. Our initial discussion, necessarily very general, nevertheless led to the identification of

some features of this process. In the first instance it was presented as involving the specification of *global entities or events* otherwise definable by *sets of local interfacial states or modifications thereof*. From this follows the idea that a control system must be hierarchical, i.e. must involve levels, since the specification of a diversity of global entities cannot be achieved *directly* from the local peripheral entities (cf. section 2.0). Indeed the whole business of degrouping (as well as grouping) can be described in terms of the problem of finding single *equivalent* descriptions of *sets* of peripheral events. (From the point of view of experience, a similar understanding proceeds from the observation that its content, both visual and motor-intentional, differs markedly from the sets of events that are otherwise believed to be the source of visual experience and the target of movement intentions, an experiential distance and discrepancy that is reflected in, and consistent with, the above more formal statement.) Now every level will specify its own object/s, each simpler than the preceding one's, until the final interfacial objects are reached. Furthermore, each level is called upon to produce its own set of lower-level objects.

While, as we shall shortly argue, it is conceivable, and much simpler to imagine, that to each level attaches but a *single* object, the fact that, at least at the highest level, one recognizes a *multiplicity* or *variation* of intentions must mean that more than that is involved. And indeed such is no doubt the case. However we already have the elements for a possible solution. Degrouping, as just mentioned, involves the process of 'multiplying' the number of objects as events proceed from one level to the next; however little has so far been said about the specification process involved in the generation of these objects. But recall from section 2.0 that this specification process rests upon *characterization*, whereby certain attributes or values are laid down by the system. Now as we saw in chapter 1, at least two types of characterizations are involved in the specification of an object, one that essentially defines it (called the *primary* characterization) and the other which specifies its particularities, i.e. the residual characteristics that are allowed within the essential one (called the *secondary* characterization); thus we saw that lines, as objects, are primarily defined according to linearity (through the linear topological arrangements of interfacial events) and secondarily according to orientation, etc., namely those characteristics of lines that are not 'touched' by the primary, or equivalently those variable features of the object that are allowed within (or by), or that follow upon, the essential one. In general, an object will be specified by a single primary, and an indefinite number of secondaries.

However a further characterization task will need to be called upon in movement production, one that does not concern the identification of the object but rather the *specification of motion* itself. Of course motion characterization is not a distinct process from that of object specification, since it will target those very same objects specified at each level. A motion characterization was defined in Lamontagne & Howe (1980) as *any change in a multi-valued characterization*, those

changes themselves *also* being specified according to a primary or essential feature, as well as according to secondary or non-essential features. Thus, for instance, a *linear movement* implies the specification not only of the *primary* linear nature of the movement, but it also requires the specification of its *secondary* features, namely direction and amplitude. (We do not in general ascribe a positional characterization to a movement.) This whole question of *the characterization of the transformation* however will only again arise in sub-section 2.2.1.1, where, having an *object* of control, the question of the nature or characterization of motion directives naturally resurfaces. For the present it is sufficient to understand that a motion characterization bears upon multi-valued, or secondary, characterizations of the objects, i.e. upon those that can 'support' a transformation.

The overall task of degrouping so far therefore involves at least *two* types of concern: [1] the specification of objects: this will be seen to be the problem of the representation of the object of control as a single global entity; [2] the specification of a change or modification of that object: this is the problem of motion characterization, attaching to *every* level's particular object. (In general from now on we will reserve the term MOVEMENT for the empirical or peripheral product of degrouping, while the term MOTION will correspond to those aspects of the control structures immediately concerned with the events that lead to movement.)

It will be observed that while this sort of distinction is completely absent in performance-driven as well as system-driven research, it is in fact only partly so in technology-driven research. Recall for example that the basic control scheme in Robotics is one where a positional characterization is first achieved for the effector as a whole, whereupon it is associated with a series of transformations of position (a path). However in general the path itself, although it has to be translated into variations of positions, is not *derived* from, or even *based* upon, the positional characterization - rather it is typically *independently* computed and then 'fed' into the controller; it was also seen that even though this could be made to work in the context of positional information, the question of generalizing the scheme to one using the dynamics characterizations of the effector was still an open one.

Following the structure of our discussion of grouping in section 1.0, we shall now be organizing the discussion of degrouping around two sets of issues, namely the strategic and the tactical issues: in the first place (under sub-section 2.2.0.0) the fundamental purpose or 'character' of a movement control system will be considered, while in the second place (under sub-section 2.2.0.1) the practical problems connected with the implementation of this purpose, that is, with the elaboration of a working movement production system, will be outlined.

2.2.0.0 Strategic views

From the point of view of day to day experience, the purpose of voluntary motor control is to

act. Actions presumably correspond to intentions, and these intentions are rarely, if ever, purely motor in nature (although they apparently can be): that is, intentions do not generally pertain to the effector system per se but rather to things that can be achieved through the *collective* movements of its parts. Sometimes, depending on the nature of the intention, different parts of the body (i.e. different effectors) will be called into play; let us however restrict the effector system to one single component, namely an arm, so that *all* intentions will automatically involve it. (We are therefore only dealing with a subset of all possible movements.) So far we have insisted nearly exclusively on the problem that arises as a result of the obvious discrepancy between the singleness (and globality) of any particular intention and the fact that in order for it to be satisfied, it had to be divided or split up into a set of local intentions or commands. We insisted much less on the now even more obvious, yet apparently contradictory, fact that movement commands are also *almost always* ultimately targeted at the *same* set or object, namely the *collection* of effector (i.e. arm) components. From that point of view it is clear that the specification of objects could not be the essential business of movement control, as there are potentially relatively few of them in the general case, and but a single one in the above; rather, the power of the system seems to lie in its ability to specify a very large number of sets of transformations for a *default* set of components. In order to see the functional implications of this tentative observation, we shall first have to step back some distance and consider anew the 'meaning' of control, and in particular movement control.

In section 2.0 it was proposed that one of the key aspects of interfacial events is variability; by this was meant the formal fact that, given the nature of the interface, a *large-potential* is associated with the periphery, a potential which a control system could be thought of as *exploiting*: i.e. the system produces a multiplicity of local events which, taken collectively, amount to something other, or more, than a particular collection of local or singular events, that otherness in fact being describable in terms of the satisfaction or achievement of a goal. As a first and very general rendering of the 'meaning' of peripheral control, this suggested that the control system was there to provide access to a *representation* of peripheral variability in terms of, or as a function of, a given goal. Turning now to the slightly more specific case of a *movement* control system, we could say, following a similar line of reasoning, that the system's goals can only be achieved in the light of a *representation of the movement potential* associated with the periphery, a representation of a more specific type of peripheral variability which we are also supposing the system itself responsible for providing. At a third, and even more precise level, if we now suppose that the interface is an effector, then providing a representation of the movement potential associated with the periphery in effect becomes the task of *representing the movement potential associated with a particular effector*.

At this point we would argue that representing the movement potential associated with a particular effector in fact presupposes the existence of, or that it can only be accomplished through,

a representation of the effector itself. Pure movement can exist of course, but only as an abstraction; for controlling the movement of a concrete object such as an effector, that abstraction must be intimately tied with another one, namely the representation of the effector. Assuming that movement is a property of an object, the above implication in effect becomes inescapable. This property of course is more variable than other properties of the same object, for instance its size or weight; this is why we can for instance speak of the movement potential of an object, and less so of its weight potential. Furthermore it depends upon two other fundamental features, namely position and time: more precisely, movement potential refers to the notion that both of these properties being indeterminate (i.e. they describe a *potential*) they can nevertheless be combined and used to define, as well as determine, a new property (i.e. they describe a *movement potential*). But these more primitive features themselves also attach to objects, objects which, in a knowledge system, are *also* defined according to features, usually of a fixed nature. In chapter 1 these attributes were called primary characterizations, while the more variable attributes were called secondary characterizations. For an effector, it should be evident that position is a primary feature for neither the interfacial components nor for the whole effector; some other basis will have to be found to essentially describe, or represent, the fixed (or, borrowing a term from performance-driven research, *invariant*) property of the effector, those which make of it an object within the control system.

However we have a problem. Not only can position not serve as a primary feature of the effector, but even if it could, the matter of representing the overall or global position of the effector has yet to be solved. This is because, as we well know, the effector is made up of a set of components, each of which can have its own, local, positional value. Thus even though from the strategic point of view we have arrived at the conclusion that the movement control system should have only one object, namely the effector, it still remains to be seen how that single object will be represented within the system: this of course is a fundamental tactical problem. Furthermore, given, on the one hand, that position cannot act as primary representation of the effector, but given also, on the other hand, the fact that position is conceptually very closely related to movement, it should now be apparent that whatever form this representation takes, it should be one that is also functionally closely related to the representation of *global* position.

In this context we might pause to make the following two sets of brief comments. It will be recalled that in chapter 1 we presented the single visual object scheme (SVOS) according to which a visual system always deals with but a single visual entity; and in particular this had to be the case for the perception of movement. The computational arguments for this hypothesis (cf. Lamontagne & Howe 1980) are quite extensive; the hypothesis probably appears somewhat surprising in view of our strong intuition, based on visual experience no doubt, that the visual system deals with *many* different objects. As opposed to this, the proposition that the movement system only deals with a

single object, the effector, seems obvious (at least when the effector system is reduced to one arm). However just as, initially the Gestaltists, and then again the whole AI effort in vision, brought home the realization that visual objects were not given as such but had to be constructed by the system, and that to suppose the existence of (one or many) visual objects was indeed but a first step in the elaboration of a working theory of vision, just so here it must be realized that even though we are supposing that there is only one object of control it is not simply given as such to the system. Indeed little else is given to the system apart from a functional access to *all parts* of the effector: but from the point of view of the system "all parts" cannot be the same thing as "a single entity" encompassing the whole of them, an entity furthermore that allows for the production of significant movements. Thus even though by reducing the number of objects of control to a single one the problem is simplified, it is far from disappearing, since that object still has to be defined. Our whole present discussion of the uniqueness of the object of control was in fact meant as a prelude to this statement; although it can appear too simple, it is deceptively so since not *just any* representation of the effector can be an interesting candidate for a movement system. We have seen above for instance that a positional representation per se cannot do. That is, the property that will bind the effector components together in a single entity should be even more primitive, in the sense that, whatever it is, it will have to *remain true* across or within all positions of the effector. It will be recalled from chapter 1 how the helicoidal ring model effectively achieved such a result for the visual objects "lines:" the primary, or invariant, feature of linearity was 'conserved' within all variations of position, length, and orientation.

In fact, the problem which led to the SVOS was *also* a problem of movement, namely movement perception. It states and argues that a final representation of movement (at a human level of performance) can *only* be achieved if a single visual object is *permanently* achieved by the system. (However it does not preclude the use of motion computation as a grouping mechanism for the derivation or definition of that object.) Hence the computational weight, or essence, of the visual system, even in the case of movement perception, must lie in *non-movement* concerns, in the establishment of the single visual object which, being extant, can *then* lend itself to final movement computation. Our analysis of the problem of movement production has led to a similar statement, namely that, in *developing* a system, concern for movement per se must be subordinated to the definition of the single object of control. In that context, it should carefully be noted that we are *not* proposing that a movement system engages in grouping: the fact is that the visual system provides, even if it is always dealing with but a single one, access to *many* objects, while the movement system does provide access to but a *single* one (in our restricted case, of a single effector; however if all effectors are considered, then that single object presumably becomes the whole body). For that reason we prefer to use the term "representation of the effector" instead of the potentially

more confusing idea of "grouping the effector," i.e. that this must be achieved through a grouping process. The essence of a movement system is and remains degrouping, and the computational weight of the system must be seen as lying not with the creation of effector representation (which in principle, if it is well done, need only be done once, and so could be hard-wired into the system), but with the problem of proceeding from this global representation, along with a (presumably also global) movement representation, to a 'solution' of the correct local movement commands implied by it.

Our second set of comments bears upon the relationship between what we have so far advanced on the necessity of a representation of the effector as a whole and some elements of the problem of movement control we encountered in our discussion of performance- and technology-driven research (cf. sub-sections 2.1.0 & 2.1.2). Even though in performance research there is little explicit concern for this problem, there is nevertheless concern for the problem of kinæsthesia, or the sensation of movement, which we believe to be closely related to it. Our thesis, formulated in kinæsthetic terms, is simply that the problem of the sensation of effector movements (as well as position in fact) can only be solved by a concern with the above *non-movement* representation of the effector. From that point of view we might again quote Woodworth (1899) whose intuition appears to us to have hit the mark: he stated that not only, for instance, is the judgement of the extent of a voluntary movement: "*a sense which is not reducible to a sense either of its force or of its duration, or of its initial and terminal positions,*" but also that it is "*the real basis on which the control is based.*" (Woodworth 1899, p. 80; our emphasis) This corresponds exactly to our argument that the *sensation* of movement, *just as its production*, must be a function of effector representation, or equivalently (as argued initially in sub-section 2.1.0) that there is no problem of kinæsthesia outside of the problem of movement production. Unfortunately, Woodworth did not disclose what he thought the nature of that basis might be.

Within performance-driven research, there has been a tendency in recent years to suggest that the basic problem of movement lies with the discovery of invariants. (cf. Schmidt et al. 1979) Insofar as at the core of a given movement command it is believed that there exists a fixed or primary shell, which would allow for, or be able to 'support,' non-essential variations in execution (the question of equivalence), we are in complete agreement. However, even though this would by definition be true for *all* movement command, the particular invariant would also *vary* according to a variation in commands; as opposed to this our use of the term invariant has a much more general purport, since it refers to the functional structure of the whole control *system*, and therefore our concern is at a level where we are not yet preoccupied with the essential or invariant feature of any *one* particular command eventually issuing from it. However we believe that the effector representation is *also* an invariant, and furthermore in exactly the same way in which the term can be applied

to any particular movement *command*. In order to see this it is sufficient to state what an invariant is: *an invariant consists of a constant relationship*. Thus the representation of the effector, as an invariant, should be built around some (as yet unspecified) fixed relationship/s among its components. With respect to particular movement commands on the other hand, it should now be evident that invariance refers to the supposed fact that at the core of any command lies a statement on the relationship to be observed among effector *components*, a statement - from our point of view at least - that is presumably made upon the representation of the *effector*. Another way of describing the relation between the two uses is to say that there is no such thing as a *general* movement invariant (i.e. that would apply to all movements) outside of that provided by effector representation, and that in fact all particular movement invariants are founded upon this single one. In fact, a movement (i.e. command) invariant will be defined as the invariant feature of the *transformation* of the effector representation (i.e. the global motion directive). This question will again be taken up at the end of sub-section 2.2.1.1.

Lastly, a few words are in order on the relationship between our hypothesized effector representation and that which lies at the root of robot systems, as apparently provided through a kinematics analysis of the effector system (cf. sub-section 2.1.2). The major feature of this type of analysis is that it is *positional*. It is based upon a few elementary spatial-physical properties of the effector system, namely the length of the components and the axes of potential movement, or degrees-of-freedom, associated with each of the constituent links. The main purpose of the analysis is to compute the position of the tip of the end-effector (or manipulator link) as a function of the position of the manipulator and that of all other effector components so that the high-level commands, expressed as *paths*, can be understood by the system. Of course insofar as the position of the manipulator tip is indeed a function of those of all the rest, it can be argued that a sort of global position is implicitly computed. However it does not follow that, in itself, the former value provides a representation of the effector as a whole. In fact the converse is true: it is the *product* of a representation of the effector, namely that *provided by the kinematics analysis itself* (i.e. the resultant set of equations of motion describing a particular effector system); *that* representation is the invariant at the root of the Robotics approach to movement. Picking up on some remarks made at the end of sub-section 2.1.2, we would now argue, paradoxically, and again insofar as kinematics analysis is taken seriously as a potential candidate for producing conjectures on a human level of control and/or performance in movement production, that, with respect to its *generality*, this kind of effector representation is too powerful: having to deal with but a single (or at most a small number of) object/s of control, humans simply have no need to access the whole analytical arsenal of Kinematics: something much simpler, and possibly tailor-made, could well suffice. On the other hand, and again paradoxically in view of its apparent power, Kinematics has not yet yielded a *human level* of

motor performance. We tend to believe that this is because Kinematics' effector representation is somehow missing the mark with respect to members of a biological class of effectors (namely arms), or, more precisely, with respect to the level of control of these effectors displayed by humans. In other words, not only is it necessary to achieve a representation of the effector, but it must be such that control at a human level of performance becomes possible.

In Robotics, with respect to the goal of achieving a human level of movement performance, the 'value' of the kinematics effector representation is never discussed *per se*; in fact it is apparently taken as the only possible one: problems of control are always defined, and solutions to them are always sought, *within* the confines of the fundamental representation. (Very similar remarks could be made about the dynamics analysis and the dynamic effector representation it rests upon.) In that context, our efforts could be described as an attempt at formulating a new kind of Kinematics, that is, at formulating a new kind of effector representation.

It was noted above that in Robotics high-level commands can take the form of paths. Now a path is a natural way of describing *objectively* (i.e. in terms of objective or formal geometric space, and more specifically the space in which movement is occurring) the most 'useful' part of the behaviour of an effector; indeed, as noted above, this is probably the major reason for initially adopting the kinematics effector representation, since in order for a robot to understand path commands, it must be provided with a representation of the effector that is compatible with them: adopting the first, namely the description of commands as paths, naturally leads to the adoption of the second, namely the kinematics effector representation. Thus in Robotics there seems to occur a sort of internalization of objective space into a 'subjective' representation of the effector, which itself apparently rests upon an unquestioned acceptance of the supposition that objective space is the *only* possible foundation for any kind of spatial knowledge (including the visual and the motor). Hence, by adopting the kinematics approach one is also accepting this implicit underlying assumption. Abandoning the kinematic effector representation does not of course mean letting go of the necessary cohesiveness between the commands and the system's representation of the effector; it does however imply abandoning reference to objective *extra-effector* space as the ultimate framework for specifying the intended results (i.e. the commands), and, by functional extension, for describing the effector as well (i.e. defining the effector representation) *in a manner that is compatible with it*.

Briefly stated, in Robotics one encounters two uses of space, one pertaining to the spatiality *in which* movement is occurring, the other referring to the *intrinsic* spatiality of the effector; since both are objective they are also compatible by definition. But supposing this to be the case for humans means supposing that for the purposes of movement production we have at our disposal a working knowledge of formalized space. However it cannot be argued that *just because* we can describe both the effector engaged in movement and the movement itself with reference to objective

space that *it must follow* that the process of movement production is based upon a knowledge of formalized space. Since from our point of view (cf. chapter 0; Lamontagne & Beausoleil 1982) the only possible foundation for spatial knowledge, *including objective space*, is the spatiality of the (visual and motor) interface, in the context of movement production a reference to objective *extra-effector* space is de facto excluded; furthermore, insofar as there does exist within the system a capacity to refer to extra-effector space, it will still have to be rooted in the more primitive spatiality of the effector. Given however that the motor interface does have an intrinsic spatiality associated with it, the above does *not* exclude reference to objective space for the purpose of describing the effector, that is, for achieving the effector representation.

In this context, it will be recalled how in the previous chapter the helicoidal ring achieved a visual representation of lines, and how the derivation was rooted in the intrinsic geometry of the retinal interface, whose structure, on the other hand, was described according to standard geometry; it was also suggested in Lamontagne & Beausoleil (1982) that, by extension, the whole of *visual space* ultimately rested upon the spatiality of the retinal interface. Similarly in the context of movement production the implication is that the *commands*, insofar as they presuppose, and bear witness to, a form of spatial knowledge, will also have to be based upon (or be a direct function of) the spatiality of the effector, i.e. upon the effector representation, a knowledge whose extent and richness will only be limited by what can be allowed by the internal representation of the effector system, upon and through which the command structure exercises its power.

This open-ended approach in fact constitutes a complete reversal of the situation obtaining within the Robotics approach: while it was argued above that the effector representation was dictated by an otherwise independently specified command (i.e. the result of movement in terms of the path of the manipulator in objective space), in our case we have argued for the opposite, namely that *the commands should be dictated by the nature of the effector representation*. This has a peculiar consequence however. By shifting the weight of the control of movement onto the question of effector representation, by making so much depend upon that representation, this new approach leads to a situation in which, contrary to the classical Robotics paradigm, it is no longer possible to start with a specific and well-defined set or class of movements to be produced by the system (cf. introduction to section 2.1). On the other hand, given a representation of the effector (i.e. a single entity standing for the whole effector) and a control system established as a direct function of that entity (i.e. a system that will effectively carry out degrouping according to the structure of that entity), one should easily be able to explore, by specifying modifications of the characteristics of the effector representation (i.e. by characterizing motion), the range of movements thus made possible. Of course in that respect we would still be studying the system's productions in terms of their result, for instance in terms of what happens in objective space to the tip of the effector.

Henceforth, the nature of the control model belonging to the Robotics approach will be referred to as being of the EXTRINSIC type, while the open-ended type of control model for which we have argued above will be referred to as being of the INTRINSIC type.

At this point let us make one further remark upon the possible psychological interest or relevance of the proposed scheme (*from representation to command*) by assuming the outlook of a newborn. We are in effect proposing that the effector representation is built-in, or innate. We are furthermore and consequently proposing that the development of skill in movement is the result of the child's exploration of the command repertoire as defined by that representation: moving is essentially a question of 'playing upon' that representation; hence moving meaningfully or skillfully becomes a problem of finding the correct transformation to be applied to the effector representation in order to achieve a given result. Thus motor 'learning' would occur from the very start; and apparently it need never stop, as shown by the fact that it can continue to occur throughout life. For us this is indicative of the intrinsic richness of our effector representation, a richness that corresponds to and allows for the human's open-ended ability to imagine and carry out a continuous flow of new purposes; i.e. in our view, such an ability is contingent upon the effector representation.

With respect to movement, at least according to the above view, a newborn's task is thus very different from a modern robot's task. And indeed by current standards the system we are about to propose would make a poor-robot. However we believe that this is so for the same reason that in terms of quality and range even a child's motor performance is incomparably superior to any existing robot's.

2.2.0.1 Tactical points

Our main objective in the present sub-section is to show that a single-object, or intrinsic, strategy *for* the system implies the adoption of a bottom-up tactic in *constructing* it. Although this might appear contrary to the main top-down thrust of movement production and control, it should be kept in mind that the present discussion, contrary to the one above in which was outlined a conception of its fundamental character, now pertains to the problem of developing a system in which it can be realized.

Following that discussion, it should now be evident that the major tactical issues arising in the development of a movement production system center upon the question of implementing the effector representation. Essentially this is the problem of specifying how the system proceeds from the top level of control, where it deals with a representation of the effector as a single entity, down to the last level of control, where it deals with each one of the effector components. Since motion characterization is a function of the diversity of features attendant upon the effector representation, the matter of computing motion, at least at the top level of command, becomes subsidiary to a

concern for this more static aspect of control.

However motion computation is not a question that arises *only* at the top level, i.e. the level at which the effector is represented as a single object of control. Observation of practically any movement of the arm reveals the obvious fact that the arm can only be involved in movement if some (or all) of its components are involved in movement. This implies that not only is motion computed for the whole arm, but that at some point it must also be computed for each one of its parts. In fact the only reason for proceeding from the global to the local level (from the whole effector down to its constituent elements) is to carry a top-level motion directive into the proper set of local motion directives. Thus the purpose of the effector representation is to allow the system to define for itself a group of sub-objects of control (corresponding to a representation of the effector components), each of which in turn will receive its own motion directive as a function of the global motion directive. *But just as the motion characterization of the global object of control is a function of its attributes, the motion characterization of the sub-objects will also have to be a function of their attributes.* Consequently, defining a representation of the object of control at a given level will also take precedence over motion computation.

But it also means that the potential for motion computation provided by the partial effector representation or sub-object obtaining at a given level is a major factor to be taken into account in implementing the final (or complete) effector representation. It is not enough to simply specify a path from the global object of control (the effector representation) to the local objects of control (the partial effector representations): since a motion directive originating at one level must eventually lead to a modification of the characteristics of the next level's objects, the specification of this course is additionally and critically constrained by the potential for motion computation associated with these objects. It is therefore necessary to know the nature of the object/s of control at one level *before* proceeding to the definition of the next level's object/s; this is what is meant by the bottom-up approach to the problem of control.

In doing so we have also in fact opted for an overall control scheme whereby movement computation occurs at every step of the degrouping process. That is, in proceeding from one level to the next not only is the object of control split up into a group of lower objects of control but motion is necessarily computed for each one of them before the system proceeds to the next level. Thus motion is 'carried through' and recomputed every time for each type of object specified in the system; just as the objects of control change from one level to the next, so does motion computation. Since motion is specified by a modification of some characteristic of the control object, the preceding implies that the next lower object of control will have to be permanently tied to that particular aspect; otherwise the modification in that particular characteristic implied by some global change in the object will not be carried through to the next lower object.

We have stated that the representation of the object of control at a given level should allow for motion computation. The physical part of the effector to which the representation corresponds has associated with it a movement potential (its kinematics). Therefore "allowing for motion computation" means that the system can access the full movement potential, i.e. the full range of possible movements, of that part of the physical effector. As successive representations become more and more encompassing with respect to the whole effector, this criterion must be continuously satisfied, i.e. *all* possibilities must remain available to the system. Of course as the complexity of the effector grows, so does the complexity of its associated movement potential, as well as the difficulty of achieving a representation of that effector. Since that complexity is a direct function of the movement potential of the constituent elements, a bottom-up approach is the only method for constructing the representation and for ensuring that at the top level of command everything is indeed possible. Thus *an object of control is the representation of a potential*, and in particular the effector representation is a representation of the movement potential associated with the effector.

But how can a potential be represented; and furthermore how can the potential of an effector be represented so that the potential of its constituent elements remains, as required by degrouping, fully accessible? In the kinematics approach to this problem, the elementary movement potentials (associated with the effector components) are represented by equations, and these equations are successively *embedded* one within the other until the whole effector is 'reconstructed' in a final equation. This is how the representation of the effector (i.e. the representation of the movement potential associated with it) is established; and it is the embedding procedure itself that guarantees full access to the 'lower' potentials within the complete description of the movement potential of the effector. However, Robotics systems do not directly act upon that representation for purposes of command; that is, the representation itself is not exploited *per se* in that it does not lend itself to the specification of commands. Rather, and conversely, as we argued earlier, the command itself is determinant for the nature of the effector representation. Insofar as it does allow for the specification of commands, they are not global in nature; this is because the only direct access to the representation is through its local constituents, namely the variables. Nothing is thereby gained since the complete specification of a command could only be done through a large set of 'independent' variables. This would not be the case if the model somehow boiled down to a 'superequation,' built from 'supervariables,' allowing the system to proceed from the latter to the much larger set of 'real' but local variables.

There is an alternative to this mode of representation however, one that is built upon the intrinsic spatiality of the effector itself but that is not driven by objective extra-effector space. Of course it still uses the same formal geometrical constructs as are used to describe objective space, but the difference is that there is an explicit representation of the object of control (as opposed to

Robotics' implicit representation) and no explicit reference to objective space in the *command*. To give a first example of what that alternative might involve, consider again the helicoidal ring, not as a perceptual system but rather, by inverting its functional direction, as a hypothetical system for *producing a moving* line segment. In that context the helicoidal ring, i.e. *the structure itself*, is a representation of the 'movement' potential associated with the system's 'effector,' namely a line. To command a movement the system would simply 'move' the effector representation; that is, it would prescribe modifications along any or all of the characterization dimensions of the effector representation, the single object of control. Conversely, *any* modification would result in movement. Now if a particular movement, say a rotation, is considered, it will be seen that the rotation command, consisting of a vertical 'jump' into an adjacent layer of the ring, concerns the line as a whole: the command does *not* specify the rotation in terms of the individual movements of the local elements of the line (the points); however these local movements do eventually occur, but they are indirectly determined by the fixed or wired-in 'degroupping' process, for which the system was arranged in the first place. (It occurs as a function of the invariant property of the 'effector,' namely linearity, of which it is the expression.) Nor is the global command specified directly with reference to 'objective' retinal space, but rather with reference to the spatiality of the system itself (intra-layer and inter-layer), even though the rotation eventually does 'happen' at the retina. Here the parallel with real effector movement is lost somewhat since the product cannot really be described independently of the interface, although it could be made to in the case of a binocular system that generated a line moving in three dimensional space instead of on the surface of a retina. However that is of minor importance here. Clearly, the nature of the command is dictated by the nature of the object of control, and not vice versa; and it would still be true even in the case of the production of a moving spatial line. This is the alternative we have explored.

Although, following section 2.0, the above does constitute a case of degroupping, with respect to real effector movement it could be misleading in another way. As outlined in section 2.0 the principles of degroupping focused exclusively upon the static case. The difference and possible source of confusion in the above example stems from the fact that strictly speaking it is the production of the line segment itself that should have been taken as the equivalent of movement and not the actual movement of the generated line segment; from that point of view the 'inverted' ring would actually produce only *one single type* of 'movement,' although it does encompass quite a large number of different 'movements.' Accordingly, the production of a *diversity of types* of 'movements' would then have to correspond to the generation of a variety of different retinal *shapes*; of course the inverted ring could be used for that purpose, as it is also apt to process a variety of retinal shapes. (It would however be impossible to move into the area of shape generation without multiplying the number of objects of control *at that level*.) As used above, the inverted ring does have a single

object of control (the 'lines'), and the control does occur through modifications afforded by that representation. This is our main point.

This will help us to bring out another major aspect of the problem of movement production: Generally speaking, a movement of the whole effector involves *different* movements of its components; but these components have the same status in the control system. One could therefore imagine that degrouping is a matter of building a representation of the effector in which the global entity is simply 'attached' to a certain number of lower primitive entities, and indeed this would formally count as an instance of degrouping. However when it came to specifying a motion for that entity, it would be realized that the homogeneous binding between the global object and the local entities only allows for movements in which all components do exactly the *same* thing. That will simply not do. In the case of the inverted ring, that is not the case of course: generally speaking, a motion of the line segment will involve many different movements of the constituent points. A motion directive in the inverted ring will generally imply that *many*, if not all, components will end up engaging in different movements; furthermore a motion directive which is specified as a modification of only *one* characteristic of the object of control (e.g. orientation) will *also* generally imply different movements of the components. It will be seen that this latter possibility does not hold for our own system. This is apparently due to the way in which the effector representation was achieved: there *had* to be secured a way, or a degrouping path, by which the global motion directive could be translated into local motion directives. The fact is that in the inverted ring example there are not any local motion directives per se; they are merely induced from the results of the production process.

Suppose however that instead of a single moving line, the inverted ring was called upon to control the motion of a linkage system composed of three jointed line segments; in that case it could no longer be argued that only a single object of control existed within the system (there would in fact be three of them) and consequently that the ring in itself would suffice to implement the control. What that representation would have to consist of in terms of the *perception* of the movement of that object, we do not know; but we believe we have some idea of what it could be for the purpose of its *production*. Insofar as the production process could then itself be inverted into the perceptual case, we would argue, and this is recognizably in line with the single visual object scheme (the SVOS), that the global characterization of its movement would have to consist of a characterization based upon a representation of the object as a whole; that is, it would first have to be grouped - according to some invariant - into a single entity, after which its movement could be characterized globally according to whatever was allowed by these characterizations. Since the nature of the global movement is a direct function of the local movements (i.e. the movements of the components), we suggest that these local movements would somehow have to be 'carried' into or

remain 'accessible' at the global level of representation; indeed it seems impossible to achieve the global representation of the whole effector (i.e. to 'group') without constantly and explicitly being concerned with the intended goal of characterizing movement.

This being said, we are now nearly ready to describe the particular control models. As a last prelude to this, the following sub-section is devoted to the definition of some interfacial constructs that will be used to describe, in sub-section 2.2.1, two simple effectors, along with the corresponding models of control we developed for them.

2.2.0.2 Some basic constructs

We now introduce four constructs pertaining to the motor interface. The first three are descriptive of the physical nature of the interface; the fourth is a theoretical concept which is more closely related to the experience of movement or effector control, but it is nonetheless introduced here as it pertains directly to the periphery. On the other hand its insertion into our current models of control has not yet been as successful or as complete as hoped, for reasons that will become apparent as we go along.

The first term is DEGREE-OF-FREEDOM (henceforth "d.f."). Although we have already referred to it repeatedly in the present chapter, we must now make it clear exactly what it corresponds to in our models of control. On first view, one might think that d.f. is but another name for variable. However a d.f. is different from, and more than, a variable in that it refers to the ability of certain *entities* (e.g. subjects in statistical inference) to *potentially* display, *independently* of all other similar entities, a certain degree of variation on a given dimension of observation. Thus for instance the variable "mass" is not a d.f. because in general it does not refer to a class of physical entities nor, insofar as it does, to a supposed capacity for each of them to potentially exhibit a variety of values on the associated dimension, namely amount of matter. On the other hand we will for instance speak of the n d.f.s associated with an arm in referring to the fact that in order to describe an arm's (the object) positional (the particular dimension) potential, n independent values are required; of course each of these values is variable, and so the key to the distinction lies in the fact that what is referred to by saying that the arm has n d.f.s is its *ability* to assume a variety of positions; one is not here interested in finding a way of describing a *particular* position, merely to indicate that it *can*, as well as by 'how much.'

In the next sub-section it is shown how a d.f. can become an object of control. In our systems it constitutes the most primitive one: this implies that it must be specified, that is, characterized. This specification is *not* a differentiation of that particular object from other potentially identifiable objects; rather specification here refers to the task of selecting or drawing from *within* the virtual (characterization) pools associated with it the particular set of values desired. In the case of effector

systems in general, and in particular in the case of our own, that specification occurs in terms of two characterization dimensions, namely a *direction* and an amplitude or *extent*. Now in Robotics for instance no such segregation is suggested: a d.f. is directly handled through a single signed variable. While the fact that our two characterization dimensions, at that level, are bound in a single 'entity' can appear to be computationally advantageous, the direction (or for that matter the extent) characterization is not 'carried' through into higher characterizations in a way that would make clear the filiation from this low level direction into a high level direction specification; are we to assume that they are two completely different things? Similarly in performance-driven research the *movement* characterizations, which in general pertain to movement as a whole, are tentatively transformed into *motion* characterizations, with very little regard to those possible 'intermediate' levels of motion characterization that would bring the high-level ones (such as direction) into the low-level (direction) characterizations associated with the peripheral components. We tend rather to believe that the specification of high-level characterizations of, for instance, direction can only be understood or built up in the light of a strong dependent relationship with the low levels of control or specification, and not the converse. While there is reason to believe that global extent and global direction can be segregated at the top level, there is on the other hand no a priori reason to think that it does not remain so right up to the peripheral stage. This of course in no way affects the possibility, and indeed the necessity, for these dimensions to 'interact' among themselves during the whole process of degrouping. Thus for example the selection of a certain direction for the whole motion might well impact on the specific directions ultimately selected at the level of d.f.s.

The second term is SUB-EFFECTOR (henceforth "sff"). As its name implies, it refers to a sub-part of an effector system. More precisely, it consists of an effector segment *plus* its associated d.f.s. Thus the forearm has two d.f.s, one for extension, the other for rotation; the forearm considered as a sff therefore refers to both the segment and the two d.f.s. While the number of d.f.s associated with a sff can vary, there is always only one segment associated with a sff. In Kinematics the physical part of a sff is called a *link* (cf. Rooney 1984), while d.f.s are connected with *joints*, a joint implying the existence of at least two links. While this holds in our case, so that the number of d.f.s associated with a segment is *always* computed individually for that segment (that is, the second link is always considered fixed), we have not, as in Kinematics, tried to describe, for eventual control purposes, the *overall* movement potential of the arm in *those* terms. (However the problem of describing the overall movement potential of linkage systems - and therefore of particular ones such as arms - is also the central problem of Kinematics.)

There is very little a system can do about the fixed physical segments of the arm per se; but it can command movement effects upon them. In our systems these are exercised through 'forces,' which in turn are mediated through d.f.s. Thus to move the forearm, a system must know *which*

particular d.f. can 'act' upon it and specify certain direction- and extent-modification values. In that way a sff also becomes an object of control, an object that is *constituted* from lower objects, namely d.f.s. In the case of the forearm, where only one d.f. is involved, the sff in fact functionally reduces (except for the identification of its specific d.f. by the system) to the lower object; however if the upper arm is now considered, which in our (second) system involves three d.f.s, then this is no longer the case. Yet we can still think of its characterization in terms of direction and extent; however the difference now is that since d.f.s can act in concert, many more directions are *possible* for the sff than for any of its d.f.s; but an *actual* motion (or command) cannot be specified in terms of the component d.f.s (even though of course once it is, the business of resolving it into the three sets of lower-level commands or characterizations still has to be taken care of): it must be done 'according' to that *new* object, namely the sff. This 'growth' in the movement potential associated with the sff derives from what might be called *df. interaction*, which simply refers to the fact that the physical 'presence' of three d.f.s at a segment gives rise to new movement possibilities for it above and beyond those of its constituent d.f.s. The whole point of degrouping is that some prior knowledge of that interaction must exist if a system is to achieve control: it is one thing to talk to each of three d.f.s *individually*, but an altogether different matter to tell a meaningful thing to all three *simultaneously*, something, we believe, that cannot be done without the existence, within the system, of but a *single* entity to which it can initially address its commands.

When we come to the next construct, namely the EFFECTOR (henceforth "eff"), the story again gains in complexity. While describing the sff, it was assumed that its neighbouring sff (by which was meant the next sff *up* - the hand's neighbour being the forearm, etc.) was immobile. At the stage of the eff, this restriction is lifted, and we are dealing with a string of single physically connected sffs, in which furthermore each sffs can have one or two neighbours. (Since the human hand involves *parallel* strings of single connected sffs (the fingers), it does not belong to the class of effs we dealt with; in our case the hand was not completely absent however, but it was simplified into a sff: only one segment, and a few d.f.s.)

Just as physically an effector is constructed from a series of sffs, in terms of control an eff will become an object built from more primitive objects, namely those that are used in the control of sffs. While at the level of sffs it is still arguably possible to 'cover' the movement potential with a simple extension of the d.f.s characterizations, it is more difficult to do so at the level of the eff, even though of course the more primitive ones are maintained: that is, it is still possible to speak of the direction and amplitude of movement as applied to the eff; however, through something we call *sff interaction*, and which is directly related to the d.f. interaction mentioned above, a much wider, and much less easily described (and by extension organized, or controlled) domain of variability is opened up, one in fact, given a sufficiently complex eff, that can be claimed to approach

the potential associated with the human arm. However there is no reason why the approach used with the more primitive components of the eff cannot be put to work on the eff itself. This implies that somehow a representation of the eff *as a whole* is necessary, as well as motion characterizations that apply to it. It means furthermore that a process of translating that single object into sff objects is also called for. However given that very little is known a priori about the potential dimensions of residual (or secondary) characterizations associated with the effector representation, there is no way of starting to work at that level first and then of pushing down the results into the sff level of control (which, on the other hand, is exactly what does happen during control); rather, the converse will have to occur. This very important tactical consideration was discussed in sub-section 2.2.0.1.

Concerning sff interaction, it will be noted that at least *part* of the complexity can come from the fact that, for a given eff, sffs *vary* in terms of their number of d.f.s and of their physical length; in terms of the eff this means that the components need not be (and in fact are not) *homogeneous*, with respect to their own makeup and consequently with respect to the spatial relations that each entertains with its neighbour/s (contrary to our example in section 2.0 where, it will be recalled, a retinal metaphor was used in our initial discussion of degrouping). The difficulty that arises in the development of a control system is that it becomes nearly impossible, having chosen the d.f. as our starting point, to treat all sffs in exactly the same manner: this is because, being physically different, sffs do not contribute equally to movements; this contrasts with retinal receptors, each one of which could be argued to (potentially) contribute as well as any other to the peripheral definition of percepts. But it should be added that we do not think that this is an unsurmountable obstacle; indeed we believe that in the nervous system the representation of the motor periphery (i.e. the most primitive component of the movement control system) *is* homogeneous; but we also believe that it does not deal with d.f.s per se. Nevertheless the basic issue remains the problem of degrouping (cf. section 2.0); however we must admit that the difficulties inherent to our choice of interface - that is, one where, the d.f. being the most primitive component, the non-homogeneity of the eff persists - had a profound impact on our work.

These three constructs, namely the *degree-of-freedom* (d.f.), the *sub-effector* (sff), and the *effector* (eff), completely describe the type of motor interface upon which our search for a solution to the problem of 'motor' degrouping was bound. In particular in section 2.2.1, where two control schemes are presented, they will be used to define the specific instances of virtual arms upon which these systems operate.

It will have been noted that our discussion of these components was nevertheless functionally loaded, that it did not remain completely untainted by a preoccupation with *movement* considerations. This is particularly evident when it is suggested that from being purely *physical* these components are also called upon to become abstract *objects* within control systems. But this should

not be surprising: for while in vision research the supposition that individual *receptors* give rise to (the most primitive) visual *objects* is universally, and most often implicitly, made (but see Lamontagne (1986) for an explicit formal statement of this assumption), in Robotics research one of the most important considerations in arm design is the complexity of control implied by the complexity of movement inherent to the physical makeup of manipulators (cf. Brady et al. 1982; Rooney 1984). While this was not the case for us, since we had opted from the start to deal with a humanoid arm, it is nevertheless true that, having done so, the identification of levels, along with their respective objects, did occur in light of the movement properties of the arm.

Having said that, the above triplet will now be augmented by one further construct, namely the LEADING SUB-EFFECTOR (henceforth "lsff"). In a sense this is the most general control concept that will be encountered here. On the other hand, as hinted at in the previous sub-section, it is not one which we have been able so far to implement in our systems in as a powerful a way as we believe it eventually can be.

The first important remark to make about it is that, even though it refers to a sff, the particular sff in question can vary. In other words, lsff is a secondary characterization for control. It is a tentative formalization of the impression, derived from movement experience, that, for instance in a reaching movement the hand is the 'leader,' and that it somehow 'pulls' the rest of the sffs along in a manner dictated by its own course. Of course at the peripheral level of execution this is quite impossible. In fact, there the converse is true: it is the rest of the sffs that, directly or indirectly, pull on the 'leading' sff. But there is no a priori necessity for all features of the control system to be apparent at the peripheral level; indeed the whole problem of movement control lies in the fact that most (if not all) of them are not (as performance-driven and system-driven research amply demonstrate). As a component of the command, the lsff would however not be a motion characterization like the others, since in general, for a given command, it would assume a single value; that is, in the motion directive, while the rest of the command would presumably be made up of a set of *modifications* of various secondary characterizations, the lsff for its part would not be. Because the choice of a lsff is critical with respect to what the rest of sffs will do in a particular movement, we believe it to be rather high up in the hierarchy of command specification. Note finally that in Robotics, the manipulator's hand is the lsff, and that the whole kinematics analysis of the manipulator is dictated by that choice. However with respect to robotics control it is the only possible one.

Secondly, at the root of this notion lies the (performance) problem of equivalence, which, it will be recalled, stems from the observation that a given movement can be produced in equivalent ways or that a given (presumably unique) command can be 'satisfied' in a variety of ways by an eff. Now equivalence is a very general *cognitive* feature (cf. Lashley 1951); consequently one can easily be tempted to believe that the key to *motor* equivalence (that is, equivalence as manifested in the

movement domain) is not to be found within the control system, but that it is rather 'imposed' by superlying cognition. However by removing equivalence into the purely cognitive domain, one does not in fact remove the problem of motor equivalence from the movement control domain, one merely transposes it; for consider that the kind of knowledge required for the control of the motor periphery in equivalent ways is inherent to the control system, not to pure cognition: that is, that in order to respond equivalently to a given goal (that could of course be provided by non-motor cognition), the control system must be able to recognize the uniqueness of the purpose; otherwise, cognition would have to tell the motor system, and in doing so of course it is no longer free from involvement with purely motor knowledge. Conversely, to the extent in which *motor* equivalence partakes of pure cognition, there is no problem in bringing it down to the level of control; the idea is not to *reduce* cognition to movement control, merely to point out that there might be reason to believe that cognition is a *continuum* having some extensions, if not some roots, in movement production.

Now what does all this have to do with an apparently simple thing like the specification of a lsff? When one reaches for a cup of coffee, we can say that the hand is 'leading' the movement. When one looks at one's watch, it is the forearm. And when one elbows on a table, it is the upper arm. Thus the system apparently has the ability to specify the 'goal' sff, the one to which the command is *primarily* directed, the rest of the sffs being 'led' by whatever the lsff will do. Now as we understand it the question of equivalence lies as much in deciding for a lsff as, once having done so, in deciding the correct motion for the rest of them (i.e. making sure that *their* movement, *along with the lsff's*, will 'amount' to the 'same' overall movement, or satisfy the same goal): the whole problem of motor equivalence boils down to the fact that, for a given purpose, changing the lsff *implies*, in general, that the correct movement for the rest of sffs will also change; or again, how can the control system, given a goal and a lsff, decide upon the correct motion for the rest? The problem of correctly deciding of course *already* exists within *any one* realization of the command; however here we are faced with the fact that *many* satisfactory realizations are possible. Because of this, and because we believe that a solution to the problem of equivalence can only be founded upon the system's capacity to correctly compute *one* motion, the problem of equivalence has been, even though it is *apparently* a high-level problem, considered by us to be a fundamental factor in the design of control systems: that is, specific control structures were always searched for with a view to 'allowing' for equivalence; in fact (see section 2.0) *insofar as a general computational solution is sought* to the first problem, namely the ability to compute *one* correct motion for every single purpose or command, the second problem, namely the ability to compute *many* equivalent motions for *one* particular purpose or command, is inescapable. Work on the development of particular control systems has, furthermore, brought to light a more primitive level at which equivalence arises:

this is simply the problem of taking into account (for purposes of 'nullifying' it) the variation inherent to the initial position of the eff. (Recall the discussion of Hebb's (1949) proposals in sub-section 2.1.0.) Thus for instance a reaching movement will be more or less ample (globally as well as in terms of the sffs involved) depending on exactly where the eff is when the movement is initiated; this problem must be solved before there can be hope of tackling the more complex cases of equivalence, which usually additionally involve a variation in lsff.

Generally speaking, what we are proposing is that the *selection* of a lsff is a key functional element in the phenomenon of equivalence. As such, and within the control system, it is a high-level concept. Intuitively, what this means is that the selection of a lsff is a decision occurring *prior* to the computation of sff motions; it is also a way for the system to decide upon which *object* the command is to bear. So far we have spoken of d.f.s, sffs, and the eff, as objects of control. From this it could have been deduced that these were the only possible ones; however that is apparently not the case: selecting the forearm as lsff for instance will give rise to a *reduced* eff, but an eff nonetheless, involving only the upper arm and the forearm. The sffs involved are *always* physically contiguous: we do not allow a reduced eff involving say the hand and the upper arm only. (The reason for this has to do with the nature of the effector representation and the fact that it is difficult, if not impossible, to specify or 'extract' from it a subset of sffs that are not contiguous. On the other hand, one can also point out that, *with respect to a given eff*, movements involving such discontinuous subsets of sffs never seem to occur.) At the level of command, it would therefore only be necessary to specify a lsff: once this is done, the system would automatically understand that the object targeted is in fact made up of the lsff and all other sffs *upward* from it to the 'end' of the eff; hence neither is a reduced eff involving both the forearm and the hand allowed in our systems. (Incidentally the above rule implies that an even higher-level decision concerning the selection of an eff has been taken; but with this observation we are introducing the problem of controlling a multiplicity of effs, an area that will not be explored further here; recall that, by default, in our systems, the eff is *always* considered selected.) The systems would *then* proceed to compute (i.e. solve for) the correct motions for the sffs (and thence for d.f.s) *according* to the prescribed goal.

2.2.1 Models for production

While in section 2.0 we had introduced and explored the problem of control and degrouping with respect to a *visual* periphery, and had formulated the basic task of a control system as that of *exploiting* the imaging potential associated with it, in the section and sub-sections that followed it attention was turned towards the problem of control and degrouping with respect to the *motor* periphery. This led to an understanding of movement production which could be summarized in the following manner: *movement production consists of the exploitation of the movement potential as-*

sociated with the motor periphery. In sub-sections 2.2.0.0 and 2.2.0.1 we argued [1] that in order to do so a system would have to be provided with a representation of the effector as a whole, [2] that this representation, by its very nature, would be the basis for motion computation, and [3] that this representation would have to be constructed from the bottom up, that is from the most primitive movement elements of the effector up to the whole effector; this is because the (global) effector representation is supposed to encompass within itself not only the whole effector but the whole of the movement potential associated with it, which itself is a direct function of the (movement potential associated with the) constituent elements of the effector. Thus, in section 2.2.0.2, we introduced the most primitive 'movement' component of the effector (the d.f.), as well as an intermediary construct (the sff), and finally the effector itself (the eff). Now these constructs pertain to the periphery, and they are of course all directed towards the most essential aspect of the motor periphery with respect to the control of movement, namely its movement potential. But having done this, one does not yet have a representation of the movement potential of the effector; one merely has a way to start thinking about its possible nature in a more precise fashion. Our own specific 'reconstruction' of the periphery was of course already guided by the need to achieve a global effector representation; although this is not evident at this point because these constructs can be taken *independently* of the problem of providing a system with a global representation, it will nevertheless become quite apparent when we proceed to define the objects to which these components are linked up in the system. As object of control the d.f. will somehow have to 'survive' into some aspect of the sff representation, and this one will also have to remain 'accessible' at the level of eff representation. (Note that the same terms, namely d.f., sff, eff, and lsff, will be used to refer *both* to parts of the 'physical' periphery and to their representation within the system. It should however be clear from the context which usage is in effect at any one time.)

Although it might appear self-evident that the most important physical aspect of the motor periphery as concerns movement control or production is its movement potential, it is certainly not the only physical property of the periphery that is determinant in that respect. In Robotics, where one is dealing with a concrete physical effector, great importance also attaches to the (mainly hindering) effect of inertial forces and friction upon the 'pure' movement potential of the effector. For us however these considerations do not enter into play, as we have chosen to deal with virtual (i.e. simulated) effectors in which they simply do not exist; there is no need to exercise a force upon a virtual effector in order to make it move. (But of course if one wished it so, it could be arranged: the practice of controlling virtual effectors patterned after real physical effectors is a current one in modern Robotics research.) The problem of achieving a global representation of energetic manifestations such as friction and inertia at work in a given effector in a manner similar to the achievement of the global effector representation appears to us formidable. (But recall that in sub-section 2.1.2 it

was suggested that the dynamics, along with the kinematics, *might also* form an intrinsic part of the effector representation afforded in biological systems. As such, their control, or their specific 'contributions' to the controlled object, would occur in a similar manner to that of the pure movement potential of the effector system.)

Two models for control will now be presented. The first one is quite simple, and although it is a model of the *extrinsic* type of control (cf. sub-section 2.2.0.1), it will serve to introduce the more ambitious second model instantiating an *intrinsic* mode of control. It is the problems encountered in attempting to generalize the first control scheme to a more complex effector that led to the second model; these problems were in fact critical to the strategic shift from an extrinsic to an intrinsic type of control scheme. There was however, beyond a complexification of its basic structure, no fundamental change in the nature of the effector from the first to the second model. The importance of the first model lies in the fact that, even though it is of the extrinsic type, the control is based upon a rather different use of geometric principles than that encountered in Robotics research; the same kind of use of geometric principles was again sought, and found, in the development of the intrinsic control scheme.

2.2.1.0 A simple extrinsic control scheme

Consider the simple eff in Figure 1 consisting of two sffs with one d.f. each. For purposes of control it is necessary to suppose that the shoulder is permanently attached at a fixed point within the two dimensional surface constituting the (extrinsic) spatial environment of the effector. (This is why it is called an 'egyptian' arm.) The elbow, on the other hand, is permanently attached to the tip of the upper segment, which constitutes its spatial frame of reference. Thus with respect to their immediate frame of reference neither the shoulder nor the elbow moves; but with respect to extrinsic space (segment 0's frame of reference) the elbow can change its position. Now these attachments (or joints) are such that they each allow for movement of the segments about the axis that passes through the joint and is perpendicular to the two dimensional surface in which movement occurs (they stick out of the page). Thus taken individually both segments have equal mobility, but taken with respect to extrinsic space segment 1 has a greater freedom of movement than segment 0. This is because the second segment's movement can be superimposed or added to the first's. More formally, the system has two d.f.s; that is, in order to specify the position of the arm two values are needed. These can be either the two local coordinates (taken with respect to each segment's axis of movement) or they can be the two global coordinates of the tip of the second segment (taken with respect to extrinsic space). Since 'actions' are described globally (e.g. "move the tip of the arm from point A to point B") but control must ultimately be exercised locally, the problem of controlling the arm boils down to the problem of correctly transforming globally described (or simply, global)

goals into locally understandable (or simply, local) goals (e.g. move segment 0 by so much and segment 1 by so much). Such a scheme will shortly be presented for the 'egyptian' arm.

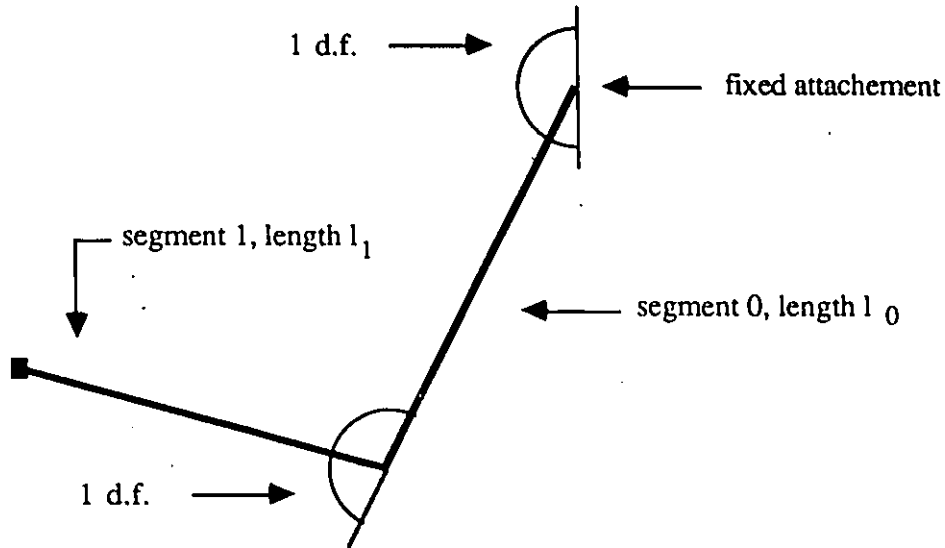


Figure 1. The 'egyptian' arm: 2 d.f.s, 2 sffs.

As suggested above, the simplest type of global command consists of a directive to move the tip of the eff from point A (the current position) to point B (the goal position). The reason for this is that all global commands can, *at the global level*, be decomposed into a series of smaller linear translations. Thus a command to move by following a specified arc (in order for instance to avoid an object) could be split up, or approximated, into a sequence of tangential movement commands, a tangential movement being a linear movement following a tangent to the specified path. However this presupposes that the 'smaller' movements themselves are linear, which unfortunately is not true in general. This is because even though one can calculate the *amplitudes* of the local movements (those performed at the d.f. level) necessary for the goal to be satisfied (namely, following the execution of the local commands, that the tip of the eff does indeed end up at point B), the actual *path* followed by the tip of the eff in going from point A to point B is *not* in general a linear path. Thus arises very early the problem of *path control*, above and beyond the problem of *global control*. What is simple at the global level does not result in something equally simple when 'directly' transposed to the local level of execution, and conversely what is simple at the local level of execution (e.g. that all d.f.s move by the same amount) does not result in simple movements in objective space. That is, taken with respect to objective space (also serving as the global command

reference) the inherent spatiality of the arm is complex. Yet from the point of view of objective space, and hence of global commands specified with respect to it, there is apparently no alternative to the fact that translations from A to B are the simplest possible movement directives.

We have *not* solved the problem of path control for the present control model, nor in fact for the next one to be presented below. It is however a problem about which we have a few things to say, and to which we shall return in the next sub-section. But at this point, we would like to make the following preliminary remarks about it. First, we believe that a solution to the problem of path control should only be sought once the global control problem has been solved, as indeed happens in Robotics: path control is subordinate to the existence of a basic level of control, upon which its solution depends (cf. Paul 1979, 1981); that is, as a more sophisticated kind of control, path control should be built on top of the global level of control. Secondly, it is well known that humans can exercise quite a precise control over the deployment of particular movements; this signals the fact that path control is indeed a basic feature of a human level of movement production. As such, it is therefore independent of the nature of the explanatory control scheme (intrinsic or extrinsic); it is an issue that must be addressed in both approaches. However, having favoured an intrinsic view of control (cf. sub-section 2.2.0.0) and having invested much more in the development of an intrinsic control model, the question of the way in which, in that context, the problem can be formulated and how it should be tackled, will be reserved for the next sub-section. For the moment we shall limit ourselves to a short comment. As stated above, the problem was not pursued in the context of the first (i.e. extrinsic) control scheme.

Briefly stated path control is the problem of imposing an overall form to movement. Recall from above that the whole issue of path control emerged from the recognition that it was not enough to simply calculate total excursions for the d.f.s, and that a whole new domain of variation was opened as a consequence of the fact that during execution *some* additional directive had to be supplied on the particular way in which these excursions would unfold. In effect the 'shape' of a movement is directly related to the way in which, within the confines of the previously calculated totals, the individual local movements develop. So again we encounter the same basic degrouping problem whereby the set of different local events must be such as to 'add up' or give rise to something that is not apparent at that level. Hence, as in vision, shape is a property of a group of local events, one that must be computed globally. It is of course at that level that a solution must be sought, and one furthermore that must somehow be 'fitted' within the extant global control. It will be seen for instance, in the second model, that there is a default shape directive which states that all calculated local motions have to occur within the same execution period (this gives rise to different speeds) and are linear; as a default, such a directive need not be explicitly taken into account at the top level of command. However supposing that it were to be, then the problem becomes, again,

one of finding the constant or invariant feature of the local execution that, for a given global command, forces it to occur, as a whole, in a specified way (e.g. so that the tip of the eff moves in a straight line). Although we are merely at the stage of formulating the problem, we believe that, given a working system in which a global level of command is implemented, it is possible to start exploring the question in a more systematic fashion. But we shall leave that question now and return to our first control model.

The basic principle behind our first control model is outlined in Figure 2. The key to the model lies in the exploitation of the constant lengths of the sffs. Here is how it works. Suppose, as above, that the global command takes the form of a required displacement of the tip of the eff from point A to point B. The problem of course is to compute motions for joint 0 and joint 1 so that following execution the tip of the eff, initially lying at point A, ends up lying at point B.

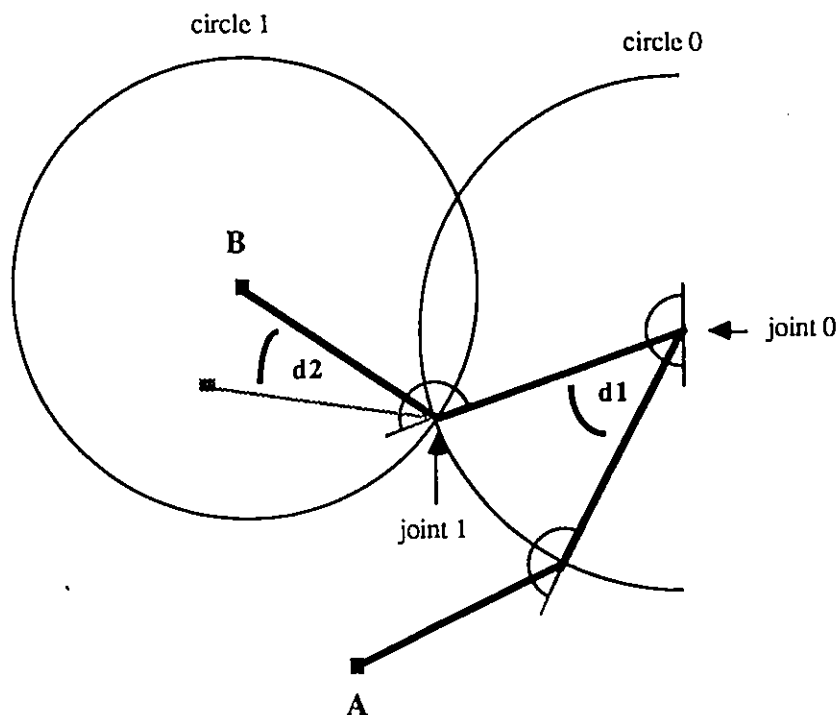


Figure 2. Control principle for the 'egyptian' arm.

Point B is used as a reference point to draw a circle (circle 1 in Figure 2) whose radius is equal to the length of the second segment. Another circle (circle 0 in Figure 2) is drawn, but this time by taking the fixed attachment of joint 0 as reference point and the length of segment 0 as its radius. In general these two circles will intersect at two points. If they do not meet, then the motions can-

not be computed, and the movement is not possible because it is out of the reach of the eff. If they meet at one point, then that single point is used for calculating motions. If the two circles meet at two points, as shown in Figure 2, then the question arises as to which one should be used to compute the motions. If the second sff could rotate freely about joint 1 then that would necessitate the introduction of a heuristic into the system for selecting the 'best' solution; however the problem never arises due to the fact that the second segment cannot extend beyond the alignment with the first segment, which it would have to do if the second intersection point were in fact a possible solution; since it cannot, there is always but one possible choice. Again, this question does not arise if there is a single meeting point. Having identified that point, the next step is to compute the signed angular difference between the 'initial' and 'final' positions of segment 0; this difference (d_1 in Figure 2) constitutes the motion commanded to the first sff. By 'moving' the whole eff to that position it is then possible to compute the second difference (d_2 in Figure 2), namely that between the 'initial' and 'final' second segments, which in turn will serve as commanded motion for the second sff. Thus, starting with a global formulation of the command, i.e. to move the tip of the effector from position A to position B, the system has computed the local motions necessary for its realization.

It will have been noted that the command takes the form of the intended result of a *movement* rather than the prescription of a *transformation* of some representation of the effector. However since it is interpreted by the system as bearing upon the whole effector and as implying a set of local motions, it can be argued that the command does in fact have global status and that degrouping does in fact occur in the model. It will also have been noticed that, as a representation of the effector, the effector (Figure 1) does *not* have global status; there is simply no need for the representation to be global in such a control scheme. Hence the globality of the command is not based upon the globality of the representation of the effector; furthermore, the command is not stated in terms of a transformation of that representation, although it does lead to a set of local transformations. As argued earlier, the globality of the representation is inherent and implicit to the whole system's functioning; it is best recognized here with reference to the control scheme itself, which exploits a *constant* geometric property of the effector, a geometric fact that remains true for all possible movements. For the representation to be an effector representation, it would need to be a single *explicit* object that somehow encompasses within itself all the components of the effector in such a way as to render apparent the movement potential of the whole; in the present scheme this is not the case since it cannot be known *at the top level of command* whether or not a movement is possible: it first has to be tried out.

As for the model itself, it was only tested on paper. Had it been simulated on a computer, we would have had to transform the procedure into the language of analytic geometry (e.g. for finding

the intersection points of circle 0 and circle 1). However that appeared unnecessary in view of the simplicity of the underlying principle, and also unlikely to be of great import with respect to the larger goal of achieving a model of control for a humanoid effector. In that context, attempting to extend that principle to a more complex humanlike effector appeared to be the most promising avenue to pursue. This is the problem of generalization.

By introducing new d.f.s and/or new segments, the original 'egyptian' arm can become more sophisticated. There are many ways to do this, and many possible intermediaries along the road to a full-blown humanoid effector. The most straightforward is through the addition of a third segment with one d.f., like the first two. Another is through the addition of two d.f.s at the shoulder, one of them having its axis within segment 0 itself while the other having its axis perpendicular to segment 0's first d.f. Finally, the third segment (or hand) can be added, one more d.f. can be added at the elbow, and one more at the wrist, bringing the total to seven d.f.s. Controlling this effector was the ultimate goal; but, following our strategic argumentation, we did not tackle it directly. Rather, we moved on to consideration of the next simplest extension after the original 'egyptian' design, namely the effector consisting of 3 sffs with 1 d.f. each (cf. Figure 3).

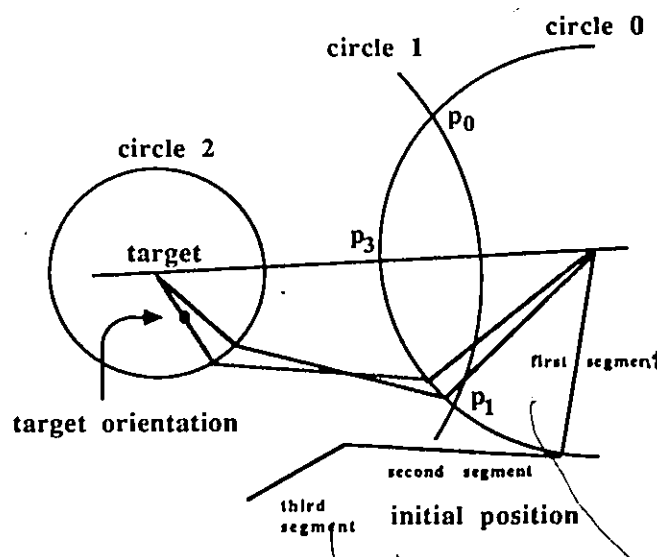


Figure 3. Problem of generalization of control scheme for the 'egyptian' arm.

Suppose the extended eff is in the initial position shown in Figure 3; suppose further that some target point ("target" in Figure 3) has been specified as the goal. As before, the problem consists of calculating the motions for the three d.f.s so that, following execution, the tip of the

third segment will coincide with the target point. Applying the control principle, one first draws a circle (circle 2) using the target as center and the length of the third segment as radius. But having done so a problem immediately arises concerning the center of the next circle, the one that is supposed to have as radius the length of the second segment ... : how will its center be chosen? But let us look for the moment to the possible fate of the first segment. Suppose that the third segment is completely extended with respect to its neighbour, the second segment, and frozen in that position. In effect this would transform the extended eff back into the original design. This fact allows us to specify a *range* for the final position of the first segment. Here is how. If a circle (circle 1 in Figure 3) is drawn with the target as center and with a radius equal to the sum of the lengths of the second and third segments, then in general it will intersect with circle 0 in two places (points p_0 and p_1 in Figure 3); again, if it does not, then the movement is impossible, etc. Now the final position of the first segment *cannot* lie anywhere on the right side of circle 0, defined by the arc traced in a clockwise direction from point p_0 to point p_1 , since from any one of those positions the 'second' segment, i.e. the second and third segments with the first being maximally extended, cannot reach as far as the target; if the movement is at all possible it must therefore follow that the final position of the first segment will lie somewhere within the complementary, or leftside, arc. This range can again be cut down to approximately half by using the fact that the second segment cannot move down below the length of the first segment, which it would have to do if the first segment's final position fell between p_3 and p_0 (cf. Figure 3). This effectively reduces the range to the arc p_1p_3 .

But we are still dealing with a range; and there is no way, by applying the same kind of consideration, in which it could be cut down further and further until a single point were found. The simple geometric truth of the matter is that there are *many* possible solutions to the local movement problem. (Two of them are shown in Figure 3.) Having exploited that aspect of eff geometry to the full, one is therefore forced to call into play *another* type of constraint in order to resolve the indeterminacy. Again many possibilities are open for consideration. The one illustrated in Figure 3 consists of specifying *within the command* something called "target orientation," which is simply a requirement that the third segment's final position should be such as to make a specified angle in (objective) space. (It could also take the form of a point, as shown in Figure 3.) This constraint is sufficient to decide upon a center for the second circle (not shown in Figure 3), whose intersection with circle 0 will provide a means to calculate the final position of the first segment. (Again, there might be unrealizable "target orientations," whose existence depends upon the possible non-intersection of the second circle and circle 0; they also cannot be known a priori.)

There are other alternatives to the above. For instance, one could seek to solve for local

motions by demanding a minimum *total* of local movements, or in fact by imposing just about any other similar constraint upon the set of local motions. But we have not explored any of them. The problem with these is that some justification must be found for the selected constraint (resolving the indeterminacy introduced as a consequence of the multiplicity of possible solutions also calling for the introduction of a further criterion). As above, this can take the form of the *optimization* of a certain relation; but such schemes can lead to complications, since other aspects of control, for instance path control, could potentially be affected by a particular choice and so must be taken into account in the justification process. However this is not the case for the "target orientation" solution.

But why does the attempt to generalize the initial control scheme lead to a problem of indeterminacy? One way to explain it is by pointing out the 'insufficiency' of the command in terms of the number of independent pieces of information it allows to be extracted from it: the extended 'egyptian' arm has 3 d.f.s, which means that *three* local values must somehow be derived from the command: however the command (i.e. the goal point) only has *two* independent dimensions. Hence the indeterminacy. This did not of course occur in the original case, where there was a unique solution from goal point to d.f.s, since the number of dimensions in the command was equal to the number of d.f.s in the eff. But here one is missing. The situation can be likened to the well known algebraic fact that it is impossible to solve uniquely a set of two equations in three variables. In order to arrive at a unique solution it is necessary to introduce a third constraint that is *independent* of the other two, just as we suggested above. In fact the *only* requirement for that third element or relation is that it be independent of the other two; outside of that, it can be practically anything. However it is clear that the particular choice of third element will directly affect the (unique) solution derived from its association with the two other relations, a different third relation giving rise in general to a different unique solution.

As it happens the deficit occurs on the side of the command. But the solution need not necessarily be sought at that level. As in the optimization scheme, it can pertain to the *execution* of the command; but the third relation can also become part of the command itself, as in the orientation scheme. Although the former is more difficult to find and implement, it has the potential advantage of leaving the command simple; the latter one is much easier to implement, but it requires a more complex command. Whichever type of solution is adopted, and for whatever reason/s, the bottom line is that there will need to be embodied in the control system *at least* as many independent constraints (or constraint dimensions) as there are d.f.s in the eff.

In actual fact this requirement is also independent of the *type* of control scheme, be it extrinsic as we have here, or intrinsic, as we shall see below. In the present model the search for a solution to the problem of indeterminacy is directed at the level of command or execution, and not at the

level of the representation of the effector: this is because given that the extrinsic model is founded upon a certain representation of the effector and that the solution is sought *within* the confines of the extrinsic model, the search cannot by definition be brought at that level. There *may* be other possible representations of the effector for an extrinsic type of control model, in which case the previous statement would no longer hold; but we know of none. (This is why the extrinsic model and its representation of the effector can be used equivalently.) Hence exploring the situation at the level of the representation of the effector entails switching to an intrinsic type of control model. This is the subject of the next sub-section.

2.2.1.1 An intrinsic control scheme

Our second model was also aimed at the control of an extended effector (Figure 4). Rather than having a third segment plus an extra d.f., it has two more d.f.s than the original, making it an eff capable of accessing three-dimensional space. In terms of complexity it therefore lies in an intermediate position, between the previous extended 'egyptian' arm and the full humanoid arm having 3 segments and 7 d.f.s. Following approximately the human proportions, the length of the second segment is equal to that of the first.

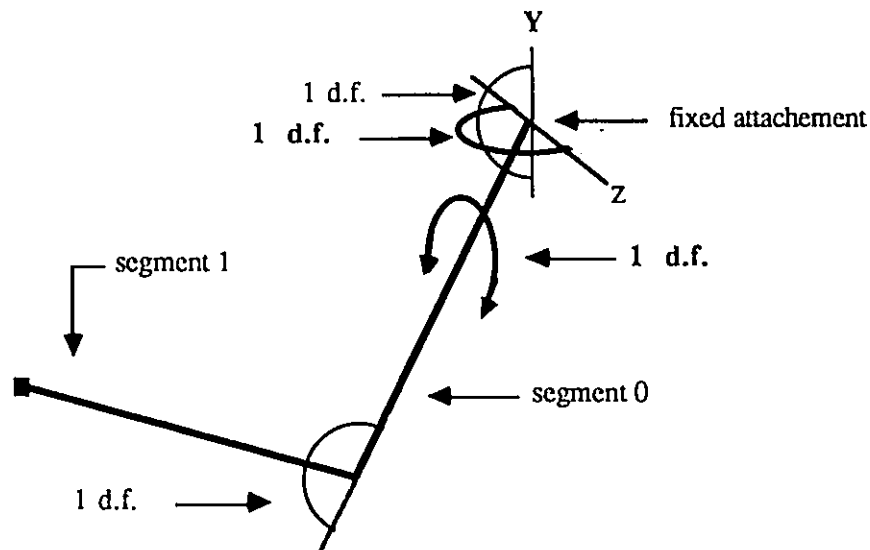


Figure 4. *Extension of the 'egyptian' arm into 3-D: the two new d.f.s are indicated in bold style. One occurs about the Y-axis, the other about the length of segment 0. The shoulder is now a ball joint. The elbow is left unchanged from the original (cf. Figure 1).*

If, by removing the d.f. occurring about segment 0, we only had three d.f.s in that eff, then it would be possible to 'extrinsically' solve uniquely for the required local translations since the command would take the form of a goal point in three-dimensional space; the number of values in the command and the number of values required to execute it would be equal. (We will not go into the details of how this could be done.) Hence the problem of indeterminacy would not arise, just as it did not arise in the original control scheme of the previous sub-section. But the new eff has four d.f.s, and the indeterminacy problem therefore still exists in the present case *if it is supposed that the command continues to take the same form*, namely a goal point, or even a path. Since the minimal humanoid arm has more than three d.f.s, and since the ultimate goal is the control of a humanoid arm, the problem cannot simply be solved by keeping the number of d.f.s down to three.

The command will not however retain that form when the mode of control becomes intrinsic; in fact *it will be such that the problem of indeterminacy no longer arises*. As argued in sub-section 2.2.0, the whole point of intrinsic control is that commands are directly based upon the effector representation. The whole point of the indeterminacy problem is that the number of dimensions of control was insufficient with respect to the number of d.f.s (or independent 'dimensions' of potential movement) in the eff. Since now the commands are to be based upon the effector representation (which it is not in the extrinsic mode of control), and since that representation 'accounts for' the full movement potential associated with the effector, then *by definition* the command cannot be insufficient with respect to the object of control. This is the first major advantage of an intrinsic mode of control. By doing so, however, the explicit and immediate reference to objective space *in the command* is lost; at the present stage this is the main disadvantage, but as we shall argue later the eventual reference to objective space is not as a consequence irretrievably lost.

Since in the present case the effector has 4 d.f.s, the representation of the effector will need to provide the system with four potential dimensions of control. But that is not the whole story. In sub-section 2.2.0.0 it was also argued that the representation should constitute a *single* explicit entity. In the previous model there was no such entity; of course there *was* a representation of the effector, but it was not available as such (i.e. as it is to us when we look at it) within the system: what *was* present was a fixed structure of computation that reflected the structure of the effector, but the eff itself remained implicit. Here also, insofar as the representation will dictate an unchanging pattern of computation, it will be fixed. However, that is the *invariant* part of the representation, a frozen core that nevertheless leaves room for variation; these sources of variation constitute the basis of command, while control itself is a function of both the invariant *and* the variations allowed within it.

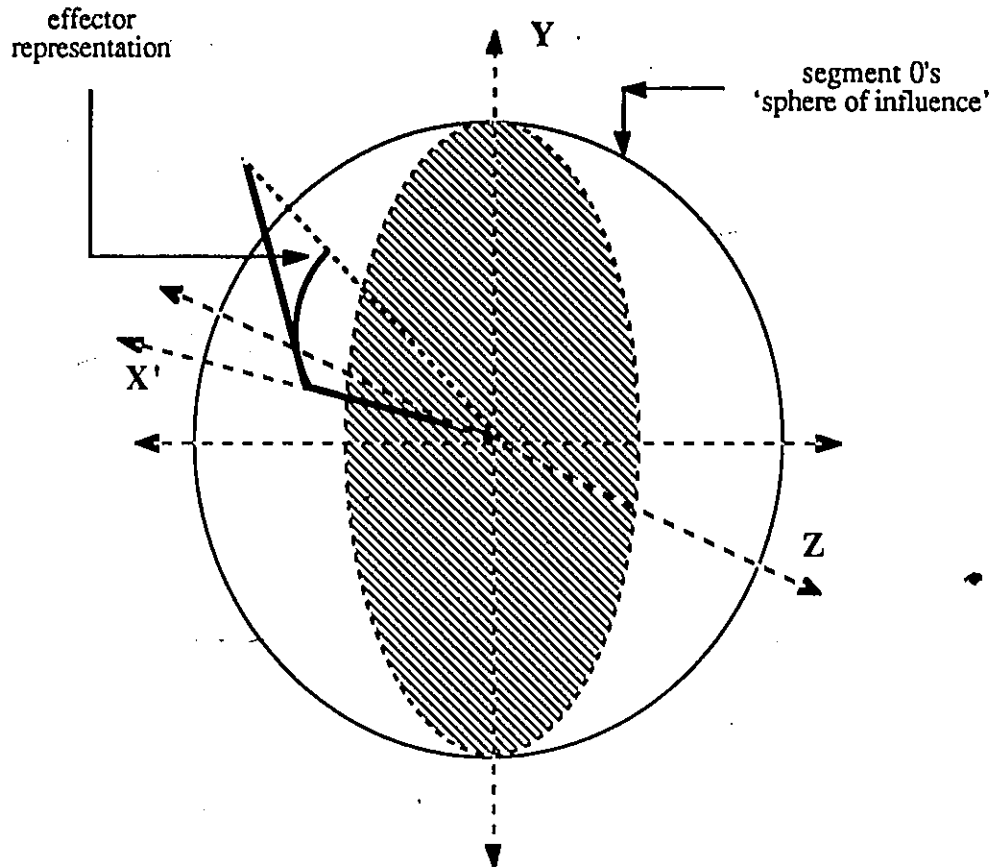


Figure 5. *Geometrical basis of the 4 d.f./2 sffs effector representation. The second segment is projected onto the first segment's 'sphere of influence,' the resulting great arc is the effector representation.*

Figure 5 describes the geometrical principle behind the effector representation for the 4 d.f./2 sffs arm. Consider segment 0 first. Its two perpendicular d.f.s (about the Y- and Z-axes) allow for the definition of what might be called its 'sphere of influence,' in much the same manner as in the first model, segment 0's single d.f. allowed to define its 'circle of influence.' Since segment 0's length is fixed, for any position of the arm the elbow will always lie on the sphere, as shown in Figure 5. Now this virtual spherical surface is used as a projection surface, with the center of the sphere, that is the shoulder, acting as projection point. Now segment 1 is projected onto the sphere. This results in a "great arc" (i.e. a segment of a great circle of the sphere); *that arc is the effector representation*; it is the 'internalized' single object of control upon which the whole system is based.

This geometrical entity has some interesting properties; let us look at some of them. First, it

is an *invariant*; that is, whatever position is occupied by the whole effector (within the bounds allowed by the d.f.s), the projection *always* results in some great arc. Now suppose that some movement occurs in the two perpendicular d.f.s at the shoulder while the two other d.f.s do not change: the arc will move on the surface of the sphere, but it won't change in length or orientation (in spherical geometrical terms). Now suppose on the other hand that everything is fixed except for the elbow d.f.: the arc won't move but its length will vary from 'zero' (when segment 1 is completely extended) up to one quarter of the circumference of the sphere (when it is maximally flexed). Finally, if only the rotational d.f. (about segment 0's axis, X' in Figure 5) is allowed to function, then the arc will rotate about the 'elbow' point on the sphere.

The point of it all is that the representation is a single entity (namely an arc) that nevertheless, through the possible variations upon its characteristics (namely position, length, and orientation on the sphere) can be associated with the state of the whole effector at one moment. These characteristics are also independent of each other, a variation in any one not implying a variation in any other.

But how can this be of any use in a control scheme? It is interesting that this constant property has been identified, but how could it be turned into a basis for command? Quite simply, by reversing the process, i.e. by having the system *produce* these variations and compute the implied motions. This is possible because of the permanent relation that exists between the characteristics of the representation of the effector (the arc) and the components of the effector; by extension, any variation (or motion) in one or more of the arc's *characteristics* should produce a movement in the corresponding effector component. In fact these changes will become the *motion directives*, i.e. the commands, in the new control system. The system of course is provided with the effector representation, as well as a means of prescribing various combinations of modifications of the eff's characteristics.

— At this point, one should easily be able to start imagining the command, and hence the movement, potential opened up by such a scheme. In effect though, it is merely a starting point, since it is not evident, even for the individual variations, what kinds of *patterns* of modifications will be meaningful: these will have to be uncovered by further study. However that is the price of an *intrinsic* model of control, the gain in 'global control' power being counterbalanced by a loss in the a priori knowledge of the objective results of the commanded motions. Yet the latter is not beyond recovery; it is merely not yet available at this stage of development. What will happen however is that a meaningful intention will take a quite different form than it takes now in terms of objective space (as described for instance in Robotics). This is because the system defines a new basis for spatial knowledge, a knowledge that must now be developed from the 'primitive' spatiality of the effector representation implemented within it. This is exactly the same kind of recursive argument

encountered in chapter 1 concerning the basis for knowledge of visual spatiality (cf. Lamontagne & Beausoleil 1982) provided by the 'primitive' analysis of line segments; and while it is conceivable in principle, we also could not yet state what form an eventually more extended knowledge of the visual world would take.

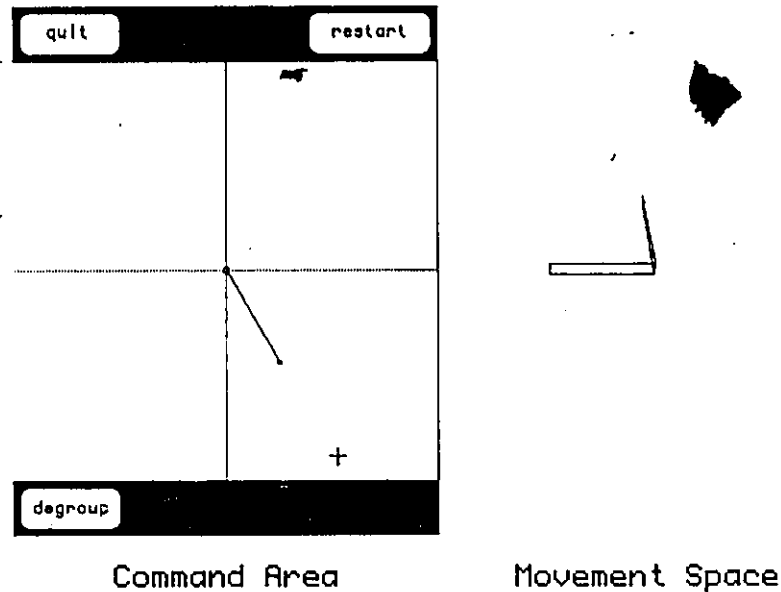


Figure 6. *Screen display of the simulated arm and control scheme for the second model.*

The control system is provided with a representation of the movement potential of the whole effector; this is the effector representation. Control and production pertain directly to that object through the specification of patterns of modifications upon its characteristics. In order to study the command and movement production potential of the system, and more specifically to get a better grasp of the "patterns of modifications" question, it is necessary to simulate the system. In effect the system as it now stands defines a sort of 'level zero' of control upon which further levels will be added. The problem of control is redefined at each step; further dimensions of control can only be directed at existing dimensions of control. This presupposes that, taken singly or collectively, the existing dimensions of control leave room for further specification of transformations. Thus, a simple modification along one characterization dimension of the effector representation constitutes a 'level zero' command; however such a modification can occur in a variety of potentially interesting ways, or according to different patterns. Before these are 'implanted' within the system however, they have to be found out. The main problem that had occupied us so far was the problem of the

effector representation. Even though it concerns the representation of the movement potential associated with the effector, it could be tackled statically, i.e. in the absence of movement; but now this is no longer possible since further levels of control no longer pertain to the effector but to the control of 'level zero' control. This is why simulation is necessary.

The first step of this process consists in programming the effector. Since the arm moves in and accesses three-dimensional space, it is necessary to use some of the techniques of computer graphics (Foley & Van Dam 1982). Appendix III fully describes the implementation. It will be noted that the simulated effector, having 3 segments and 7 d.f.s, is the complete humanoid arm. However here the effector representation, and hence the control, are limited to the first two segments and four d.f.s; this does not interfere in the simulation of the control scheme: the three remaining d.f.s and the third segment are simply 'ignored' by the system, receiving null motion directives. Lower down we will turn briefly to the problem of the effector representation of the full humanoid arm.

The second step of the process is the implementation of the control scheme. The only real difficulty that this part of the simulation presents centers about the nature of the "user interface," i.e. the type of interaction with the system dictated by the command mode. These concerns are not critical for testing the model; but as argued above they can become quite important for a further study of its extensions. So far for instance there are no facilities in the programme for generating patterns of modifications, as would be required for a study of the problem of path control. A secondary problem of the implementation is speed; the calculations are too lengthy to obtain real-time movement: everything has to be pre-computed. This limit to the ease of use is also an important impediment to an investigation of the full range of commands available in the system. Finally, since it is better to see *both* the command and the resulting movement simultaneously, the screen is split in two areas; but this reduces somewhat the quality of the 3-D display.

The implementation of the control scheme does not follow exactly the geometric principle (Figure 5) underlying it. For instance the *length* of the projection segment is made to correspond to the degree of *extension* of the forearm: the larger the length, the larger the extension; geometrically, the converse is true. Furthermore, the *orientation* of the projection segment is interpreted by the system as a value for the rotational d.f., not as a position for the forearm; in actual fact these are 90° apart. But the *position* of the 'lower' tip of the projection segment is 'correctly' interpreted as corresponding to the tip of segment 0. Of course neither of the above deviations encroaches upon the basic principle of the singleness of the object of control since, following the rule of representation, there is no a priori necessity of a strict allegiance to the 'objective' situation, merely that the system's *interpretation* of that object (i.e. the association of its features with certain parts of the effector) be unchanging.

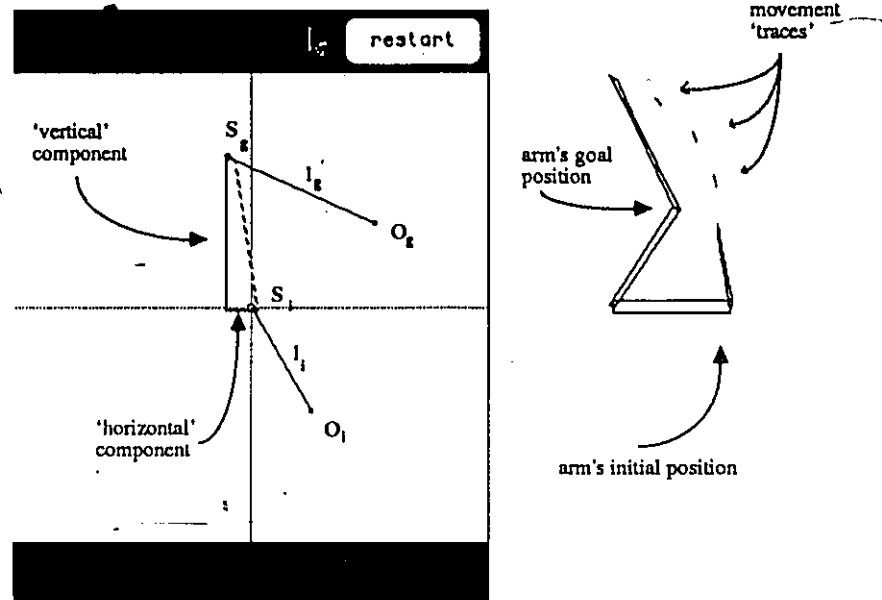


Figure 7. Example of command (on the left) and the resulting movement (on the right). The "i" subscript indicates the initial values of the control dimensions, while the "g" subscript indicates the goal values. See text for further details.

Let us look more closely at the implementation of the control (Figure 7). At the beginning, the eff is in the "initial position," the same as that indicated in Figure 6 above. There is no special reason why that particular position is chosen as the initial position. The "user interface" for commanding motion is set up in such a way that the first specification (notified by positioning the pointer (the small cross at the bottom of Figure 6) and "clicking" the "mouse" (the piece of hardware used to move the pointer on the screen)) will be the new position of the tip of segment 0; this corresponds to the small circle in the Figure. Next, the pointer is moved to another position and the same procedure is repeated: this time however the specification concerns the projection point of the tip of the second sff, represented by a small dark square in the figure. The programme then joins the two "points" together by a straight line, thus completing the process of command. The new 'line' consists of the *goal position*, while the original one is of course the initial position (see Figure 7). At this point the programme takes over and starts to degroup, i.e. to compute the local motions on the basis of the commanded modification of the effector representation. (The "degroup" button shown in Figure 6 is there because the programme also allows for more than one transformation specification (up to twenty in fact), in which case one needs to tell the machine that the command sequence is complete and to start computing. However it is not yet possible to edit a

given command once it has been set down. Clicking the "restart" button interrupts the execution of a command sequence; it signals the system to start afresh.) A default value of five positions, including the initial and final positions, are computed by the programme in order to show the movement of the effector as it progresses from one to the other; this is displayed on the right half of the screen, but only, as mentioned before, once they have all been computed.

Here is how the system computes the local motions. Let us suppose, in order to simplify the exposition, that it computes the complete motion, from the initial to the final position, in just one step. The programme has in hand three pairs of values, namely $\{s_i, s_g\}$, $\{l_i, l_g\}$, $\{o_i, o_g\}$. In order to extract from the first pair of values the two motions directives for the two d.f.s at the shoulder, it decomposes the vector $s_i s_g$ into its two independent components, one 'vertical' component for the 'vertical' d.f. at the shoulder, and the other 'horizontal' component for the 'horizontal' d.f. at the shoulder. These lengths are directly interpreted as angles; each one will be 'fed' as such into its corresponding d.f. Next, it proceeds to compute the difference between l_g and l_i , a difference that is again interpreted as an angle: the interface does not allow for the specification of a goal segment whose length surpasses 180 units. This value is interpreted as the commanded increment for the elbow d.f. Finally, the orientations of the 'command segments' are compared, i.e. $o_g - o_i$ is computed; this is the motion directive for the rotational d.f. at the shoulder. As in the previous case the interface does not allow for the specification of a change in orientation beyond 180° ; in fact it is only possible to position the command, or goal, segment so that its orientation falls between $+90^\circ$ and -90° . This reflects, as in the other cases, the physical limit of the rotation d.f., but in this case it also reflects the position of the d.f. in the arm's design. The decision to place it there was again made on the basis of the human arm's structure.

In effect then, the simulation implements the three levels of the whole control scheme: [1] the global level of the single eff, [2] the intermediate level of the two sffs, and [3] the local level of the four d.f.s. Although the commands themselves *seem* to be specified at the level of the sffs, it should be remembered that in fact *they bear upon the characteristics*, namely the position, orientation, and length, *of the global object*, namely the segment; the *only* way in which the movement potential of the 'local' components or lower objects of control, i.e. the sffs and the d.f.s, could be represented is through (secondary or *variable*) characteristics of the global object of control; as for the primary or invariant nature of that object, it is based upon the constant relationship between the components of the eff (as described and 'captured' through the projection principle) which is embodied in the immutable sequence of computations described above.

The local (d.f.) motion directives are computed as a change in specific attributes of the global object of control. Now the geometry of the effector projection is non-linear. More specifically, the

relation between the extension of the forearm and the length of the projection segment is non-linear. Nevertheless the system interprets it as a linear function; thus, for *any* position of the forearm, a given increment (or decrement) of the command segment's length is interpreted as implying the *same* displacement in the extension d.f., while in fact it should not. (The demonstration will not be given here.) For the purposes of command this has little consequence; but it could not be ignored in a more systematic study of the relationship between the spatiality of the effector (as represented in the command area - cf. Figure 6) and the behaviour of the effector in objective space (as represented in the movement space - cf. Figure 6); this would require the introduction of an inverse mercator 'correction' (i.e. a projection from a plane to a sphere) in the system. However even if the relation between the command segment's lengthening and the implied displacement at the elbow d.f. were established correctly in the system (i.e. strictly according to the geometrical principle of projection, in which there is a non-linear relation between the elbow's angle and the projection segment's length) the non-linearity of the effector tip's movement in objective space would still persist. This is because the non-linear relation between d.f. excursions and effector tip movement is an inescapable fact of effector kinematics.

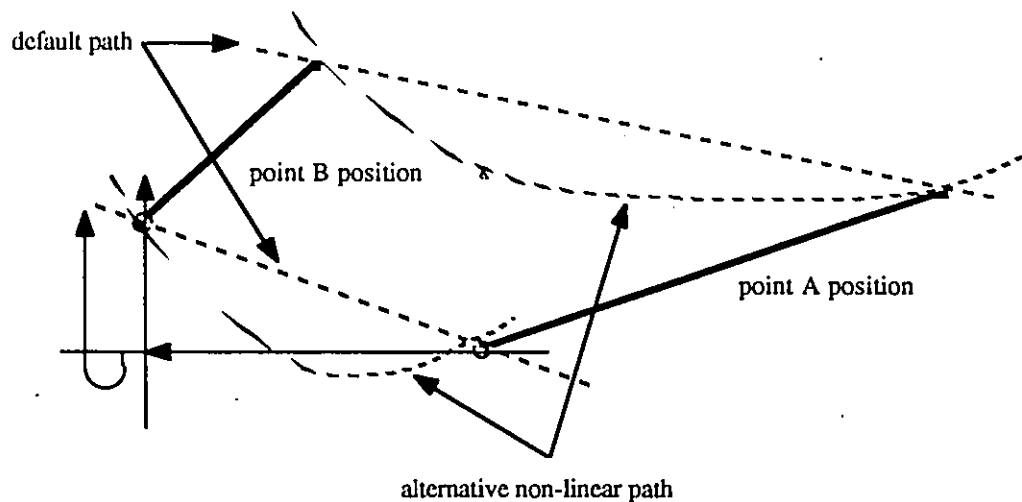


Figure 8. *Path control. The default is a linear computation of the intended change, leading to constant velocities for d.f.s. An alternative non-linear path is outlined; note how this can imply a reversal in the direction of movement of a d.f. (here the vertical one). All movements still occur within a fixed interval of time.*

The crux of the problem of effector control however is not that this relation exists but that it should be 'captured' (or represented) and exploited in such a way that, with respect to objective

space, a *given* result can be commanded. In its current state the model allows for a very large number of results to be achieved. However they all have the same status; they are undifferentiated and unorganized. On the other hand, the control schema leaves a lot of room open for the introduction of further constraints or additional dimensions of control needed to achieve such a differentiation. One of the most evident concerns the possibility of path control.

At this stage of development of the model the local motion directives are computed as simple *linear translations* in time, i.e. as constant velocities. (However since in general the calculated displacements are unequal but the time interval for the whole movement is fixed, the (constant) local velocities vary from one d.f. to another.) This is a default result; constant local velocities are computed simply because there is as yet no alternative. But what would happen if they were allowed to change (within the confines of the fixed overall movement time)? The answer is that the system would now have more control over what happens *between* the initial and the goal position.

Let us take up the case of a movement (in objective space) from point A to point B (Figure 8); that is, let us suppose that the initial position of the effector is such that the tip of the arm is at point A in objective space (point A position) while the tip of the arm in the goal position is at point B in objective space (point B position). As described above, the model will compute a linear change for each d.f. (the dotted straight lines in Figure 8; orientation differential not explicitly indicated); it will then proceed to compute different constant speeds for each of them. This will provoke a certain displacement of the whole arm that will effectively bring the effector tip at point B. But these computed changes are only one among many potential ways of 'getting from A to B.' It is possible for instance to travel a different (no longer linear) 'path' from point A position to point B position (the dotted curves in Figure 8). What would happen if it were followed? *The arm would still end up at point B position, but it would get there by a different route; furthermore each d.f. would no longer have constant velocities.* Thus is opened up the way to a study of path control. Now let us reverse the situation. Suppose for instance, following the robotic model, that somehow the arm is externally taken from point A position and brought to point B position *through a linear translation of the tip of the effector*. Suppose furthermore that at a large number of constant smaller time intervals the arm is projected onto the virtual sphere associated with the first sff. The resulting series of projection segments would 'fit' between point A position and point B position; and it would follow a certain non-linear *pattern*. If now this sequence is taken as a command sequence, then by definition the path of the tip of the effector will (ideally) follow a linear course in objective space. Reverting now to the initial situation, imagine that this were done for another linear translation of the tip of the effector. A certain pattern would again emerge. And so on for any number of trials. The question that now arises is: will there be something *common* to all members of the resulting set of projection sequences? Is there a '*linear command*' invariant, or are all linear transla-

tions independent and otherwise unrelated to each other? We believe that a linear invariant can indeed be found; but we have not yet identified it.

Recall that our initial task, namely the definition of the effector representation, also centered about the problem of specifying an invariant. As global object of control, it had the property of seizing the movement potential of the *whole* arm. From the point of view of production, this meant that the representation would remain consistent with any and all movements, since it retained its essential or primary character *across* all possible positions of the effector. The problem of the linear command invariant can be defined in exactly those terms, except that now it concerns not the effector representation but the *transformation* of the representation (which is why it is called a *command* invariant). Of course, whatever that transformation turns out to be, it will *also* need to be independent of position, i.e. remain true *across* all possible (starting) positions of the effector. This is where the idea of lsff (i.e. leading sff) could well enter critically into play, for the problem is essentially defined with respect to the movement of the tip of the second sff. It is true that such a linear translation is impossible with the first sff alone; however supposing that the eff consisted of three sffs, as the full humanoid arm does, then the choice of which 'end' sff's tip, the hand or the forearm, would translate in a linear fashion must be made. In that sense alone, the lsff would be called into play as a further important motion characterization. This complicates even more the problem of linear command equivalence, since the transformation will have to be applicable to both possible instances of lsff.

One clue to a possible approach to the problem of the command invariant is already present in Figure 8. In general the linear command invariant (for instance) will involve transformations of *all* control dimensions. So far, and this is true in the simulation as well, the commands are specified along the *individual* dimensions of control. Taken as a whole, these particular modifications amount to a certain pattern of course, but one that is not well characterized, since no explicit relationship *between* the various changes is called for or imposed. In order to approach the problem of the command invariant, we must [1] be able to specify (i.e. characterize) such collective constraints, and [2] have a way to proceed from them to the computation of the implied local changes (i.e. the changes at the level of the individual control dimensions). In effect this means that a new *level* of control is called for, one whose object is the more primitive level of control with which we have dealt so far. What we want, for instance, is a way of describing something like the "alternative non-linear path" of Figure 8. As stated earlier, we are only yet at the stage of analyzing this problem, but there is one aspect of the of the situation which already deserves to be singled out.

There is a potentially interesting aspect of the effector representation which we have so far not explicitly noted, and which can be stated in the following intuitive manner: because the representation is a single 'bound' entity, it is possible to 'take a hold of' the *whole* object on any one of its

dimensions (i.e. its 'lower' end-point, its length, and its orientation). Thus if one seizes the 'lower' end-point (the end-point marked with a small circle in Figures 6, 7, & 8) and 'pulls' on it, the whole segment will follow; similarly for the other two (although less evidently or clearly for the last one). This suggests that the specification of a pattern of transformation could be 'constructed around' one outstanding dimension of the object, whose transformation would be explicitly described, while those of the rest would simply 'follow,' or assume default transformations. For instance, one can imagine that in pulling the 'lower' end-point of the segment in a certain direction, the segment's orientation would automatically align with it, while the length would remain unchanged, or also change as a function of the chosen direction of motion.

Even though this is true for all dimensions of the object of control, some of them appear to be more interesting than others, in particular those that correspond to a sff rather than simply a d.f., such as the orientation. As it happens, focusing the motion directive/s on the lower end-point amounts to selecting the first segment as the lsff; furthermore, *that point corresponds to the tip of the segment*. Now what 'part' of the object would lead to selecting the second sff as lsff? *It can be neither the length nor the orientation*, because even though the length of the object corresponds to the extension d.f., which is the second segment's defining d.f., the orientation, even though it properly belongs to the first segment, has its effect upon the second segment, *and none on the first*. (Incidentally, this is also the case for the rotation d.f. in the forearm, which exercises its effect upon the hand rather than upon the forearm itself, with which, of course, it is by definition associated.) In order to make the second segment the lsff, it is therefore necessary that this d.f. also be taken a hold of, along with the extension d.f. *The part of the representation that would make the second sff the lsff is the upper end-point* (marked with a small square in Figures 6, 7, & 8). *This is the projection point of the tip of the second segment*.

Of course since the forearm is attached to the upper arm, the second segment cannot, in general, be commanded alone, as is the case with the first segment which, in our model of the effector is attached to none other. But this is not a problem. The main point about making the forearm the lsff is that the command can be focused on it, while the upper arm will simply 'follow up' on its motion. Indeed *this seems to comply with our intuition that the lsff actually 'pulls' on the rest of the effector*, an intuition that is in apparent contradiction with the fact that, at the level of peripheral execution, it is the lsff that is actually pulled, directly or indirectly, by the rest of the sffs. Such an observation on the problematical difference, or apparent clash, between the nature of a motor intention and the nature of peripheral execution has of course been made time and again; our model however allows reconciles them.

In the context of the full humanoid arm (with 7 d.f.s and 3 sffs), the problem of selecting the hand as the lsff and of 'focusing' the command about it will also arise. However before it can be

tackled, the effector representation for the humanoid arm must have been defined. Again, this is the problem of generalization (cf. sub-section 2.2.1.0), about which, at this stage, we have only a few preliminary remarks to offer.

One immediate avenue to exploit is of course the principle of projection; the idea of projecting the *two* sffs, namely the forearm and the hand, upon the first segment's 'sphere of influence' is a natural extension of the earlier situation. However this leads to some problems, chief among them being the fact that it is no longer possible to associate quite as clearly as in the above case those 'parts' of the new representation with the actual eff. Consider for instance that this projection would in general give rise to an additional projection segment connected to the 'upper' end-point of the original projection segment. Now as the hand moved about, the length and orientation of this segment would vary. But neither the length nor the orientation of this little segment unequivocally 'reflects' the three d.f.s of the hand, namely the two 'wrist' d.f.s and the forearm rotation d.f. (which, as argued earlier, although it belongs to the forearm, is functionally associated with the hand): that is, a change in orientation could just as well result from a movement at the wrist as from a rotation of the whole hand; similarly for the length. (In actual fact, a change in *both* length and orientation will result from the action at *both* the 'wrist' and the rotation d.f.s.) But there are other possibilities; for example the actual connection point, having two d.f.s, could be associated with the two 'wrist' d.f.s, while, as in the previous case, the orientation could be tied with the rotation d.f.; concerning the first pairing, a problem immediately arises in that it would be impossible to command a movement of the hand alone: this will not do. Instead, the *other* end-point could be paired with the two 'wrist' d.f.s. Even though this would probably work, it is somehow unsatisfactory, or 'unnatural.' For instance one of the possible problems with such a scheme is the actual characterization of the global object of control. Although it would be difficult to substantiate, at this stage we tend to favour a simple representation of the effector, and to put the computational weight of control on the command side rather than on the representation.

An alternative way of applying the projection principle to the humanoid arm is to project the *hand* onto the surface of a virtual sphere associated with the *forearm*, the center of that sphere being at the elbow. However, because of the structural difference of that part of the effector with respect to the earlier upper arm-forearm effector, the same ambiguity as described above still exists here: that is, motion in the 'wrist' d.f.s will (in general) produce a change in the orientation of the projection segment, but a motion in the forearm rotation d.f. will also lead to the same result. Thus, the hand has *two* orientations, one being 'artefactual,' i.e. not directly commanded by the system, and consequent upon motion in the 'wrist' d.f.s, the other being 'real,' i.e. commanded directly by the system, and consequent upon motion in the forearm rotation d.f.. These two will somehow have to be dissociated in the representation.



It will be noted, finally, that the alternative scheme would apparently give rise to a *partitioned* object of control; that is, contrary to the first scheme, the representation would no longer consist of a single (geometric) entity. But this is only apparently the case, since the second 'part' of the representation is linked up with the first through the forearm, which is common to both 'parts.' On the other hand, as argued above, the unity of the representation is founded upon the embodiment in the system of a *fixed sequence of computations*, that sequence being dictated by the nature of the representation, in which is assured the singleness. That will presumably not be different for the eventual representation of the full humanoid arm.

With these comments, our presentation of the intrinsic control model for the extended effector is complete. We have redefined the general problem of movement production, placing it in a new context in which the major 'local' problems have been clearly identified. First and foremost among them is the achievement of an effector representation, since everything else depends upon it. The effector representation itself then sets the stage for amplifying and deepening control, from a straight positional change command mode to one where a path, i.e. a pattern of positional changes, can be specified. In effect each level of control constrains the implementation of the next one. We have only hinted at the full potential of control opened up by the present intrinsic scheme. However the problems are clear, and we believe that they are worthy of further attention.

To conclude the present chapter we shall next take up briefly certain key points of our model in the light of issues that were raised earlier in performance- and system-driven research.

2.3 Extensions of the model

Very few objections would probably issue from performance-driven, as well as system-driven, research concerning the idea that a model of movement production should be of the intrinsic type, even though many of the models proposed there liberally make use of the reference to objective space, either with respect to the movement itself or the result of the movement (cf. Schmidt et al. 1979; Fitts 1954); the problem however lies not in the reference itself but with the nearly total absence of concern for the *functional* nature of its representation. (This concern is of course at the heart of technology-driven research, but its models, as we have seen, are of the extrinsic type, a feature that might in part explain the apparently low level of relevance found within performance- and system-driven research for the robotic models.) Indeed the most general problem in performance- and system-driven research derives from the task of defining a theoretical context in which intrinsic production could be conceived. *This, we believe, is what the intrinsic scheme described above attempts and achieves to a certain extent.* The primitivity of the control model itself should not be seen as limiting; on the contrary, as stated above, it opens up clear avenues for further enhancement. We are not yet at the level at which the productions of the model can equal or rival the

performance of humans; yet it is already conceivable that they eventually *could*. Unless some critical fault is found with the present scheme, the major problem is no longer the definition of a context for intrinsic production; rather the major problems now lie with the further development of the model in the direction of a human level of performance. Nevertheless this is by no means a trivial matter.

The pivotal conceptual component of the intrinsic scheme is the idea of a global effector representation. While it grew out of a preoccupation with the construction of a *working* (or computational) model of production and not out of a primary preoccupation with performance, although reference to performance played a major rôle in its development, we would like to point out that this idea of a global effector representation can, in the context of mainstream performance theory, be interpreted as constituting a hypothesis. Thus, it *might* be possible to derive testable empirical consequences from it. For our part, we believe that this should be based on the model itself, which is an attempted, albeit incomplete, instantiation of the principle; but the model is probably still too primitive to allow for that. However this does not preclude the possibility of bypassing the particular model and deriving empirical consequences from the principle itself, or indeed of formulating a different instantiation of the principle.

With respect now to the principle of the instantiation, namely *projection*, a similar statement about possible empirical consequences could be made for system-driven research. It concerns anatomy. What we are about to propose does not derive directly from the model itself; rather it is based upon a possible application of its underlying principle, namely *projection*, to movement production as defined in the context of natural systems. As above, therefore, it is an extension of the model, rather than a strict attempt at linking the details of the model with the details of the specific problem-context.

One of the most remarkable aspects of periphery-bound processes in natural (i.e. nervous) systems, especially the visual and movement systems, is the fact that the whole set of local peripheral components is at some stage *mapped* into the cortex. Although in vision some functional reasons to explain the nature of the retinal mapping (cf. chapter 1) and some ideas on the possible functional exploitation of its properties have been adduced (e.g. Schwartz 1980), very few such propositions are encountered in system-driven research about the "motor homunculus." One does come across the idea that the reason for which the thumb has a cortical representation as large as that of the forearm and upper arm put together is because the thumb is involved in more intricate and fine movements than the two other segments, the implication being that more work, and hence more cortical matter, is needed to 'compute' these fine movements and engage in these minute corrections. This could well be the case, even though little in the way of a convincing demonstration of this idea has been forthcoming.

But there is another possibility. It is known that the retinal surface is not homogeneous with respect to the density (as well as type) of receptors found on it, the central part, or fovea, having a much richer colony of sensory cells than the outlying region. It therefore makes sense, since there are more cells in the fovea, that in the retinal projection to the cortex the fovea should occupy a larger area than the outlying regions. Thus it is not necessary to conjure a *functional* reason to explain why the fovea is better at fine or intricate perceptions; the *anatomical* argument suffices. However the story would be different if the visual cells *were* homogeneously laid out on the retina and the projection of the central part *still* had the largest representation: in that case, some central functional reasons would need to be put forth, ideas probably very similar to those encountered in system research about the homuncular mapping. However it is not clear whether or not the homuncular projection is in fact homogeneous.

We would simply like to suggest, contrary to the retinal case, that there might be a way to explain the particular distortions of the homunculus *without* calling into play the non-homogeneity of the mapping *nor* the classical reason of the larger functional weight necessitated by the production of intricate and fine movements. The idea rests upon the fact that the retina is a *surface* whereas an effector (or at least the disposition of the muscles about it) is a *three-dimensional* structure. In the previous sub-section we described a particular mapping of the effector onto a virtual sphere, a procedure that introduced some distortion in the image with respect to the original projected object. *This distortion was a function of the mapping procedure alone.* Now suppose that instead of mapping volumeless linear segments, such as the ones we used previously, we were to map cylindrical linear segments; and suppose that the projection 'lines' had to travel *inside* the effector components, i.e. inside the cylinders. If the projection occurs on a *flat* 'cortical' surface then one can imagine that distortions will again be introduced, the exact nature of which will depend on two things: first, the topology of the effector and second, the projection rule. (Concerning the latter, the fact that the projection lines occupy a certain volume is especially important.) Given this intuitive picture of the situation, one could expect for instance that the *more* intricate their inherent topology, the *more* distorted their projected image will be. (Thus the topology of the hand's fingers is more complex than that of the foot's five toes; while the toes are all parallel and joined along a 'straight' line, this is only the case for four of the five fingers, the thumb being displaced downwards.) Of course this is merely a suggestion; but it could be substantiated through a simulation. On the other hand, it might be possible, by modifying the mapping rule/s so as to *reproduce* the motor homunculus, to discover the nature of the corticomotoneuronal mapping.

Finally, we would like to briefly comment upon one aspect of the problem of representation as it pertains to all three approaches to movement. In order to do so we shall center our discussion around the question of path control. As was often mentioned in the course of the present essay,

given the fact that movement is a *physical* phenomenon, the temptation is great to suppose that the representation of movement (in the sense of the internal knowledge presumed to be necessary for its production) is also of a 'physical' nature. In that respect, the relation between one (objective) description of the product and the supposed internal prescriptions (i.e. the representation afforded by the system) needed to produce the same movement is a one-to-one relation. Thus, one would seek to uncover the nature of the representation by studying the physics of the objective product. However this assumption of a one-to-one relation is a very strong assumption; after decades of intensive effort, it has still not opened the door to the nature of the 'higher' levels of movement representation and command. But we are forgetting something crucial about representation. The only absolute requirement for a representation is that it establish a relation, directly *or indirectly*, between its own 'behaviour' and that of whatever it represents; it is *not* necessary for the representation to 'behave' in strict accordance with a description of that which it represents nor to have a *direct* access to all of its aspects. (cf. Craik 1943)

It will be recalled from the above description of the second control model, how, in the default path mode, the system produced (in general) varying *speeds* for each component. However there is no explicit 'component speed' directive in the system; rather, they arise as a *consequence* of the structure of motion computations: the system merely treats the space and time dimensions in such a way that [1] speed is produced, and [2] *different* speeds are produced. Stated another way, 'speed' is built into the system, but it is not explicitly available to it; nevertheless it is an important dimension of the objective description of the produced movement. (Presumably, the fact that we can obey a 'speed' directive is based upon our ability to translate the requirement into the corresponding constraints on execution; however the directive would issue at a high level of command, and probably give rise to different speeds at each component: but what is the nature of a high level "fast movement" command, and how does the system derive from it the collective speeds of the components?) If we now turn to a higher level of path control (see Figure 8), we observe that not only are there different component speeds but that in general they would *also* vary; again however, nowhere in the system would there be found an explicit acceleration directive. Yet in the same manner as above, accelerations *will* be observed in the product; only in that indirect sense can it be argued that the system has access to acceleration.

Concerning the neuronal representation of movement, the same kind of comment could be made. It is well known that the physics of muscular action is quite involved. Are we to believe that the motor brain consists of an 'implementation' of this physics? Certainly, the brain knows something about muscles and their control, but it is probably quite different than what we know, even though, at the level of muscular control it can masterly exploit what we otherwise describe as its physical potential.

Conclusion:

Grouping & Degrouping

Having reached the end of the present essay, we will now recapitulate the thesis' central elements by briefly reexamining the issue of the relation between grouping and degrouping. We shall be recalling attention to the idea of 'theories' that was touched upon in chapter 0, but in the light of the embodiments of the principles of grouping and degrouping which we have seen in the two preceding chapters. We shall also be discussing grouping and degrouping in the context of the unique nature of periphery-bound knowledge processes.

Loosely speaking, our understanding of grouping and, by extension, of degrouping, is such that their instantiation in working models seems to hinge upon one fundamental requirement, namely that of providing the systems with some built-in knowledge sufficient to accomplish a given knowledge task. We have seen, in both the case of the helicoidal ring model (chapter 1) and the second control model (sub-section 2.2.1.1), that such knowledge can take the form of a certain structure (or underlying principle) along with a certain set of related operations. Once these models are specified and shown to work, one could say about them very much the same thing as one would say about any other machine, namely that they do but a single thing (for instance to group selected sets of points into lines, or to degroup a given transformation specification into a set of local motions). Now consider the following: suppose we were to carry on from where the helicoidal ring model left off and develop the visual machine to an extent in which it could be argued to constitute a complete implementation of grouping (in the sense of having reached a human level of visual performance); similarly, suppose that the control model were expanded to the point of totally implementing degrouping. Would it *still* be true to say at that point that the systems accomplish but a single task, namely to "group," or to "degrouper?" We believe that in fact it would. But formulating the question in this fashion tends to overshadow a most important aspect of even the present incomplete machines' workings: when we maintain that they do a single thing we are characterizing the whole systems' functioning; but when we consider them locally, they actually do many *different* things: for two distinct retinal lines, the helicoidal ring will give rise to two separate events (and yet the *same* grouping process is responsible for them); similarly for the second control model, two distinct commands will provoke two dissimilar movements (and yet they are treated in the *same* way by the system). However this is the whole point behind the principles of grouping and degrouping; *it is the only way of capturing within a finite structure an apparently infinite number of products.*

On the other hand, taken in the light of the fundamental achievement of any *scientific theory*,

this is not at all surprising. But what can this have to do with the above proposition? Two more comments are in order here. First, achieving a full working theory of the principle of grouping or degrouping would certainly count as a scientific theory, for armed with such a realization one would certainly be able to explain all of its particular instantiations. That is, following the fundamental prerogative of any scientific theory, it would account for a 'universe' of possible concrete 'facts.' However, and this is the second comment, these 'facts' themselves, that is the specific biological (or technological) instantiations of the principles, have associated with them their *own* 'universe' of facts, which are, for instance in the case of humans, the whole diversity of visual experience, and the whole range of possible movements. That is why they are knowledge phenomena; that is why a model, of visual perception or movement control, was called a 'theory' in chapter 0. And just as there can be no a priori reason why there should be a limit to the range of facts explained by, say, a physical theory, there can be no a priori reason why it should be impossible to arrive at better and better 'approximations' to man, and this, again, in the face of the apparently paradoxical fact that we are finite machines endowed with a finite number of fundamental (computational) abilities.

The filiation between the principles we advanced in chapter 2 on the production of movement and the principles we outlined in chapter 1 on visual perception is not one which we pursued beyond an acknowledgement of the 'importation' into the area of movement control of certain key concepts, especially the definition of motion computation. Now the problem of movement production does not necessarily easily lend itself to a simple transposition and inversion of the perceptual principles. We did however (in section 2.0) introduce the concept of degrouping in the context of such a reversal of visual perception. But when the context shifted toward the control of an effector, this was no longer possible mainly because of what was called (in section 2.2) the non-homogeneity of the motor interface, whose basic property, namely that it moves, led us even farther away from the 'ideal' perceptual interface with which the discussion had been initiated.

On the face of it, there are some discrepancies between what could have been expected to be the nature of the instantiation of degrouping following the arguments in section 2.0 and that which was actually developed in the context of movement production (especially the second model). Again this is due to the nature of the respective ultimate objects of control, i.e. the two different peripheries. The basic property of the visual interface, taken as an object of control, is of course not a movement potential but an *imaging* potential. Consequently, if we now transpose the basic tenet of our conception of movement control to that earlier case, we would have to say that a model of imaging (i.e. image production) has to embody a representation of the imaging potential of the whole visual interface. In a way however, this is exactly what a visual system does: it captures that potential in the sense that it's ability to characterize, at any one moment, the totality of the interface's state must be a function of the system's capacity to characterize *any* interfacial state (or sequence thereof,

since time is an essential dimension of interfacial events). This is why the discrepancy is only a function of the nature of the interface.

The process of degrouping has two facets, one being static or fixed, the other being variable or changing. The fixed aspect of the process of degrouping is found in the connections between characteristics of the representation of the effector and the effector components; this was called a path of computation in section 2.2. The representation is always 'interpreted' in the same way. But the nature of the representation is such that there is room for variation, which in the control model takes the form of the specification of modifications of the variable or secondary characteristics of the representation. In visual perception the situation is not different: the fixed aspect of the grouping process is found in the detection of the primary characteristic, which (in general) leaves room for the attribution of secondary characteristics. For example, and following the above, for the helicoidal ring model of chapter 1 we would say that the state of the interface is always 'interpreted' in the same way by the system; the system always seeks for the existence of a linear relationship among sets of retinal events. However this interpretation occurs within a variety of 'contexts' defined by the combinatorics of the secondary characteristics of line segments. Supposing one were to use the helicoidal ring for the purpose of producing line segments, all the variable (i.e. secondary) features would have to be explicitly stated, while the fixed aspect, namely the constant linear relation, would not need to be. In the same manner, the invariant '*interpretation*' of the effector representation never needs to be explicitly demanded in the command, while the *modifications* must always be explicitly dictated.

Thus, even though because of the fundamental difference in the nature of the motor and visual peripheries a degrouping process for a movement production system cannot be defined through the simple inversion of a grouping process for a visual system, the same fundamental logic has been shown to be relevant to both situations. It is the logic of periphery-bound knowledge processes.

o o o o

APPENDIXES

CONTENTS

Appendix I

Helicoidal Ring Simulation	158
0 Introduction	158
0.0 The helicoidal ring and its simulation	158
0.1 Specifications	159
1 Steps of the simulation	161
1.0 Projection	161
1.1 Tremor	164
1.2 Inhibition	168
1.3 Some stimuli and outputs	171
2 Wrapping up	185
3 Retinal size, field size, and orientational resolution	185
Addenda 1-7	191

Appendix II

Note on the 'hexconv' Function	197
--------------------------------------	-----

Appendix III

The Simulated Arm	202
I: Logic of Implementation	202
II: The Programme	207

Appendix I

Helicoidal Ring Simulation

0 Introduction

0.0 The helicoidal ring and its simulation

The helicoidal ring (or "pile") is a basic computational unit, or primitive, originally developed by Lamontagne (1976) to serve in the construction of a particular visual motion detection system. It has features that allow it to be used at different computational levels and for a variety of computational purposes. Some of these were explored and illustrated in Lamontagne (1976) while in a further study (Lamontagne & Beausoleil 1982) the helicoidal ring was explicitly contextualized within the larger issue of visual spatiality (understood as part of the visual world) by tying it, in a conjectural manner, to the visual modality, and drawing from this rapprochement specific consequences for both aspects of vision.

Of course all this rests on the fact that the helicoidal ring does work, that it effectively solves the particular problems that it addressed initially. There is however no easy way to prove this yet with the kind of generality that would make of it a true formal primitive. The nature of the computations involved in the ring is very simple, but the structure of the ring, partaking intimately of its computational characteristics, is not as simple. This makes for two fundamental difficulties, both partially allayed by simulation. First, even though the functional components of the pile are easily apprehended, their simultaneous workings and interactions are not easily followed even in a particular case, let alone for the general one. Second, there exists, to our knowledge, no 'natural' mathematical treatment which can be applied to such a unit (structure & functioning) in order to convince ourselves quickly of its competence; the ring was developed with the explicit desire to avoid as much as possible recourse to sophisticated mathematical operations (by definition, anything beyond addition and subtraction), and the attempt was very successful in this respect. In a way, therefore, it has traded computational simplicity with structural complexity. Mathematical treatment of the pile's structure and function is not impossible of course but the point is that there is more appeal for a 'direct' demonstration of its functioning than in a (possibly tedious) formal treatment that calls for the mathematization of the structural elements of the ring, a procedure which, for the sake of power, loses much intuitive appeal, or lacks 'naturalness.'

We assume that the reader is familiar with the helicoidal ring's

workings, as set out intuitively in Lamontagne & Beausoleil (1982) and chapter 1 above, and more specifically with the unit's basic four-stage computational sequence, namely, *projection, summation, inhibition, and reduction*. The present simulation has been concerned with the first three stages only, the fourth (reduction) being critical only with respect to an eventual integration of the ring's output within a larger computational scheme, and not presenting any foreseeable difficulty. The second step (summation) being nearly trivial, most of the work consisted in the implementation of projection and inhibition.

0.1 Specifications

The present simulation was carried out on a Cromemco™ Z-2 microcomputer, running on a Z80 microprocessor, with 128K of RAM, a Dual Disk Floppy Drive, and an RCA graphics monitor (model TC 1212) driven by a Matrox card (type ALT-256AS). The programme was written in machine language via the Cromemco CDOS Macro-Assembler, version 2.12.

There are three parameters attached to the pile: retinal size, field size, and orientational resolution. A consideration of the capacity of the machine at hand led us to adopt the following values for those parameters: the (square) retina holds 96x96 receptor cells, its fields were set at 3x3 receptor cells, and the rotation of any field cell layer (the actual layers of the ring) relative to its neighbour in the downward direction is 3°, the axis of rotation being perpendicular to the layer's center. Retinal size, field cell layer size, and orientational resolution are closely related since for instance one does not want orientational resolution to surpass retinal representational capacity, or conversely for retinal representation to surpass orientational resolution by too much: for the first case, consider a pile with a 1° resolution equipped with a retina that cannot go below 5° in discrimination; and for the second, consider a retina with a 1° capacity attached to a ring with only 5° of discrimination - the first leading to unwarranted redundancy, the second to unwarranted paucity of analysis. Retinal representational capacity however *must* exceed layer resolution in order to implement rotational tremor. (Rough calculation using human performance data in orientational discrimination and Hubel & Wiesel's (1979) figures concerning the possible number of "cortical units" would suggest that an ideal pile would have an orientational resolution of 1° and a 'retina' of about 200x200 receptors; field size remains unspecified in this context - but see section 3 for a further discussion of this point.) The hypothetical nature of these calculations however remains consistent with the nature of the pile: for not only is 1° perfectly representable on such a matrix of receptors, but 0,5° is also, thus allowing for the operation of rotational tremor; furthermore these figures are close to the limit, the 'next' step down (i.e. 0,5° resolution,

with a rotational tremor of $0,25^\circ$) pushing the effect of rotational tremor *outside* the retina since rotations of the retina for 1° are 'visible' at the 29th receptor from the *middle* one onward, for $0,5^\circ$ at the 57th, and for $0,25^\circ$, at the 115th.

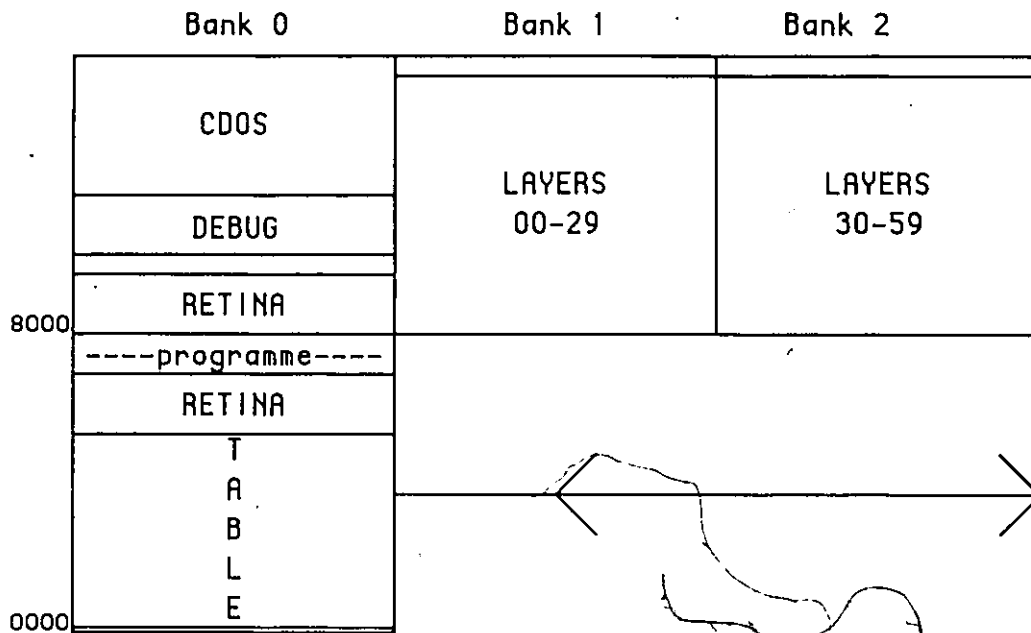


Figure 1. Memory organization in computer for simulation. CDOS is used mainly for loading the various lookup tables (all dumped into the TABLE are), while DEBUG is necessary for loading the programme, placing it at 6E00H and running it. The higher RETINA is the original one, containing the stimulus; the lower one is either a copy of the original or a result of the tremor-transformation. This is the one used during the projection phase, 'travelling' with the programme into Banks 1 & 2. LAYERS of the ring are organized sequentially; the orientation of any particular layer is equal to three times its number.

There are 60 layers in the simulated ring. This means that characterization of orientations beyond 180° are not effected directly by the pile, but rather are computed in a manner that takes position of the layers into account. Thus the pile will not 'see' a 225° angle, but two segments that are at a 'distance' of either 15 or 45 layers from each other depending on which direction is searched among the layers. There is no real advantage to be gained in terms of orientational characterization by

'twisting' the ring even more (180° more to be exact) since then *one* line segment would generate *two* complete outputs from the analysis (in the case at hand, one pair at the first and 60th and at the other pair at the 15th and 75th layers, counting layer 0 as the first), and thus involve further considerations at least as complicated as those implied by a restriction to 180° . Furthermore any extra information drawn from such a scheme is already implicit in the usual one. Note finally that in Lamontagne & Beausoleil (1982), the representation shown in Figure 5 does suggest that the 360° scheme is standard, but this should not be understood.

Memory in the microcomputer is organized in blocks of 32K. Since *one* byte represented *one* (either retinal or field) *cell*, 9K were necessary for the retina (96×96) and 1K per field cell layer (32×32); the whole pile required 69K of RAM to start with (60 layers + retina). At any one time memory access is limited to a maximum of 64K only, so that we could not have everything at our disposal simultaneously. We solved this problem by splitting the 60 layers in two groups of 30 and assigning each 30-layer block to the upper part of Bank 1 and Bank 2, while the block in Bank 0 (which held the retina and the programme) could softwarewise 'travel' freely to Banks 1 & 2 (see Figure 1).

1 Steps of the simulation

1.0 Projection

One of the original aspects of the present simulation with respect to previous ones was the use of lookup tables for various purposes, chief among them being the projection (phase I) of the retinal stimulus to every oriented layer of the ring. The main advantage of this approach is that time-consuming operations are replaced by the very simple calculation of an address offset; of course the cost is space since the tables are held in memory to minimize accesses to the disk. But previous experiences at simulation with high-level languages (POP-II, SIMULA, LISP) made this a very attractive proposition, in view of the wish that results should be available within reasonable delays. This is in fact what happened: what previously was a matter of hours now took but over twenty minutes of processing.

These tables contain 9216 pairs of coordinates (stored in separate 18K files) arranged in an orderly fashion corresponding to a left-right/top-bottom cell-by-cell sweep of the retina. The coordinates themselves are those of the particular field cell of a particular layer of the ring to which the given retinal cell projects (see Addendum 2 for an extract drawn from the table for 3°). Thus there should be 60 tables in all, one for every layer. However a very simple transformation of the coordinates for the layer at μ is used to generate the coordinates at $\mu + 90^\circ$ [(X, Y) becomes $(-Y, X)$], so that the number of tables is effectively reduced to 30; it could

actually have been cut down to 15 but the required transformation being a bit more intricate, and only extra disk space being taken up by the set of tables (just one being 'active' at one moment) we decided not to do it. On the other hand, we found out (when tremor was implemented) that reading each table three times (once for the nine projections at each of the three rotational tremor positions) did become rather time-consuming.

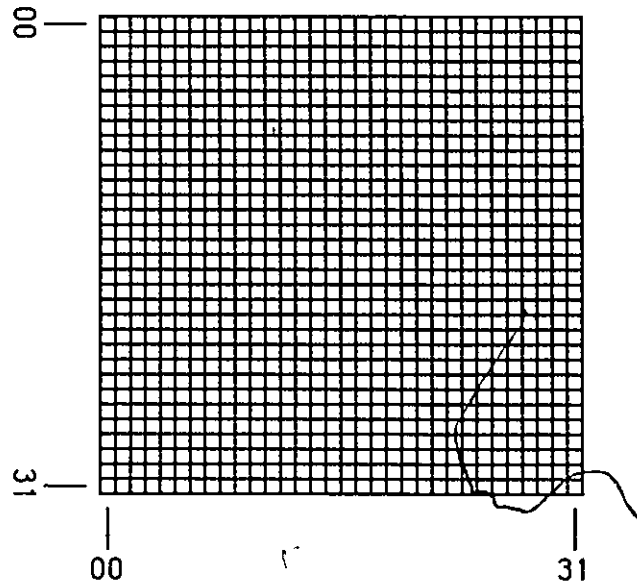


Figure 2. *Layer coordinate scheme used in projection tables.*

A coordinate is a value between 00 and 1F (in hexadecimal representation, or, in decimal representation, between 00 and 31) corresponding to field cell displacement along the X- and Y-axes according to the scheme shown in Figure 2 where the origin is at the top left. To obtain the *address* of the field cell to which projects a particular retinal cell, we needed simply to calculate $20H \cdot Y + X$ and add the result to the address corresponding to the 'origin' of the particular layer with which we happened to be dealing. The value of the 'origin' is updated sequentially by 400H (1K bytes) starting at 8000H, to move from one layer to the next in the projection. We could of course have pre-calculated the *addresses* (which also occupy 2 bytes each) and put *them* in the table instead of the relative coordinate scheme that was finally adopted; this would have had two side effects: first the location of the layers would have been fixed to one place in memory, and second, we would have had to re-write the transformation allowing mapping $\mu + 90^\circ$ from the *address* for

the projection at μ instead of from the coordinates, as it now stands. However the extra computational cost, if any, of the present setup is not really significant enough to render it essential. As for the retina we dealt directly with the address so that, contrary to the above proposition, and in conformity with the existing scheme at the level of layers, it can be anywhere in memory, only the initial value (or 'origin') having to be modified in case of a displacement. This address is stacked during projection.

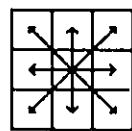
There remain just two points with respect to the question of these tables. First, why should *one* be needed for every layer? Offhand one would think that, in principle at least, just one table is necessary for the whole ring, that would be recursively applicable to the *result* of the projection on layer μ in order to generate the *next* one down (i.e. $\mu+3^\circ$). What prevents this is the field structure imposed on the retina, i.e. the fact that *retinal* cells are mapped onto *field* cells, so that it is not true in general that two successive field cells (relative to a 3° rotation) are mapped onto by the *same set* of retinal cells; this in turn can be seen when one considers the unitary or discrete nature of the process of projection; it is *cells* and not points that are being mapped. In fact if one considers *only* the retinal cells corresponding to the center of the field (i.e. the middle cell) at layer 0 it is true that a single transformation table would work and that consequently only one would be necessary; however for the *other* retinal cells the mapping at $\mu+3^\circ$ can fall in one of the 8 adjacent field cells to the one predicted by a transformation of the field cell at μ , and this is so because even though the *retinal* cell is *inside* the area which corresponds to the field cell at 0° , it does not remain so for *every other* field cell on every other layer to which the 'central' retinal cell maps. At one point it 'jumps' outside of it, onto a field cell next to the one predicted by the supposed single (rotation) transformation of field cells. There no doubt exists some systematic way in which this occurs, and it is probably a function of layer number with respect to retinal radial distance and field size, but we have not worked it out. Note also that if the field size is an even number (e.g. 4×4) there is *no* central cell, so that the argument is even more critical in those cases. - The second point concerns the generation of the tables. This was done in BASIC (see Addendum 1 for the listing). The retinal 'cell' was defined as the middle coordinate of 1×1 squares in the cartesian plane, while the coordinate of the field cell was obtained by rounding off the result of the transformation. This can be seen clearly in the listing (at lines 70,80).

Phase II (summation) it is a very simple matter. We chose to summate on the rows instead of the columns of the layers as a matter of computational convenience.

25

1.1 Tremor

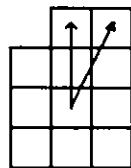
A slight multidirectional movement of the eyeball produces small shifts in the retinal image. Therefore implementation of tremor involves shifting the stimulus around in the largest possible number of different directions with relatively strict limits on amplitude. Now different *directions* are represented by related changes of *position* and since there is so little room to manoeuvre (and this not only remains physiologically plausible but also computationally effective since too large a tremor endangers the identity of segments lying close to each other, the idea being that a search for the best-fitting projection that ventures too far away from the locality of the given segment can be confounding), very few different directions are available: with an amplitude of one (1) receptor, only eight (8) are possible (13% of the total number (60); see Figure 3a). Now the tremor component was originally added to improve projection, once the field structure was imposed, which itself was an aid to projection. Since the field size we chose was so small (3), we originally assumed that these projection problems would also be minimized, and that a tremor reduced to a 2x2 'field' would suffice (see Figure 3b).



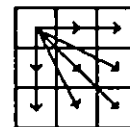
(a)



(b)



(c)



(d)

Figure 3. (a) 'standard' translatory tremor: amplitude & direction
 (b) 'reduced' tremor
 (c) 'short' and 'long' tremor; extra directions
 (d) tremor scheme in a previous simulation

It turns out that this is not the case at all for the very reason stated above, that is, tremor has a direction *and* a distance component that are intimately related. By dropping the size, we effectively also dropped

more than half the available directions, and since it is necessary by definition to wholly conserve the direction component we soon realized that we had to revert back to what might now be called a 'standard' scheme of tremor (Figure 3a). Of course the reverse is also true: by allowing for greater amplitude one can 'visit' new directions (Figure 3c, right arrow) and also travel farther in an 'old' direction (Figure 3c, vertical arrow). It is with respect to this 'long' tremor that we noted above a possible danger of confusion, for two distinct parallel retinal projections could end up as two identical adjacent segments on a layer, the larger one eventually 'replacing' the shorter one after a 'within tremor' search for maximal projections. By itself this is not necessarily catastrophic but it defies the original assumption regarding the pile's intended output, according to which such a situation should not arise. Therefore even though extra directions are thus opened up, an attractive possibility in itself, we decided to abandon the long tremor and stick with the short one.

It should be added however that following experimentation these extra directions in fact yielded little or nothing new, and so did not really improve the mapping. Why this is so is not really very clear, but field size certainly has something to do with it, as by definition stimulus width is equal to field width, so that for a larger one, resolution being enhanced accordingly, what we might call the 'imperfection' of the stimulus is much less of a problem, rendering the original projection (i.e. projection of the initial stimulus) rather good. For example in a previous simulation, before this point was cleared up, and assuming one considers the 'long' tremor and extra directions to have been redundant, it turns out that a 'reduced' tremor sufficed (see Figure 3d); now in that simulation field size was 4x4 (but orientational resolution at 5°, possibly explaining the redundancy of the extra directions), but the stimuli were line segments, whereas we worked on this question mainly with a stimulus circle, so that the argument remains more suggestive than conclusive. Yet it is clear enough to predict that in a future simulation where field size would be larger than any so far tried, say 5x5 or 6x6, the effective improvement in mapping brought about by our 1-receptor amplitude translatory tremor would be nearly negligible (in our case the amelioration was very evident). However one should not conclude that at one point the necessity for translatory tremor will vanish completely, since for a threshold set at more than half the field 'area' (e.g. 12 out of 16 in the case of a 4x4 field size), it will always be possible for a stimulus segment to fall half on one field cell and half on one of its immediate neighbours, rendering necessary a 'push' in either direction in order to jump the threshold in any one of them. Reducing the threshold at or below 50% can of course entail considerable 'double detections.' All this remains consistent with the discussion of translatory tremor in Lamontagne & Beausoleil (1982).

Contrary to the translatory components of tremor, which could easily be implemented by various linear translations of a copy of the stimulus in memory (placed in the upper RETINA of Figure 1), rotational tremor was implemented by a transformation table, previously computed and stored on disk along with the other tables (see Addendum 3 for the programme, and Addendum 4 for an extract from the table). The principle involved here is identical to that for the projection tables. We computed the rotation of the retina by $+1,5^\circ$ and $-1,5^\circ$, effected it on the stimulated retina and projected the result along with all translatory components in turn. Theoretically, if a stimulus segment's orientation is restricted to integer multiples of 1° , for a pixel resolution such as we have (3°) a tremor of 1° amplitude is sufficient to 'push' a wayward segment in the closest template (e.g. if a segment is at 4° then $4^\circ - 1^\circ = 3^\circ$, which fits exactly, while a push of $+1^\circ$ fails to do so). Consequently a question arises concerning the extent of rotational tremor. Now it so happens that retinal resolution does go down to 1° (in fact to $0,5^\circ$, but not to $0,25^\circ$). So on the one hand, given the above argument, a 1° rotation is attractive, while on the other a $1,5^\circ$ rotation allows for a more marked effect: at $1,5^\circ$ the effect of retinal rotation becomes 'visible' at a distance of 19 cells from the center of the retina (which is also the center of rotation), while at 1° 'visibility' is pushed out to 29 cells. Depending on the specific retinal input, one extent might be more useful than the other. This is precisely what happened, i.e. the 'long' rotation seems more appropriate for curves while the 'short' one appears better suited to straight line segments (see further below for an example).

Therefore the projection phase, the longest one in terms of processing, involved three complete successive mappings on the layers, one at rotation '0', one at rotation '1/1,5' and finally one at rotation '-1/-1,5'. For every projection, at every stimulus angle and for all (9) positions, a call was made to a subroutine which chose the best-fitting segments between the current mapping and the result of the previous choices. [This aspect of the analysis has no structural counterpart in the original model of the helicoidal ring; as it now stands it could not directly be transposed into a wiring diagramme. Furthermore it is still incomplete, that is, it does not generate *all* maximal mappings.] This was effected in each memory Bank using the extra 2K at the bottom of the block; 3K were involved: the layer matrix (1K) containing the current projection, the F800H matrix (1K) containing the previous cumulative best mapping, and the RESULT matrix (1K, placed immediately after, at FC00H) in which were inscribed the results of the current analysis, this matrix being moved to F800H after completion for the next projection. (Again, since addresses were identical for the μ and $\mu+90^\circ$ layers, the BESFIT call analysed the *two* target layers of the projection in one pass.) Once a complete cycle was done (for the nine translatory positions at a given rotational position of the stimulus) the F800H matrix was moved back to its

corresponding layer, the analysis was then pursued on the rest of the layers, after which the 'next' rotation was imposed on the stimulus, whereupon the whole process was repeated. This was strictly a programming device that saved having to read the rotation tables for every pair of layers individually. Had enough 'live' room been available to either hold three copies of the retina, or, as we had, two copies, plus the rotation table, then projection cum tremor could have proceeded directly and the analysis been done completely for one (in fact two) layer(s) at a time.

Contrary to translatory tremor, rotational tremor does not 'continuously' affect all retinal cells. Geometrically speaking it does of course, but the discretization process necessary in mapping, and the rounding off procedure which ensures it, gives rise to a circular area in the middle of the retina which remains unaffected by rotation. The radius of this area is given by $[(1.5 \cdot \cos(\mu) - 1) / \sin(\mu)]$, where $\mu = 1^\circ / 1.5^\circ$, the 'abscissa' of the first transformed cell (at $y=0.5$) that 'jumps' into the next cell down (or up). It is easy to see that the finer the orientational resolution of the retina, the larger becomes this radius; intuitively it follows from the fact that retinal resolution is a function of its size. In general therefore, just as in the case of translatory tremor, the contribution of rotational tremor components to mapping should tend to diminish as the pile's resolution (orientational and retinal) increases. However, contrary to the translatory components, their effect is non-homogeneous. [The homogeneous effect of translatory tremor is due to the uniform distribution of receptors within the retinal array; thus it is not absolutely necessary, but peculiar to the present 'realization' of the pile.] Thus any segment small enough to fit inside the 'untouched' area will not benefit from rotational tremor; and in general the finer the resolution of the pile, the longer can be this segment. (This is interesting, as supposing the *amplitude* of rotational tremor to be *less* than the orientational resolution of the pile places a lower limit on this segment's length for a given retinal size.)

What happens then inside this area, where resolution is *less* than the total retinal resolution? A segment in this area is forcibly orientationally ambiguous and one cannot count on rotational tremor to push the segment in the closest 'slot' since it has no effect on the stimulus; so it will be mapped maximally on at least *two* successive layers; if these projection segments are not sufficiently different in size to destroy the other, then the pile will 'see' two *different* segments. It is difficult to be more precise here since the actual *length* of the (small) segment will affect the mapping, the extreme case being the single point, which would map maximally on *all* layers. [This is a particularly clear example of the way in which 'pile geometry' deviates from pure analytical geometry, where a line segment can be assigned a unique orientation independently of its non-zero length.] For a given segment falling inside the unaffected area it is therefore hard to predict exactly where it

will map maximally, except the immediate area (see further down for more details on an actual case). However when retinal size is increased, as it is in fact in our case (with respect to the special sub-retinal field in the center), and for a sufficiently 'large' stimulus, orientational tremor *does* start to play a rôle. For instance we dealt with a circle (radius 43 retinal units), which was placed in the 'affected' area, and noticed that addition of rotational components of $1,5^\circ$ ensured a maximal mapping in 20 new layers (10 for each component). Yet retinal size cannot be increased indefinitely if one is to remain close to physiological evidence; computationally of course there is little restriction against it. However if one thinks of a retina involving tens of millions of receptors, a $1,5^\circ$ tremor builds up quite a bit at the periphery (for 100 000 000 receptors, to a sweep of 44 retinal cells); this can be problematical with respect to the search for maximal mappings since for a given field size successive layer mappings will start occurring farther and farther away from the original. One way of achieving 'large' field size while conserving orientational resolution could be through a modification of the mapping structure of the retina or its receptor layout (from a homogeneous to a non-homogeneous distribution) in such a way that this peripheral swing is compensated for or reduced to 'normal' size; however for the moment it remains a suggestion that has not been explored fully.

1.2 Inhibition

The inhibition phase entails, for every field cell of every layer, choosing a vertical (i.e. perpendicular to the layer) linear path through the immediate neighbouring layers. The cell values encountered there are compared to an inhibition weight which is a function of the 'distance' of the inhibited layer from the inhibiting cell's and also of the 'force' of the inhibiting cell. This function is simply the (discretized) cotangent of the particular angle corresponding to the 'distance' of the inhibited layer (see Figure 4). In the programme (see Addendum 7 for the first page) it is defined as a vector (COTAN, line 14) made up of the integer values of the cotangent from 3° to 18° at 3° intervals. (Note that these values are not rounded off to the closest *lowest* integer value, but to the closest *largest* integer value; thus, $\cot(3^\circ)=19,081$ is rounded off to 20, not 19; see section 3.) Formally speaking it should be carried through to 45° , where the cotangent becomes 1; but an inhibition weight of 1 is completely useless, since there is no possibility of a segment of length less than 1 occurring in a layer; furthermore since segments of length 1 and 2 have been removed after projection and summation (such short projection segments being very difficult otherwise to purge out of the system completely and also 'uninteresting' with respect to the ring's purpose) it follows that the shortest segments 'surviving' projection being of length 3 and more, we need only inhibit 'as far as' a weight of 4. Were it the

case that an inhibition weight could 'kill' a cell value that is equal (as opposed to strictly smaller) to it, then of course we would have extended the vector one more layer. However it was decided against that since otherwise, for example, a stimulus segment covering a length of 3 fields becomes self-destructive, an undesirable result; more generally, for 'imperfect' detections, ones that involve multiple, and possibly equal, detections, it seems more preferable that they should all survive than that they should mutually destroy each other; it is also critical for the analysis of curves where short segments of equal length and in close layer proximity are called to survive.

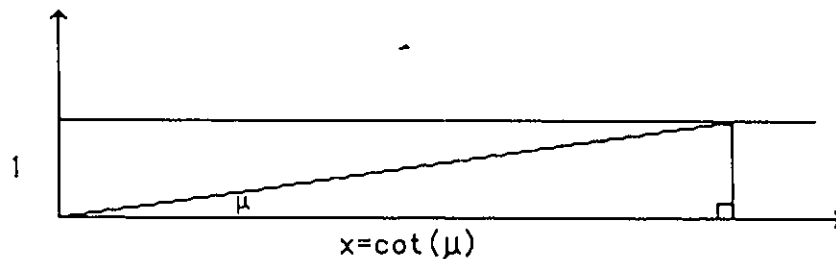


Figure 4. Geometry of a projection segment at layer μ ; because summation is done row-wise, we are interested in the length of the abscissa, which becomes simply $\cot(\mu)$ since the ordinate is 1 by definition.

In line 15 of the programme we encounter another vector, called INHVEC. This vector is tailored to individual field cells and holds the progression of inhibition weights for the six critical layers; it is used for both 'upward' and 'downward' inhibition since the values are identical in both directions. It is constructed in the following manner: the current weight of the field cell is compared successively with the values in COTAN and assigned to INHVEC until a lesser value is encountered, whereupon it assumes it and the rest of the values of COTAN. These values are then used in the inhibition process.

Again for the inhibition step we used a table previously generated in BASIC, in which were included sequentially, for every field cell of one layer, the required twelve pairs (6 in the 'upward' and 6 in the 'downward' directions) of coordinates (see Addendum 5 for the listing, and Addendum 6 for an extract of the table). Afterward it was just a matter of reading a field cell, generating the corresponding INHVEC and then using the appropriate coordinates to compute the address of the currently inhibited cells. Concerning the nature and use of this table the following important point has to be made. Since the pile is only twisted 180° , special cases arise in the inhibition process whenever the extreme layers

(layers 0 and 59) are encountered, that is whenever, in 'travelling' in the 'downward' direction, inhibition has to be carried from, say, layer 58 to layer 59 and *then* to layer 0, layer 1, etc., or conversely, whenever in 'travelling' in the 'upward' direction, inhibition has to be carried from, say, layer 1 to layer 0 and *then* to layer 59, layer 58, etc. Now layer 0 is 183° 'down' from layer 59 ($177^\circ + 183^\circ = 0^\circ$), and so the coordinates in the part of the table for the 'downward' direction (the first six pairs) cannot be used directly. However, using the fact that $\sin(180+\mu) = -\sin(\mu)$ and that $\cos(180+\mu) = -\cos(\mu)$, and also that the table does contain the transformed coordinates corresponding to the latter (positive) values, we can readily compute the correct address for the particular inhibited cell in layer 0 simply by 'negating' the relative coordinates found in the table (by subtracting them from 31). Thus the first half of the table is used in two ways in the inhibition process in the positive or downward direction. Now the second half of the table, containing the six coordinate pairs computed for a transformation of 177° , 174° , etc. (see Addendum 5), that is, the 'upward' direction part of the table, is similarly used in two different ways: for the inhibited layers *below* layer 0 the 'negation' process is carried out using the fact that $\sin(-\mu) = -\sin(180-\mu)$ and that $\cos(-\mu) = -\cos(180-\mu)$. But for an inhibition path that must travel beyond layer 0 into layer 59 the table coordinates are used without modification, since for example layer 58 is 174° 'up' from layer 0 (and not -6°).

In fact inhibition concerned not only this computed 'central' cell but also (via the HFIELD vector, line 16) its eight immediate neighbours, which we treated (i.e. inhibited) in the same manner. This was done as a general palliative procedure for the cases where, for example, one field cell 'sits' on top of (or below) three or four corresponding field cells of another layer; although it can be made mathematically clear exactly to which unique field cell inhibition must be carried it cannot be argued to be sufficient with respect to the model: to make this point clearer, we might say that inhibition of a cell plus its surround is the counterpart, among the layers, of the mapping problems which occasioned the creation of retinal fields.

Care must be taken in the implementation of the inhibition phase; since we want *every* (non-null) cell to exercise an inhibitory effect, it should not be 'killed' before it has had a chance to do so, a situation which can easily arise in a straight sequential procedure. To avoid this problem we simply *marked* the inhibited cells, and deleted them *only* once the whole process of inhibition was completed.

There was no need to generate an inhibition table for every layer of the pile since they are all organized identically with respect to each other. We could, in fact, have reduced the table even more, at the cost of a repetitive 3° -transformation on every 'visited' layer, but again it seemed simpler just to list them directly. The table of course needed to be used just once.

1.3 Some stimuli and outputs

By definition a legal stimulus for the helicoidal ring is one whose width is equal to field size width. In our case they spanned 3 receptors. It turns out, because of the digitalization process, that this value inescapably assumes the nature of an *average*, i.e. for a given stimulus the effective width ranges from 2 to 4 retinal cells. Now this is the source of 'imperfections' in the stimulus, a problem that can probably largely be offset by increasing field size. In the case of the circle (see Figure 6) they were pointed out by the layers on which a maximal mapping could not be obtained after the 27 tremor-positions projection cycle had been carried through. The reason is that the curvature of the circle is not continuously smooth, and breaks down in certain areas; these form a symmetrical pattern (see Figure 5) which however is not readily predictable from a simple examination of the stimulus (but see section 3 for a possible explanation).

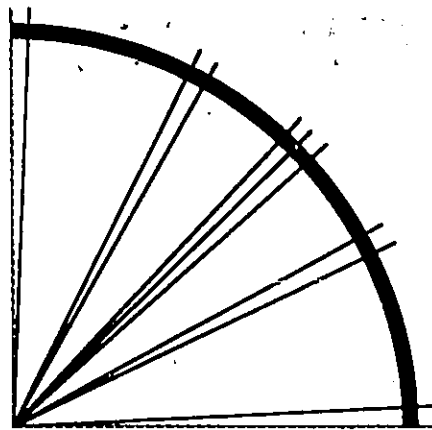


Figure 5. Areas of 'imperfection' in the circle; the pattern is replicated identically in every other quadrant. On these layers the segment length (tangents of the circle) attains only 9, whereas on every other it reaches 10 (or OA hex; see Figure 7); these layers correspond to the following angles (with respect to the X-axis): 3, 27, 30, 42, 45, 48, 60, 63, and 87.

This is expressed, after completion of the projection phase, by a slight difference of length in individual pairs of tangents that never exceeds 1 (i.e. non-maximal mappings are two tangents of length 9 instead of the maximum 10). These layer contents of course are 'killed' by adjacent segments of length 10, that in fact inhibit at this strength for

the *two* neighbouring layers in any direction (for a total of four). So not surprisingly no immediately neighbouring 'lesser' segment survives this inhibition. It is interesting to note, in view of this result, that the pile would 'see' the circle as something resembling a disjointed 'polyhedron' of curved segments of unequal length; it is even more interesting to speculate that for a much smaller circle (and one that would not benefit from rotational tremor) the pile would start seeing something between a circle and a square, i.e. an ambiguous figure. Be that as it may, the output from analysis of the circle was pretty close to the expected one, a typical 'ideal' example being shown in Figure 7. Out of the expected 120 equal tangents, the pile generated 84, or 70%. Recalling that the retina involves only 9216 cells (of which only 7235 (78,5%) are 'effective'), and that field size is 3×3 , this result seems remarkable. It is in fact our best one, obtained with a $1,5^\circ$ rotational tremor. At 1° of rotation the drop in performance is considerable, only 64 segments (of length ranging from 9 to 11) emerging from completed analysis. This quantity is the same as that yielded by an analysis *without* rotational tremor, so that at 1° of magnitude it contributed, except for a few previously unencountered larger segments of length 11, no overall improvement. We linked this degradation to the imperfection of the stimulus, and assumed that a rotation of $1,5^\circ$, by providing a more extended effect, reduced it considerably.

The two line segments shown in Figure 8 differ in every respect: position, orientation, and length. They are, moreover, intersecting. Lastly, not only does the orientation of segment B correspond exactly to a layer (number 11), but its projection benefits from rotational tremor (since it falls partly *outside* of the retinal area untouched by it). No real surprise therefore to find it mapped maximally in only one layer (Figure 9); and no further unexpectedness in watching it 'clean out its 'environment' completely during inhibition. Segment A is more problematical, and more interesting. It is nearly completely inside the 'untouched' area and its orientation is off from that of any layer. Projection (Figure 11) with a $1,5^\circ$ tremor tells us that it is mapping maximally on three layers, and that it is 'seen' as close to layer 34 as it is to layer 35, which makes sense.

The length of these segments is also affected by the peculiar orientation of the A segment, the *expected* lengths being calculated from projections at *precise* 3° intervals. This is the first critical observation. Next, and most importantly, we invoke a somewhat different version of the stimulus imperfection argument. This was not evident from the start. Close examination of the A segment (at an angle, looking from the bottom) reveals a certain thickening of its upper part; it turns out that this is enough to explain the discrepancy, for the extra 'moving space' allows for a (nearly) maximal mapping in exactly the direction observed (layer 36) but not in the other. We believe this 'distortion' to be part of a

larger pattern inherent in the discretization of the segment and only partially visible here. Our first idea was that this artefact would be sensibly reduced (if not completely surmounted) upon *lengthening* of the stimulus segment - and indeed this idea is *not* really different from the previous one concerning increasing directional ambiguity with decreasing length, or the conclusion we finally arrived at for that matter, the main problem having been that we thought strengthening one characteristic would be sufficient to conceal the defective other. So we tried this, and obtained a result very similar to the first one! This was surprising if only for the fact that the segment now overflowed much more largely into the tremor region, thus discrediting somewhat the importance of our earlier argument. In fact we did more than that: we also systematically thickened the B segment without altering either its length or orientation, and produced an ambiguous or 'curvilinear' output (three maximal segments surviving inhibition). It seemed then, by elimination, that stimulus imperfection must have been the cause.

A mistake in our initial calculations had led us to believe that a 1.5° rotation was the best we could do, and although it turned out that our argument concerning the appeal of a 1° tremor was vindicated, as we shall see presently, we would probably otherwise have missed the 'extent argument' and not realized immediately that the analysis of the circle could thus be largely improved. Unfortunately, as noted above, the argument does not carry over to the circle and we are left at this point with the somewhat disagreeable segregation of straight lines and curves: it is in this context that stimulus imperfection appears to weigh more and more heavily, for it seems the most promising path to ultimate reunification.

Figure 10 shows the final projection segments on layers 34-36 with a 1° rotational tremor. As can be seen, layer 35 is selected as the one corresponding best to the stimulus segment's orientation: no doubt this is because, through a 1° shift in the stimulus we 'attained' 105° (the orientation of layer 35), while no other layer was likewise perfectly 'met' ($104-1=103$). Of course the inhibition weights generated in layer 35 are sufficient to kill both neighbouring projection segments, so that the desired result is achieved: only *one* segment survives analysis.

This result is coherent with the above justification, for it appears that an imperfection of the second kind (i.e. in terms of stimulus width) was indeed concealed here by a more restricted rotation of the retina, allowing us to suppose, a posteriori, that imperfect width alone was not responsible for the ambiguity of detection, the extra push contributed by a too large tremor also being required.

There is one more surprise in stock. First note that, as concerns segment B, the expected outcome was generated irrespective of the size of rotational tremor: this is not surprising in view of the fact that segment B has a 'perfect' orientation. So it is quite tempting to genera-

lize this to all such straight line segments. Yet it would be incorrect to do so, and the one counterexample we know of is the most unexpected, because among the most 'perfect' stimuli of all: the horizontal line! This stimulus is *triple* detected by the pile: a maximal segment (length 32) survives on layers 0, 1, and 59 *irrespective* of the size of rotational tremor. So again the segregation we talked about is not as complete or as simple as we thought. This result in fact brings us back down to the fundamental issue of the nature of the retinal array, for the peculiarity of *this* perfect stimulus is that it lies along one of the axes of symmetry of the *retina*. Strangely enough, it is as if the pile had reversed the staircase problem! Could it be, simply that the 'real' retina has *no* axes of symmetry?

This concludes the present section. We have not so far carried out any more systematic investigation of the pile's behaviour in terms of a larger portion of the class of straight line segments and circles, or attempted the study of the pile's competence in new ones. There is much interest in doing so, especially in the area of 'small' stimuli, as it could possibly lead to further critical predictions.

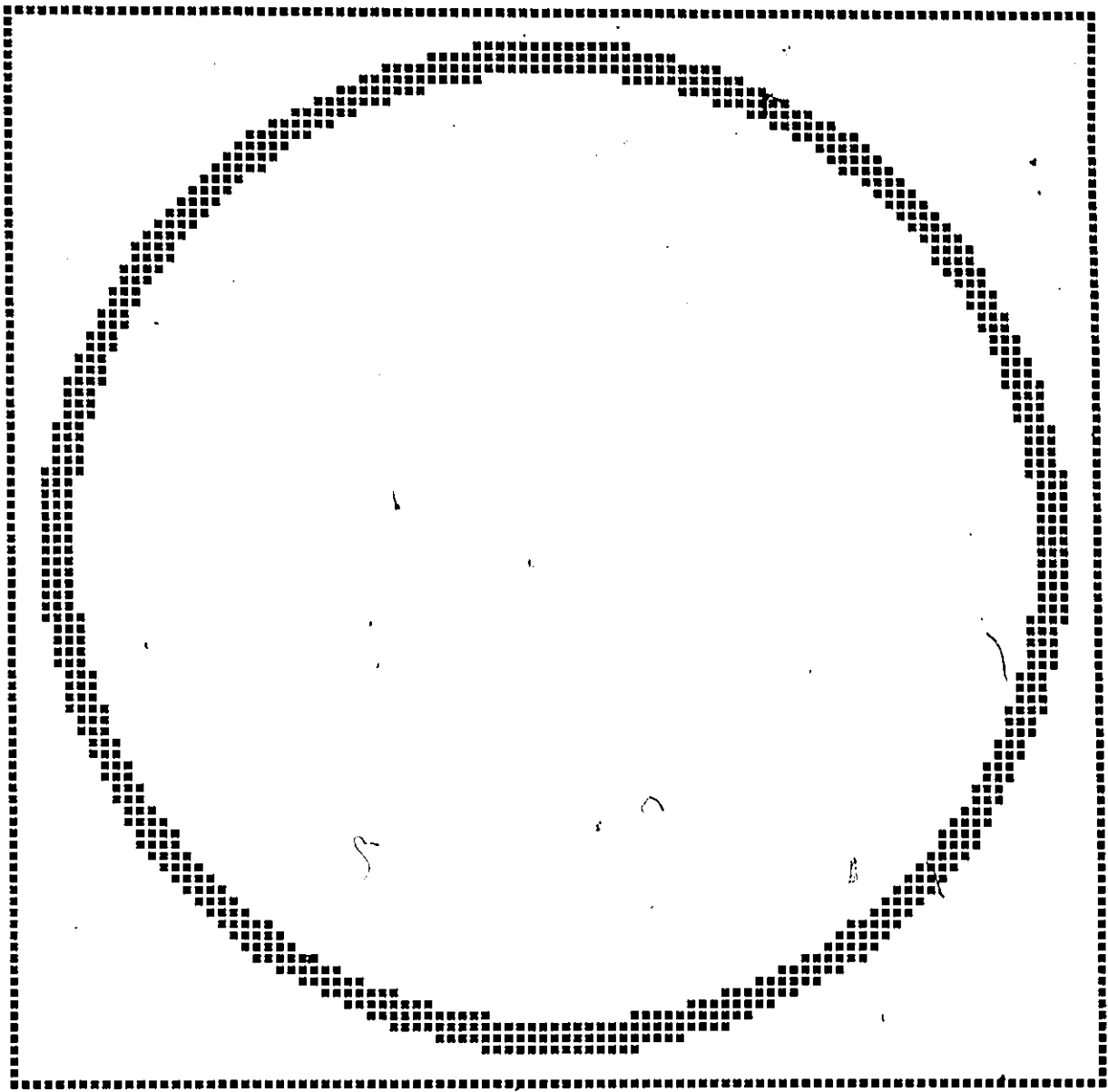


Figure 6. Stimulus circle on retina ('actual' size)

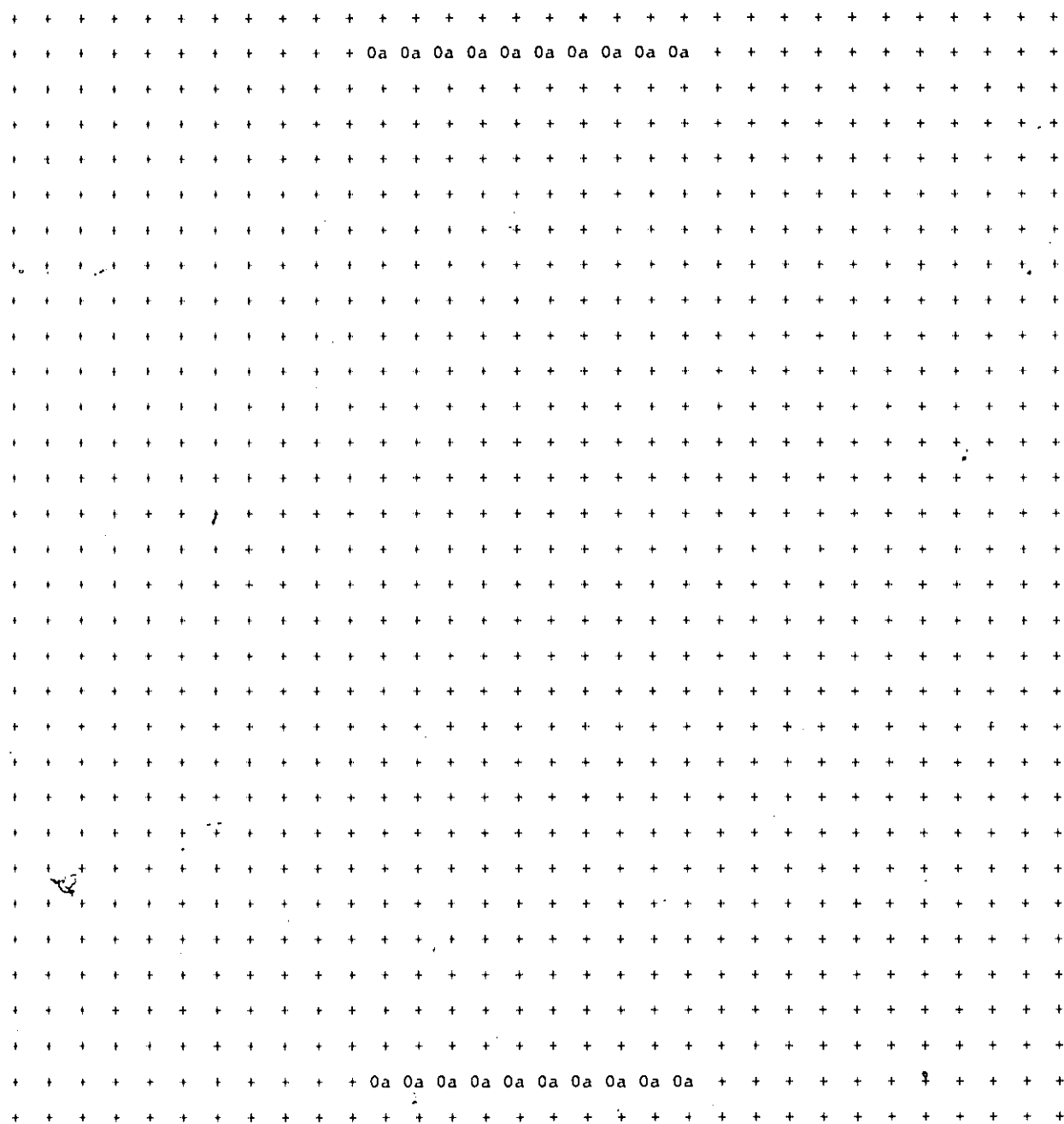


Figure 7. Residue on a layer following circle analysis

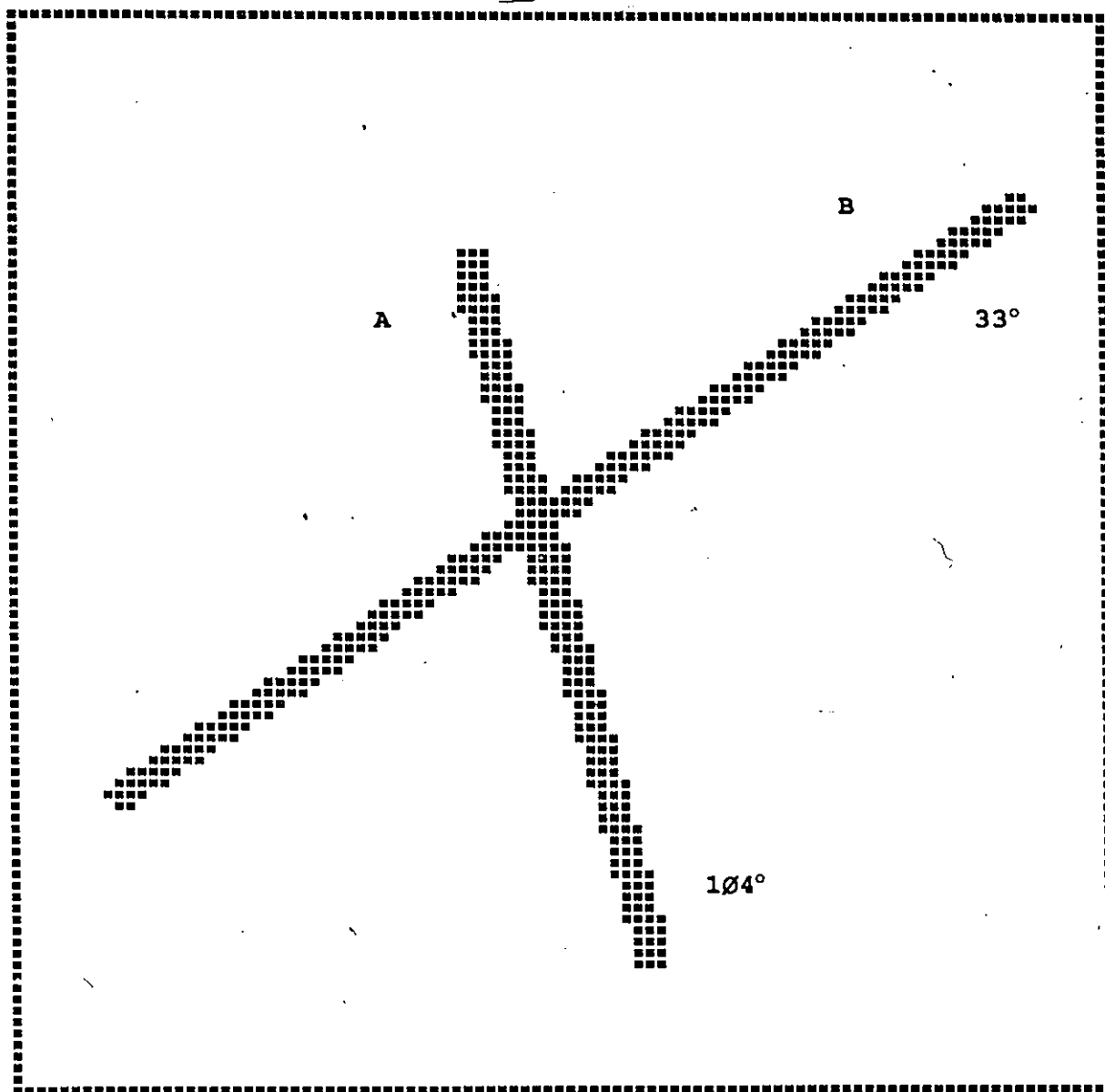


Figure 8. Two stimulus line segments on retina ('actual' size)

(

3

[illegible]

segment: B layer: 11

[illegible]

segment: B layer: 12

[illegible]

segment: B layer: 13

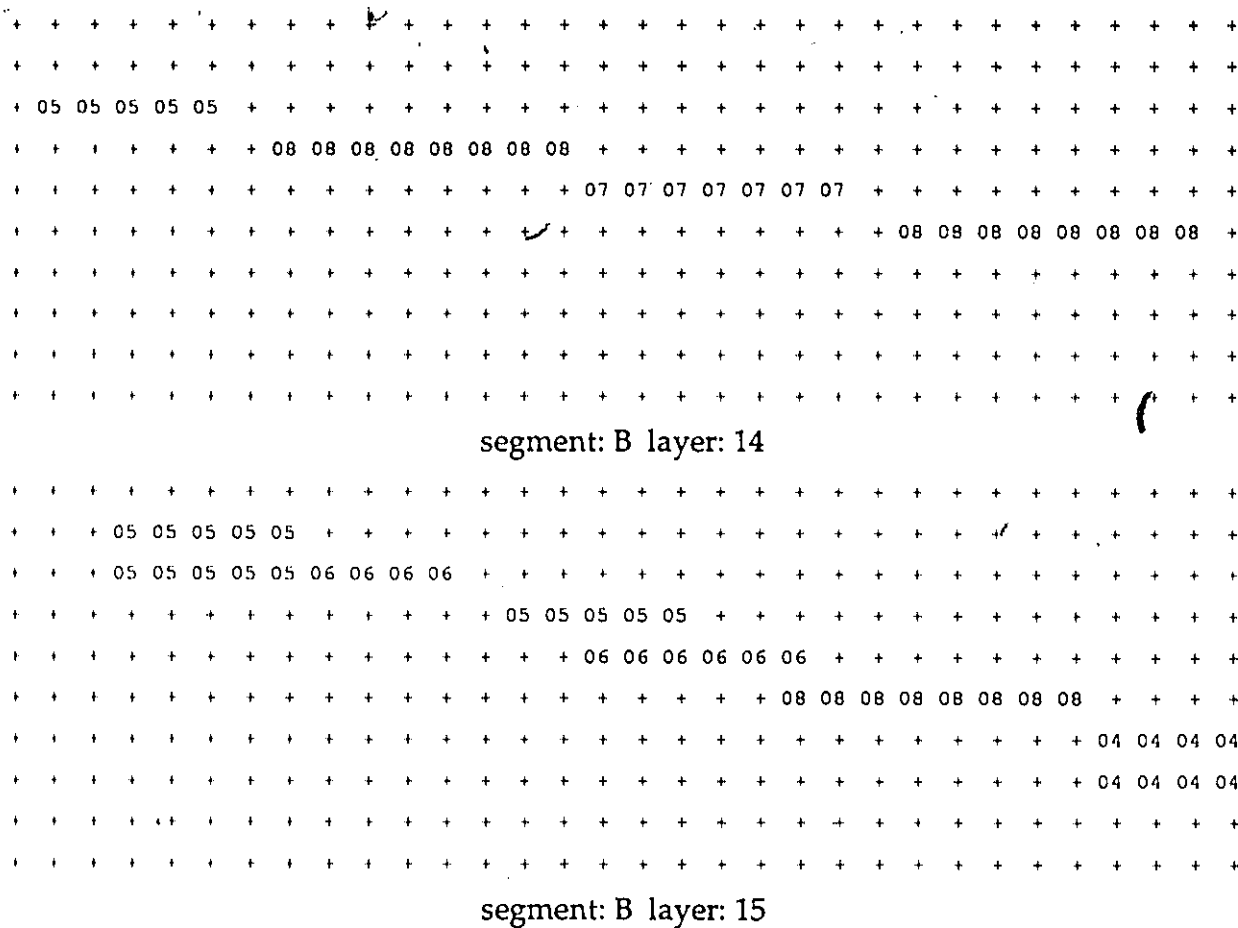
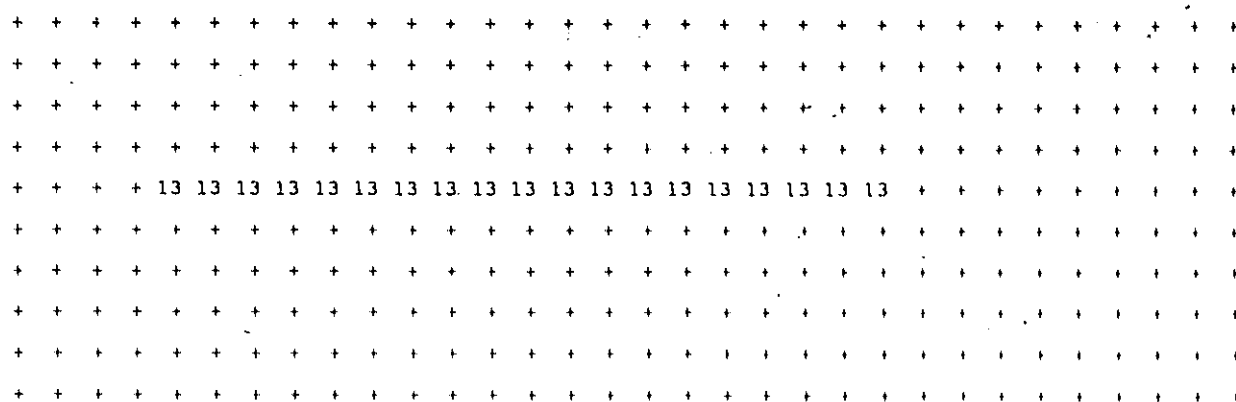
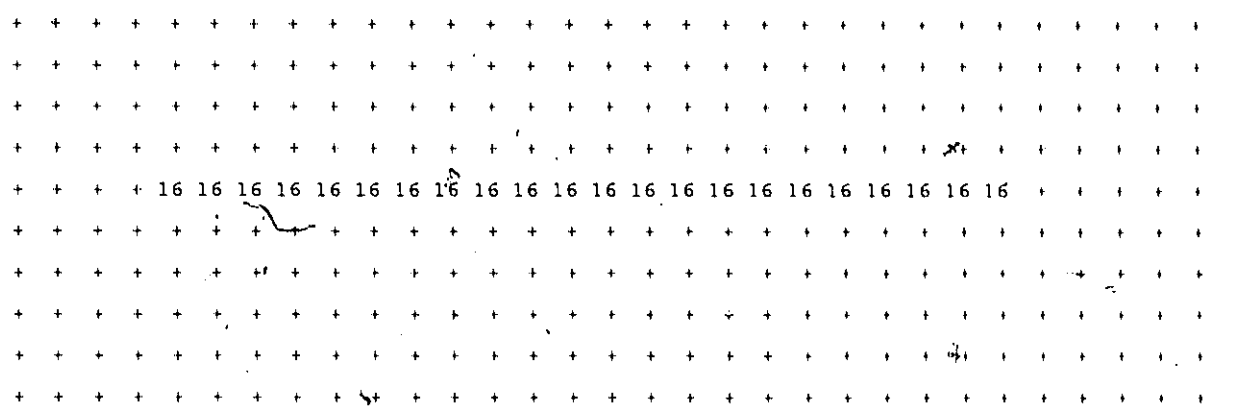


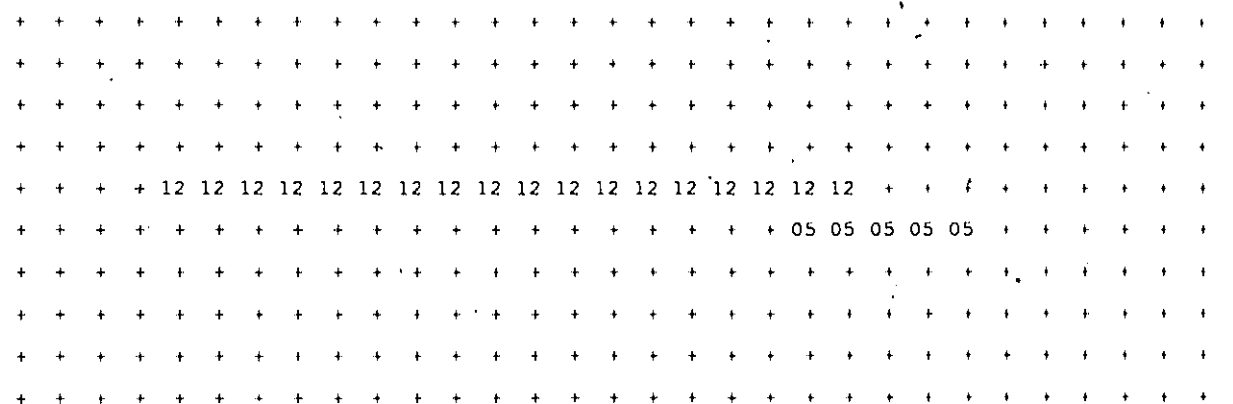
Figure 9. Maximal mapping layer for segment B and its vicinity (layers 8-15). Note that only the middle ten template 'lines' of the layers are shown.



segment: A layer: 34



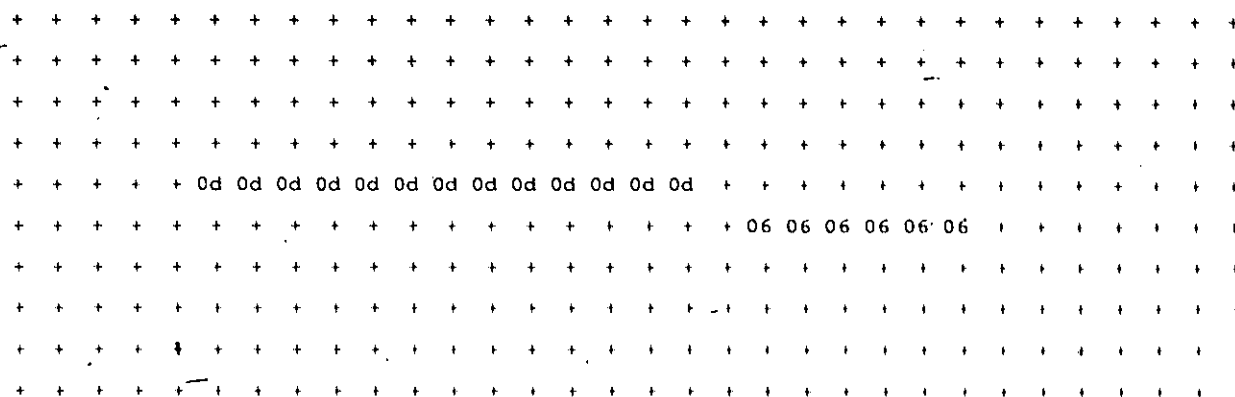
segment: A layer: 35



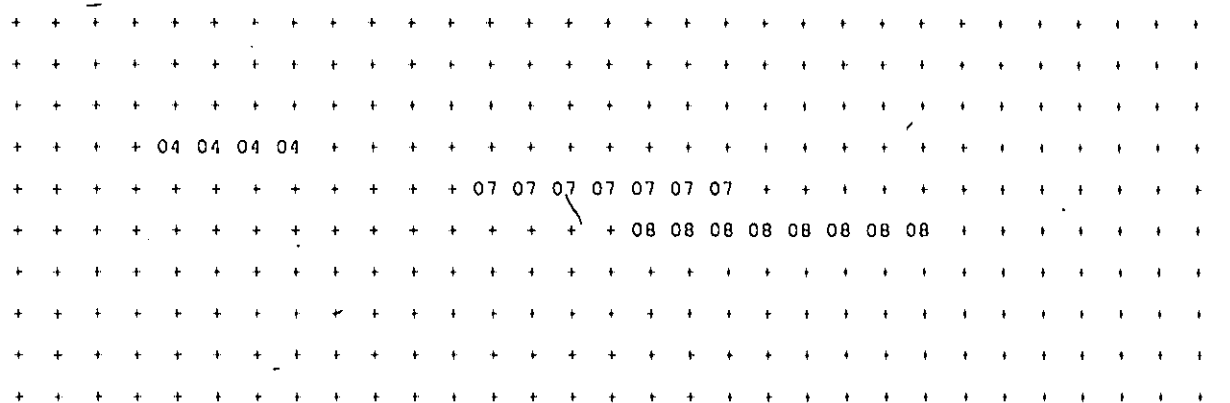
segment: A layer: 36

Figure 10. Maximal mapping layer for segment A and its immediate environment. Projection with a rotational tremor of 1° . Only the middle ten template 'lines' of the layers are shown.

[illegible][illegible]



segment: A layer: 37



segment: A layer: 38

Figure 11. Maximal mapping layer for segment A and its vicinity (layers 31-38). Projection with a rotational tremor of 1.5° . Note that only the middle ten template 'lines' of each layer is shown.

2 Wrapping up

Although the present simulation is far from achieving a general proof of effectiveness for the pile, it has not only demonstrated the workability of a particular realization, but also indicated and clarified certain general aspects of the structure, which to date had not been stated explicitly: the two most important have to do [1] with the extent for tremor (translational and rotational), and [2] with the values and interrelationship between retinal size, field size, and orientational resolution. Since biological relevance for this abstract structure is only to be found within a certain range of the possible conjunctions of parameters, it is important to work out more precisely the optimum values required for one parameter given constraints on one or two of the others: this is the subject of the next section. However as concerns the size of the retina there is an immediate split between the biological data (especially primate) and the ideal requirements of the helicoidal ring: Hubel & Wiesel have gotten us used to the idea of a 'partitionned' retina, while of course the pile only has one, and insofar as it shows some affinity with the "cortical units" then its retina has to correspond to Hubel & Wiesel's retinal field. This in no way endangers the pile's relevance; it merely puts it clearly in context and points a possible way into yet deeper levels of visual processing: the pile could contribute to the forging of a path into them.

3 Retinal size, field size, and orientational resolution

The problem is the following: given that one wants to construct a pile with a resolution of (say) 1° , what are the *minimal* requirements for retinal size and field size?

The first observation to make is trivial. It is that the layers, whatever their size, will have to be displaced (i.e. rotated) in 1° steps from 0° to 179° , so that the pile will have 180 of them.

The second point is more important. If the pile is to discriminate 1° differences of orientation on the *retina* these differences must be *representable* on it: this means it must have enough cells so that a stimulus that is rotated by 1° will not again activate the *same set* of receptors; otherwise it is clear that no mapping difference will ensue (of course this cannot be guaranteed for *all* stimuli; however as a rule we consider the minimal requirements for a large 'central' stimulus). One can readily calculate the range in number of retinal cells necessary for the task: for 1° it means a retina between 56×56 and 114×114 . However more than that is required in order to implement rotational tremor. In the case of a 1° -pile $0,5^\circ$ must *also* be representable on the retina, and this can only be achieved by enlarging it. Again, there is no clear-cut way to calculate this enlargement, for, just as in the above case, $0,5^\circ$ is more or less *well* representable on the retina; now to help us decide what is the 'best' size,

we might adduce a criterion of the following sort: that retinal size should be enlarged sufficiently for $0,5^\circ$ to be very well representable but *not* $0,25^\circ$, the 'next' step down; that is, within this restriction, we choose the maximum dimension allowed. In actual fact this means a retinal grid of 228×228 since $0,25^\circ$ is only 'visible' at the 115th receptor from the middle, implying a retina of at least 229×229 to achieve minimally the next resolution level. This number is extremely interesting for it corresponds rather closely (within 14%) to a similar value of 200 obtained quite independently (see 0.1 above). Now it might be asked why $0,5^\circ$ *has* to be possible on the retina: to answer this question one must imagine a case where rotational tremor is absolutely needed. As an example one can think of a full horizontal line placed in the middle of the retina but with a short piece of one of its extremities 'bent' in a step: this corresponds to a segment whose orientation is somewhere between 0° and $0,5^\circ$; it follows that rotational tremor is needed to maximize the mapping.

This basically settles the matter of minimal retinal size. The rest of the present discussion will therefore assume that it is fixed.

We now move on to the next, and much more involved, question of field size and layer size. Given a retinal size each automatically determines the other.

Again the first consideration is rather trivial. What we want the pile to achieve ideally is a non-ambiguous mapping of a straight line stimulus on the layers so that, through inhibition, only *one* projection segment will survive. Now we will see shortly that this kind of behaviour cannot be ensured for all lengths of straight line stimuli: segments whose length falls below that minimum will be at least doubly detected. This is inescapable because the cotangent drop (see below) is *discretized*; the best we can do is choose the range of length which suits us. For example in our simulation the minimum length is 4 (i.e. $4 \times 3 = 12$ retinal cells); in principle any segment of length less than that will map maximally on at least two successive layers at its length before starting to decrease slowly and with much repetition of identical values, the more so the farther (in either the 'upward' or 'downward' direction) its mapping from the best-fitting layer: for example when it reaches length 2 it stays there for four or five successive layers before dropping finally to 1. (This is why *projection* segments of length 1, 2, and 3 can be discarded: since no *stimulus* segment below 4 is allowed, such mapping results become irrelevant (except possibly for the analysis of curves, but again they are not part of the stimulus) whenever they appear.) We say in principle since those calculations are done with 'central' stimulus segments, i.e. segments crossing, or lying very close to, the center of the retina. There may be an 'off center' correction to add for the general case, but it would be a reduction factor.

If only *one* segment is to survive, it means that we have at least to

obtain *different* mapping lengths on successive layers. Now the least permissible difference is 1, but this is very easy to obtain because the drop of the cotangent is quite rapid, even across 1° (for $\mu=1^\circ$ to 18° , $\cot(\mu) = 58, 29, 20, 15, 12, 10, 9, 8, 7, 6, 6, 5, 5, 5, 4, 4, 4, 4$). To be more precise here, the actual drop in the pile is a *constant* times the cotangent, and for a fixed threshold, the larger the field size the smaller will become this constant. In general the threshold is rather high (66% for 3×3 fields, and 75% for 4×4 fields); the effect of heightening the threshold is similar to that of enlarging field size: it reduces the constant and thus the projection lengths.

Here is a more rigorous demonstration of these properties. We wish to calculate the various projection segments' sizes across the basic orientational resolution increment of the pile. To obtain these values we first use the basic transformation equation giving the new ordinate $Y' = -X \sin(\mu) + Y \cos(\mu)$, where μ is the increment; in our case it will be 1° . To find the desired value one has to 'travel' along a line delimiting the *threshold* of the field until this line crosses over the zero into negative values in the new coordinate system, the one corresponding to the particular layer of projection. Suppose the threshold is such that this line is at $y=1,5$: it is sufficient to seek the fate of this 'mid-line' since the cell's ordinate is represented in our calculations by this middle value, so that if *all* retinal cells forming our hypothetical stimulus are supposed active then it is exactly at the computed point (the zero-crossing) that threshold fails to be reached, and this is what we want, for the 'length' (further corrected for field size and symmetry, as we shall see presently) can then simply be calculated (see Figure 12). Solving for $Y'=\mu(=1^\circ)$ at $y=1,5$ gives $x=1,5 \cdot \cot(1^\circ)$. Two further corrections are then made. The first one is to multiply the result by 2 to take care of the symmetry of our segment about the 'origin,' while the second is a division by " p " (field size), in order to transform the obtained value " x " into a layer length " L ."

Putting these considerations into the equation we obtain:

$$L = \frac{3}{p} * \cot(\mu) \quad (1)$$

For a more general formulation, suppose the threshold is such that the 'travelling' line is at $y=t$ (t will be an integer value $+0,5$). Then

$$L = 2 * \frac{t}{p} * \cot(\mu) \quad (2)$$

It is now easy to see that as the threshold increases, for a fixed field size, L decreases since t decreases, and that as field size (p) is increased, for a fixed threshold, L also decreases. Now the whole point of this discussion is that one does not want the L 's to be too small, and so will choose t and p accordingly. [It so happens, because $t=1,5$ in the two cases,

that (1) works both for 4x4 fields with threshold at 75% and for 3x3 fields with threshold at 66%, but not for example for 6x6 fields with threshold at 66%.] We will get back to this further down.

The generated value, L , will be a decimal which needs to be rounded off, since it will become a component of the ideal inhibition vector: given that one is sure to encounter such a projection length for a sufficiently large stimulus and that inhibition has to be strong enough to 'kill' it, it is a good idea to round off to the *largest* closest integer value.

Using equation (2) with $p=4$, $\Delta\mu=1^\circ$ and $t=1,5$ (threshold 75%) one calculates the following values for L as μ is incremented: 43, 22, 15, 11, 9, 8, 7, 6, 5, 5, 4, 4, 4, 4. As one can see ambiguity sets in at $L=5$, while if field size were at 3, it would set in earlier at 6 (compare with previous list); this loss is minimal however as one also observes that the *range* of 'lengths' is much more restricted in the present case.

With this last item we are now ready to bring all the pieces together. For, now we have a natural constraint allowing to discuss 'interesting' criteria for the selection of layer size. We had indirectly touched upon it with equation (2). There we saw that if field size is much increased then L will decrease - in fact we had known this intuitively from the start. Now given that retinal size is fixed, increasing field size will necessarily *reduce* layer size (and so L). This is exactly where the argument can be brought to play. For consider that nothing can stop one from drawing up a 1° -pile with a 'proper' retina (say 200×200) having large fields (say 10×10), *except* that length discrimination in the pile will be much reduced - in fact to 'lengths' ranging from (say) 4 to 20 only. Now just as one is interested in a good range of *orientational* discrimination, one can equally be interested in a good range of *length* discrimination (and hence 'positional' discrimination): this is the business of layer size, for characterization is only as good as a layer can provide (through summation). Although a range such as 4 to 20 might appear sufficient, there is reason to believe that much higher precision than that can easily be allowed.

In the earlier sections we have alluded now and then to the idea of stimulus imperfection and that it was a function of field size, the bigger the field the less the imperfection. Intuitively it corresponds to the idea that a 'smoother' representation of the stimulus can be obtained through larger fields. To understand this fully we have to tie it to the original problem which brought about the implementation of fields. This was the problem of deciding to which layer cell a given retinal cell mapped when the retina was conceived as structurally identical to a layer. The solution consisted in 'replacing' every retinal cell by a *group* of retinal cells, thus increasing considerably retinal size, and setting a threshold of activation on the layer cells. In fact this solution is nothing other than a *digitalization* (or approximation) procedure for a decision that is *decimal*. In other words it had to be done that way since calculation with

decimals was not allowed in the pile. Now the question of field size, viewed in this light, amounts to the question of how much one has to 'subdivide' every retinal cell in order to get a good mapping decision: a degree of success might be defined formally as the percentage of 'discrete' mappings which correspond to the 'mathematical' mappings, i.e. ones based on the calculation of *area* and a decision based on a threshold also expressed in terms of area and not in terms of number of receptors. It should be clear intuitively that if the threshold is set at 75%, then it does not matter very much, overall, in terms of correct mapping decisions, that field size should be set at 10 or 20 or 50, the *difference* in terms of error of these various schemes amounting to very little. [If one imagines field size to 'grow' infinitely then at the limit mathematical precision would be reached.]

So indeed we can start to see on these grounds that field size should ideally, again, be in the lowest end possible, but not too low - we believe that 3 is too low, i.e. that at that value we are beginning to flirt dangerously closely with and picking up some of the mapping problems which the field structure originally tried to correct. To this can immediately be linked the idea of 'smoothness' of retinal representation: what it amounts to basically is the fact that if field size is 'large' there are not only more units involved in a given stimulus (the stimulus having, by definition, field width) but more importantly that this fact allows for better mapping decisions. Now in the case of straight line segments this might not seem very important but if we wish to enlarge the stimulus domain into a wider variety of curved segments then it does indeed become important - remember the circle. Again one could calculate exactly, say for a certain (non-zero) layer, the percentage of successes and choose to 'correct' field size accordingly. [It is possible, for a given field size, for the degree of success in mapping to vary from layer to layer: this would explain Figure 5.] We have not done this, but it seems a good bet that past a field size of 5 the increment in success rate will be negligible for unary increases of field size. So although from the point of view of mapping or projection field size should be *increased*, the cost in terms of length discrimination can get too high; on the other hand, *decreasing* field size allows for a finer-tuned metric, but endangers the quality of projection itself: an optimum balance has to be sought amidst these interrelationships.

To conclude, let us draw a tentative picture of the 'ideal' 1°-pile: a retina of width 228 receptors, a field size of 4 (and *possibly* 5) and consequently a layer size of 57/45. As for threshold, somewhere between 66% and 75%.

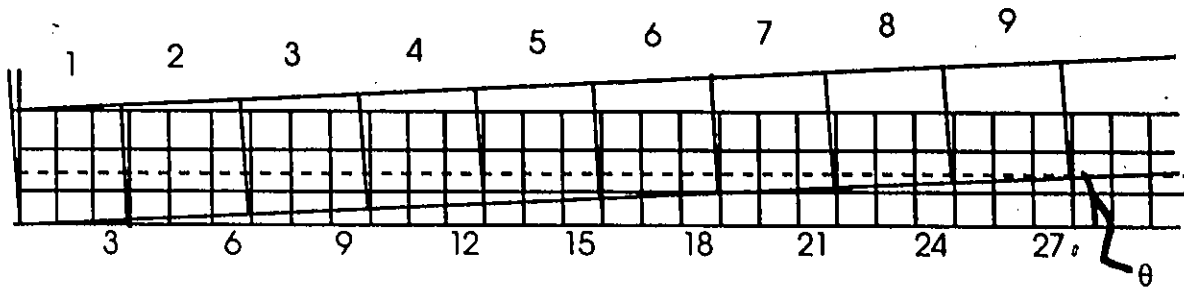


Figure 12. Geometrical demonstration for the calculation of L with $\theta=3^\circ$, $p=3$, threshold 66%. Shown superimposed is the transformed coordinate system corresponding to the layer of projection at $\theta=3^\circ$. Transformation involves a rotation to which is applied a reduction factor associated with field-size, in this case 3. Since threshold is set at 66%, the 'travelling' line (dashed) occurs at $Y=1.5$, the threshold line, since the points (or cells) on this line, whenever they and the cells 'above' them within a given field, all map onto the same layer-cell, ensure that threshold will be reached by that target cell. The reason for considering the cells 'above' them is simply because rotation occurs in that direction, and generality is maintained by the fact that projection lengths are sure to reach a value of 1 at or before 45° of rotation. For another stimulus segment the particular projection layer shown would no longer be the 3° -layer but the first layer below the best-fitting template, the one corresponding to the stimulus' orientation. Thus, supposing ALL retinal cells shown are active, one can see that threshold will be reached in all layer-cells up to the 9th since the six cells delimited, inside every corresponding retinal field, by the 'travelling' line will all 'fall' within the particular layer-cell, whereas this fails to happen past the 9th. To calculate precisely the coordinate of this point one solves the transformation equation for x at $y=1.5$ when $y'=0$, that is, where the 'travelling' line crosses the zero-line in the new coordinate system. Since the reduction factor is lost in this calculation it has to be reintroduced as a correction. For the present case, the solution is $x=1.5 \cdot \cot(3^\circ) = 28.627$; reintroducing the correction factor yields $x=1/3(28.627)=9.542$. Now a last correction has to be made. For convenience we have shown just half the picture, with the center of rotation as reference point or 'origin,' clearly an identical situation is present on the left side of the 'origin.' Therefore to obtain the LENGTH (and no longer the coordinate, which in the simplified case are one and the same thing) of the projection segment we must compute $L=2 \cdot x$, which amounts to 19.08. Note that in the context of OUR simulation the complete adjustments to the cotangent yield a value of 1, i.e. the ideal projection lengths are just the cotangent values. This is not the case in general however. So although Figure 4 happens to be correct, it should be understood as a demonstration of the cotangential nature of projection lengths rather than as an algorithm for their computation. And finally we wish to indicate that this is a (mathematical) IDEAL, and that a particular stimulus may have to be 'helped along' by tremor to reach it. Nevertheless these are the values used for inhibition.

Addendum 1 BASIC programme to generate projection tables 00-2A.
 (Tables 2D-57 generated by editing this programme.
 Lines 62,63 & 70,80 were introduced to correct a bug in
 our version of BASIC. Normally INT is sufficient.)

```

1  DEG: INTEGER M,E,X2,Y2
2  DIM F$(29),A$(10) : M=0 : A=0
3  A$="TABLE00.CRD"
4  F$="000306090C0F1215181B1E2124272A"
10 X=-47.5 : Y=47.5
11 A$(5,6)=F$(M,M+1)
12 PRINT"File ";A$;" is now being generated ..."
20 CREATE A$
30 OPEN\1\A$
40 C=COS(A) : S=SIN(A)
50 X1=(X*C+Y*S)/3
60 Y1=(-X*S+Y*C)/3
61 I=16 : J=15
62 IF X1<0 THEN I=15
63 IF Y1<0 THEN J=16
70 X2=I+SGN(X1)*INT(ABS(X1))
80 Y2=J-SGN(Y1)*INT(ABS(Y1))
81 IF X2>=0 AND X2<=31 THEN 83
82 E=-1 : GOTO 100
83 IF Y2>=0 AND Y2<=31 THEN 90
84 E=-1 : GOTO 100
90 E=(256*X2)+Y2
100 PUT\1\E
110 X=X+1
120 IF X=48.5 THEN 140
130 GOTO 50
140 X=-47.5
150 Y=Y-1
160 IF Y=-48.8 THEN 180
170 GOTO 50
180 CLOSE\1\
190 A=A+3 : M=M+2 : IF M=30 THEN 210
200 GOTO 10
210 END

```

Addendum 2 Extract from the TABLE03.CRD file, the projection table for layer 1 for the first line of the retina; the value "ff" indicates a mapping falling outside of the layer matrix. First pair of coordinates (X,Y) is italicized.

```

ff ff ff ff ff ff ff ff ff ff ff ff ff ff ff
ff ff ff ff ff ff ff ff ff ff ff ff ff ff ff
ff ff ff ff ff ff ff ff ff ff ff ff ff ff ff
ff ff ff ff ff ff ff ff ff ff ff ff ff ff ff
ff ff ff ff ff ff ff ff ff ff ff 00 0d 00 0d 00 0d
00 0e 00 0e 00 0e 00 0f 00 0f 00 0f 00 10 00 10
00 10 00 11 00 11 00 11 00 12 00 12 00 12 00 13
00 13 00 13 00 14 00 14 00 14 00 15 00 15 00 15
00 16 00 16 00 16 00 17 00 17 00 17 00 18 00 18
00 18 00 19 00 19 00 19 00 1a 00 1a 00 1a 00 1b
00 1b 00 1b 00 1c 00 1c 00 1c 00 1d 00 1d 00 1d
00 1e 00 1e 00 1e 00 1f 00 1f 00 1f ff ff ff ff

```

Addendum 3 BASIC programme to generate the rotation (tremor) table of the retina. (A corresponding table, TREM2.CRD, is obtained for rotation in the reverse direction.)

```
10  DEG : INTEGER E,X2,Y2
60  CREATE "TREM1.CRD" : OPEN \1\ "TREM1.CRD"
70  A=1.5
80  X=-47.5 : Y=47.5
120 C=COS(A) : S=SIN(A)
130 X1=X*C+Y*S
140 Y1=-X*S+Y*C
150 I=48 : J=47
160 IF X1<0 THEN I=47
170 IF Y1<0 THEN J=48
180 X2=I+SGN(X1)*INT(ABS(X1))
190 Y2=J-SGN(Y1)*INT(ABS(Y1))
200 IF X2>=0 AND X2<=95 THEN 220
210 E=-1 : GOTO 250
220 IF Y2>=0 AND Y2<=95 THEN 240
230 E=-1 : GOTO 250
240 E=(256*X2)+Y2
250 PUT \1\E
260 X=X+1
270 IF X=48.5 THEN 290
280 GOTO 130
290 X=-47.5
300 Y=Y-1
310 IF Y=-48.5 THEN 330
320 GOTO 130
330 CLOSE \1\
340 END
```

Addendum 4 Extract from the rotation table: mapping of second line
of retina. First pair of coordinates (X,Y) is italicized.

00	01	00	02	00	03	00	04	00	05	00	06	00	07	00	08
00	09	00	0a	00	0b	00	0c	00	0d	00	0e	00	0f	00	10
00	11	00	12	00	13	00	14	00	15	00	16	00	17	00	18
00	19	00	1a	00	1b	00	1c	01	1d	01	1e	01	1f	01	20
01	21	01	22	01	23	01	24	01	25	01	26	01	27	01	28
01	29	01	2a	01	2b	01	2c	01	2d	01	2e	01	2f	01	30
01	31	01	32	01	33	01	34	01	35	01	36	01	37	01	38
01	39	01	3a	01	3b	01	3c	01	3d	01	3e	01	3f	01	40
01	41	01	42	02	43	02	44	02	45	02	46	02	47	02	48
02	49	02	4a	02	4b	02	4c	02	4d	02	4e	02	4f	02	50
02	51	02	52	02	53	02	54	02	55	02	56	02	57	02	58
02	59	02	5a	02	5b	02	5c	02	5d	02	5e	02	5f	ff	ff

Addendum 5 BASIC programme generating the inhibition table.

```

10  DEG : INTEGER E,M,X2,Y2
20  A=3 : M=0
30  X=-46.5 : Y=46.5
40  CREATE"INHTAB.LYR" : OPEN\1\ "INHTAB.LYR"
50  C=COS(A) : S=SIN(A)
60  X1=(X*C+Y*S)/3
70  Y1=(-X*S+Y*C)/3
80  I=16 : J=15
90  IF X1<0 THEN I=15
100 IF Y1<0 THEN J=16
110 X2=I+SGN(X1)*INT(ABS(X1))
120 Y2=J-SGN(Y1)*INT(ABS(Y1))
130 IF X2>=0 AND X2<=31 THEN 150
140 E=-1 : GOTO 180
150 IF Y2>=0 AND Y2<=31 THEN 170
160 E=-1 : GOTO 180
170 E=(256*X2)+Y2
180 PUT\1\E
185 IF M=1 THEN 220
190 A=A+3 : IF A>18 THEN 210
200 GOTO 50
210 M=1 : A=180
220 A=A-3 : IF A<162 THEN 240
230 GOTO 50
240 A=3 : M=0 : X=X+3 : IF X=49.5 THEN 260
250 GOTO 50
260 X=-46.5 : Y=Y-3 : IF Y=-49.5 THEN 280
270 GOTO 50
280 CLOSE\1\
290 END

```

Addendum 6 Extract from the inhibition table: these are the twelve coordinate pairs (Y,X) specifying the inhibition path of a layer cell at (00,0c).

```

00 0d  00 0e  00 0e  00 0f  00 10  00 11  ('downward')
1f 14  1f 15  1e 15  1e 16  1e 17  1d 18  ('upward')

```

Addendum 7 First page of helicoidal ring simulation programme

```

1          DS  5E00H          ;100H-48FFH holds TABLE** (18K)
2                                     ;4900H-6CFF is RETINA (9K)
3                                     ;8000-A3FF is storing space
4                                     ;100H-60FFH holds INHTAB (24K)
5 BUFFER    EQU 0A4C0H        ;housekeeping buffer
6 FLNAME     DB  'B:TABLE'
7 TABNUM     DB  '00.CRD/'
8 ROTTRM     DB  'B:TREM0.CRD' ;rotational tremor table
9 INHFIL     DB  'B:INHTABL.LYR' ;inhibition table
10 STIM      DB  'B:CIRCLE.SEC' ;circle table
11                                     ;approximate (integer) values of
12                                     ;cotangent form 3 to 18 degrees
13                                     ;at 3 degree intervals
14 COTAN      DB  20,10,7,5,4,4
15 INHVEC     DS  6            ;adjusted cotangent/inhibition vector
16 HFIELD     DW  -33,1,1,1EH,2,1EH,1,1,7
                                   ;inhibition 'field' vector
17 TRMVEC     DW  -20H,1,20H,20H,-1,6
18 TREM       DB  0,1,2,60H,62H,0C0H,0C1H,0C2H,8
                                   ; 'standard' tremor vector
19 HPOINT     DW  0            ;pointer to next 24-vector in INHTAB
20 BANK       DB  2            ;current bank of inhibited layer
21 LYRST      DW  8000H        ;head of inhibited layer
22 COUNT      DW  0FFH
23 LYRMEM     DW  8000H        ;head of inhibiting cell layer
24 BKMEN      DB  02          ;bank of inhibiting cell layer
25 FCB        DS  33
26 BUFCUR     DW  100H        ;initial value of write buffer
27 TABADD     DW  100H        ;current address in TABLE**
28 LAYER      DW  7C00H        ;layer starting address
29 LAYRNO     DB  -3          ;(address of curr scan pos is stacked)
30 CTR        DB  00          ;'switched' analysis counter (BESFIT)
31
32 START      LD  HL,0E8H      ;100H-18H (24)
33             LD  (HPOINT),HL ;initialize pointer
34             CALL CLRET      ;clear retina
35             LD  HL,STIM     ;loading CIRCLE.SEC
36             CALL MDUMP

```

Appendix II

Note on the 'hexconv' Function

A retina, made up of discrete units called cells arranged in an orderly two-dimensional manner, can easily be simulated in a computer through the use of arrays. Simulating an activated retina, on the other hand, requires a description of the stimulus compatible with this representation. Thus it requires the transformation of a (mathematically) continuous description of the stimulus into a discrete one: we call this partitioning, and it is achieved by sampling some points along the stimulus and calculating the corresponding cell (or array) coordinates. Usually, but not necessarily, the borders of these idealized retinal cells are associated with the *integral* values on the axes of the cartesian plane (0,1,2,3, ...). These values are then used as *indices* in the retinal array in the computer; thus $RET[x,y]$. So whenever a cell at (X,Y) is activated, $RET[X,Y]=1$, where X is the rounded value of x, and similarly for Y, (x,y) specifying a point on the stimulus; otherwise $RET[X,Y]=0$. This works well for what might be called a 'cartesian retina.' The transformation from *continuous* values of the stimulus to the *discrete* values of the associated cells will involve a *first* discretization of the continuous stimulus (in the x-axis for example): the rule is simply to cut up or partition the axis in pieces that will preserve continuity in the cells, assuming it exists in the stimulus of course. Thus we would not sample the continuous stimulus at every x+2th value since then the sampled points could 'jump over' some cells, leaving undesired wholes in the array; on the other hand sampling must occur below a width of 1 to ensure that two successive sampling points will fall either within the same cell or two adjacent ones (an adjacent cell being defined as any member of the eight surrounding cell); in practice, this implies that some redundancy is bound to result from the conversion, but it can be minimized by selecting a sampling width that is close to 1. Usually also, as indicated above, one will start with a first approximation on the x-axis, generate the corresponding y-values, and discretize both; in some cases (e.g. with asymptotic curves in the y-axis, if the case ever presents itself) the reverse order would be advised: the first approximation would then be done along the y-axis, then the corresponding x-value generated, etc.

Most graphics cards work with a (discrete) cartesian grid. Thus a 'cartesian retina', corresponding to this sort of grid, seems very natural since it simplifies the interface to the computer programme. Needless to point out that arrays are equally 'cartesian.' This was typically the case in the simulation of the helicoidal ring (see Appendix I), that, incidently, was first conceived with a cartesian retina. However the natural structure of a retina of receptors is much closer to a hexagonal grid than to a

cartesian grid, a hexagonal arrangement arising when one packs equal-sized chips or pennies or receptors on a plane. This would be interesting but of little value were it not for the critical fact that the neighbourhood relations are much less ambiguous in a hexagonal arrangement than in a cartesian one, where the problem of corner cells arises: in the hexagonal pattern, each cell has six neighbours with equal (geometrically) clear status: each shares one of the center cell's *sides*.

It is of course a trivial matter to map a hexagonal grid into a two-dimensional array in the computer, *once* the hexagonal coordinates are known. As for neighbourhood relations, much as in the cartesian case, they can easily be represented by a simple formula involving the character of the row-value (oddness or evenness), etc., allowing for access to adjacent cells within a programme. It is also possible to transform a cartesian display (on a graphics monitor) into a hexagonal one: this involves a reduction in resolution however since *many* points are needed to represent *one* (hexagonal) cell.

The 'hexconv' function, written in C, takes as input an ordinary cartesian coordinate (x,y) and finds the hexagonal cell (hx,hy) into which it falls. This is achieved in five steps.

The basic feature of the algorithm, the thing that renders it a comparatively simple matter, is the positioning of the cartesian axes onto the hexagonal grid (see Figure 1): it can be seen that the cartesian unit points (0,1,2, ...) are placed such that they fall, on the x-axis direction, on the middle point of a (unit length) hexagon, whereas in the y-axis direction the 'unit' marker is set down directly in the hexagonal mode, the geometry of axis superimposition placing the first marker at $\sqrt{3}/2$, the second at $\sqrt{3}$, etc. The algorithm, contrary to the cartesian-grid method, does not proceed directly to compute, by simple rounding, the coordinates of the cell into which a particular point falls. Two intermediate steps are needed to achieve these final values. The first consists in confining the x-value inside a unitary integral interval, defined by the boundaries and the middle of the x-axis hexagon, which happen to occur at multiples of $1/2$. Thus if for example, $x=1.4$, the upper boundary will be 3 (since $1.4 < 3 \cdot (1/2)$) while the lower boundary will be 2 (since $1.4 > 2 \cdot (1/2)$); a similar procedure is carried out for the y-value, but this time of course the units are no longer of length $1/2$ but $\sqrt{3}/2$. This computation for y only cuts the possibilities down to two; for the x-value there is a similar but less obvious ambiguity, stemming from the fact that in the y-axis direction the hexagonal cells are not aligned but 'crooked.' More on this 'side effect' below. The second intermediate step resolves both of these ambiguities, according to a rule based on the observation of a simple pattern among these intermediate pairs as they relate to actual hexagonal cell coordinates: they 'go together' according to oddness and evenness, and since every pair, being made up of two successive integral values, must have one of each kind, it is easy to

couple them by 'type' (odd with odd, even with even), thus defining *two* adjacent cells as most likely candidates for inclusion of the given point (following the previous example, and Figure 2, the two pairs are [2,4] and [3,3]). Then it is simply a matter of reverting back to the cartesian coordinates of those cells' center and of choosing the cell whose center is closest to the point.

The first step (cf. programme below), by rounding, computes the intermediate 'hexagonal' x-values; x_1 is always the lower value: this works because, for example, x is declared type double whereas x_1 is declared type integer so the assignment itself forces the conversion. Then the y-value is also rounded up and down to the first approximation. Next, the first true hexagonal x , h_x , is calculated, using the fact that it will be $1/2$ of the intermediate unit x falling on the center of the cell (x_1 is either odd or even, and must be tested). Thirdly, the pairs of x 's (x_1, x_2) and y 's (y_1, y_2) are coupled into *center* coordinates so that the values are matched according to oddness or evenness: this is done because the centers of the *hexagons* are characterized, in terms of their (x, h_x) coordinates, by such combinations (see Figure 2). There are four cases to consider in all and step three simply switches the x 's, if need be, or exchange the y 's. Testing for oddness/evenness involves doing a bitwise AND on the values with 1. Having trapped the point (x, y) within the area defined by the two centers, we still have to decide which is closest. This is done in step four, where it should be noted that the y -values are corrected for the scaling factor of $\sqrt{3}/2$, and the x -value for $1/2$, because the context reverts back to a cartesian one. Finally step five adjusts h_x , if necessary, according to whether h_y AND (logical this time) the least x (anciently x_1 , but since it might have been switched around with x_2 we had to remember it in $x_1\text{mem}$) are odd: this correction is necessary because although our first approximation to h_x was correct in itself, because of the non-linear arrangement of hexagonal cells one sometimes has to start counting $h_y(=0)$ at the left-side cell of the chosen h_x . We should make it clear that starting at the baseline ($h_y=0$) one moves in the hY axis first to the *right*, then to the *left*, then again to the right, and so on. This is indicated in the right of Figure 2.

Although as shown here this function works for just one point (x, y) it can be modified to accept a *pointer* to an array of such coordinates, whereupon the function will transform all of them.

(x, y) can also range over the entire cartesian plane; $\text{hexconv}(x, y)$ will work correctly for negative values as well as for positive ones, as illustrated here. It has been used however in a context where the input coordinates were already discretized. It turns out that transforming a set of such values will cause the appearance of wholes in the hexagonal grid, i.e. discontinuities in the stimulus. This is taken care of by a further function that tests for it and produces a filling-in cell in the hexagonal transcription. This function is not reproduced here.

The hexagonal grid, like the cartesian one, places the origin at the top left corner. This means not only that only positive coordinate values are used to access it, but also that the direction along the y-axis is reversed: increments in y involve a *downward* displacement along the axis. Again the function described here does not take care of this further transformation; it is up to the programme using it to do so.

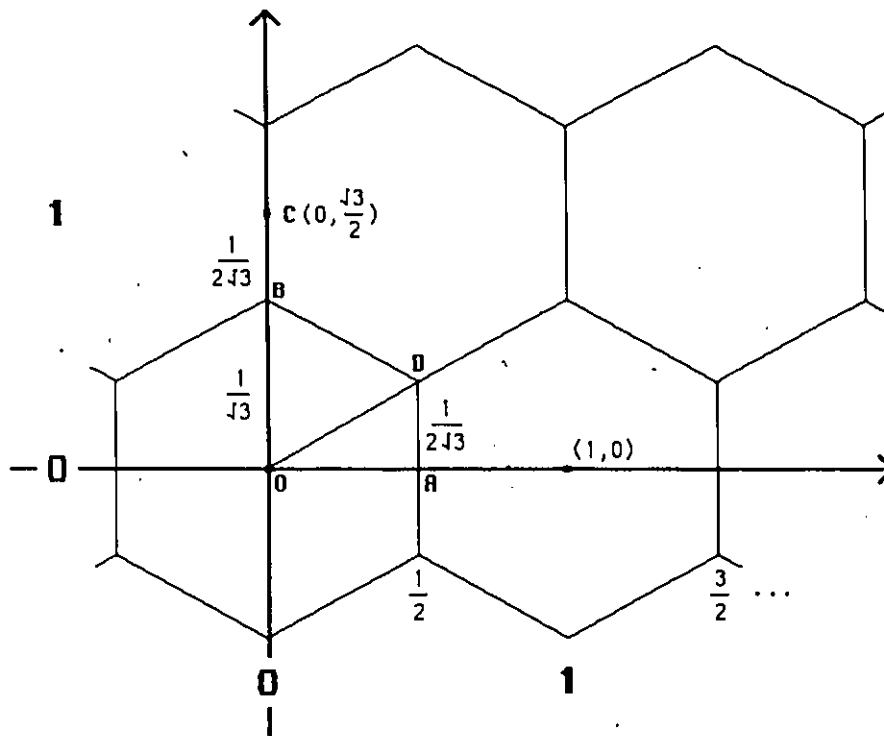


Figure 1. Geometry of cartesian coordinate system superimposed on a hexagonal grid; choosing $OA=1/2$ we obtain $OB=OD=1/\sqrt{3}$ and $BC=AD=1/2\sqrt{3}$. Thus $OC=OB+OD=1/\sqrt{3}+1/2\sqrt{3} = 3/2$, i.e. $y=1$ is at $(0, \sqrt{3}/2)$.

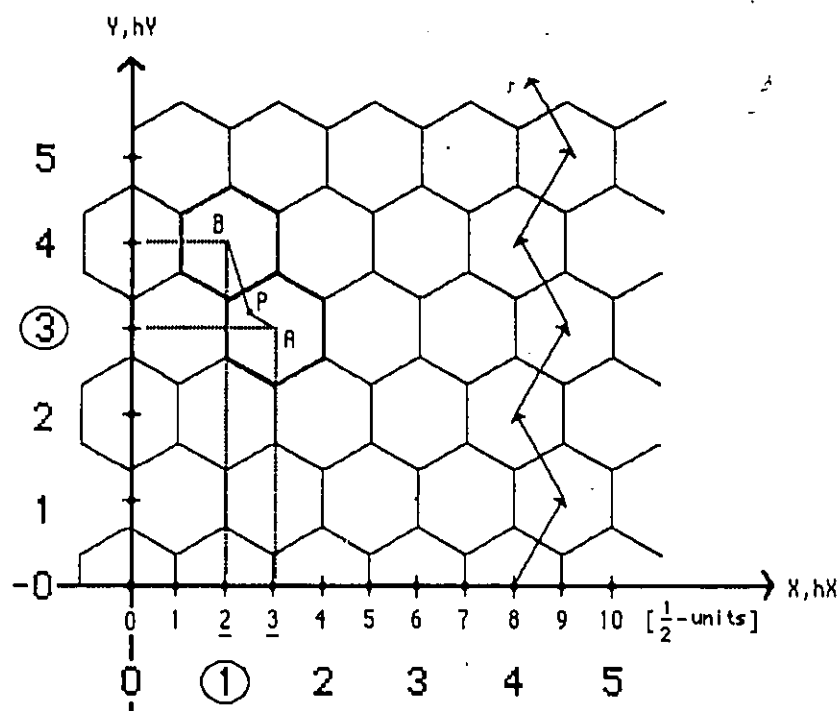


Figure 2. Superimposition of cartesian coordinate system over hexagonal grid. Both 'X' and hX units are indicated, while only hY is marked. On the left is shown an example of a point P boxed-in by the 'X' and hY coordinates. Once the two closest centers are identified (A & B) it is a matter of choosing the closest one to P, and, if necessary, to apply a correction to hX - this would not be the case for P, whose hexagonal coordinates are 1,3 (circled).

Here is the function listing.

```
hexconv(x,y)
double x,y;{
int x1,xlmem,x2,y1,y2,hx,hy,temp;
double c=1.7320508075688773; /* sqrt(3) */
x1=2.0*x; x2=x1+1; /* round x down and up wrt 1/2-units */
y1=2.0*y/c; y2=y1+1; /* same for y, but wrt hY */
xlmem=x1; /* remember least x */
if(x1 & 1) hx=x2/2.0;
else hx=x1/2.0;
if(y2 & 1) {if(!(x2 & 1)) {temp=x2; x2=x1; x1=temp;}} /* switch x's */
else if(x2 & 1) {temp=y2; y2=y1; y1=temp;} /* switch y's */
if((sqr(x-0.5*x1)+sqr(y-y1*c*0.5))<(sqr(x-0.5*x2)+sqr(y-y2*c*0.5)))
hy=y1;
else hy=y2;
if((hy & 1) && (xlmem & 1)) --hx; /* adjust for 'side effect' */
}
```

Appendix III

The Simulated Arm I: Logic of Implementation

In the first part of this third Appendix we provide a detailed technical account of the computer implementation of our "particular case of control," namely, a virtual anthropomorphic arm with seven (7) degrees of freedom (d.f.s). The second part (see below) consists of a listing of the programme. The programme is based upon some elementary geometrical principles that are instantiated in the context of the specific requirements of the implementation. These requirements can be recalled briefly. We wanted a visible, dynamic, human-like, effector with seven d.f.s that could easily be driven at that level, and that could eventually be hooked up to a control structure; moreover the output had to be as realistic as possible. This was achieved through a 3-D graphic display of the precomputed successive positions of the arm's segments. These intermediary positions were separately calculated in a series of four steps.

The arm consists of three segments: the upper arm, the forearm, and the hand, with an associated 3-2-2 scheme of d.f.s, starting with the shoulder (Figure 1). That is, much like the human arm, each segment being taken singly, the upper arm's motion occurs in three independent axes while that of the forearm and hand do so in two. Segments are elongated rectangles whose lengthwise proportion is in the ratio of 2:2:1, again approximating human anatomy.

Degrees of freedom differ in the axes with respect to which they can be expressed, and hence must be computed. However all of them involve *rotations*, either about an axis in an adjoining segment or about the axis of the considered segment itself (cf. Figure 1). The hand's d.f.s can be associated with a yz-coordinate system affixed to the forearm's extremity (the wrist), the forearm's axis being the x-axis: thus rotations occur, and are computed, about both the y- and z-axes for the hand. Movement of the forearm occurs with respect to its own axis as well as about the z-axis of the upper arm's coordinate system. Finally the upper segment is driven by rotations about three axes, one of which, again the x-axis, is the segment's own, while the yz-coordinate system, as opposed to all others, belongs outside the effector system: it corresponds to the screen's frame of reference, with the z-axis perpendicular to the viewing surface.

Driving the arm is accomplished through the specification of three sets of values: [1] a single parameter, N , the total number of steps to be calculated and displayed, [2] a set of seven values for each d.f., defining

the effector's initial position, $IP=(ip_j | j=1,...,7)$, and [3] a corresponding set of seven amplitudes of displacement, or extents, for d.f.s; corresponding elements from these two groups are paired. Thus the programme is called in the following manner:

arm N $ip_1 \delta \cdot ext_1$ $ip_2 \delta \cdot ext_2$... $ip_7 \delta \cdot ext_7$

Initial positions range from -89° to $+89^\circ$, while extents, also being signed ($\delta=\pm 1$), and allowed within the $[0^\circ, 179^\circ]$ interval, provide for the specification of direction. N is limited to a maximum of 89 and a minimum of 0. Overflowing, or the computation of a d.f. motion (sum of current position and specific increment) leading to a position beyond the allowed boundaries, is checked for at each step: whenever it occurs further increments are inhibited there, that is, the d.f. remains stationary at the extremum while the rest of the d.f.s proceed through their preordained paths.

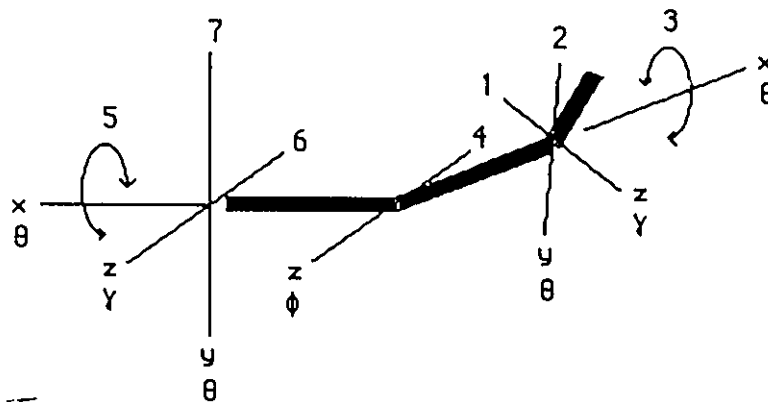


Figure 1. The simulated arm (dark). From right to left, the three segments are drawn (hand, forearm, upper arm) along with three additional pieces of information: (i) the axes of rotation for the d.f.s (x,y,z); (ii) the order of computation of movements (numbers); (iii) the angle names used in calculations (greek letters). The order of prescription within the command (see above) is 2,1,3,4,5,7,6.

We have argued in section 2.2 that a basic characteristic of biological movement, namely its naturalness, also derives from a nearly synchronic activity in sffs, giving rise to unequal speeds at the various d.f.s involved. In the programme this is implemented quite simply in the

following manner. Since the whole movement must occur within the specified number of steps, N , each extent is first divided by N , i.e. broken up into N parts; given that in general these values differ from one d.f. to another, it follows that the computed fractionary motions on intermediate steps of the movement will vary in amplitude also; hence the different speeds.

All movements are computed through rotations operated on the 12 (x,y,z) points defining the segments that, initially, are 'piled' on top of each other on the x -axis. Simulating rigid-body motion is then achieved by transforming groups of points making up either one, two, or three segments (the whole effector). To construct and preserve the links between the segments throughout these various transformations, we proceeded in four stages, each of which comprises three operations. In the first stage, the group of four points constituting the hand is transformed according to the second through the fifth parameters pairs provided on calling the programme; second, the hand is translated up the x -axis just beyond the forearm; finally, the four new points are stored back in the temporary matrix $\text{SegM}[12][3]$. These procedures are then repeated for the transformation of the eight points of the hand-forearm, with a further translation of this pair just beyond the upper arm, the update occurring now on these two segments' new positions. The third and fourth passes split up the management of the remaining three d.f.s, which affect the whole effector now, in a 1-2 sequence, the first involving only the rotation at the shoulder, the second taking care of the remaining pair of d.f.s. Note that since the function $\text{init_tr}()$ expects two angles, the second one must be 0.0 in the third call.

Once the final position of the arm is computed for one step we proceed to a final transformation for purposes of 3-D display. This involves the well-known projection equations (cf. Foley & Van Dam 1982) that map the 3-D arm, i.e. its twelve (x',y',z') points in $\text{SegM}[12][3]$, into a projection plane selected according to the available screen size and in such a way that no (undesirable) clipping is necessary. These viewing coordinates are stored in a new matrix, $\text{PM}[N][2]$. Whereupon the loop repeats for a next step until all N steps have been run through. Note that each step restarts at what is called "position 0," the *internal* initial position which we described above (not to be confused with IP, the one provided by the caller). This leads to two consequences, one of minor, the other of major, importance. First, the "if(j)" conditional, which ensures execution of the increment operation on every step except the first, allows for a smooth integration of the calculation of IP within this framework of movement computations: the arm is simply moved from the internal initial position to IP. Secondly, by 'reinitializing' the arm's position each time, and calculating its next position from there instead of from the *previous* one, we greatly simplified the required transfor-

mation matrices: in view of the fact that such matrices are already much more involved for the case of a *single* rigid body, the computational gain can be appreciated even more here given the fact that we are dealing with a *composite* one. Of course nothing of this "internal business" shows in the final product. But even with this simplification, the relatively large weight of the computations prohibited 'real-time' display of the intermediate positions; this is why they were all pre-calculated and stored beforehand. Thus the final programme segment simply reads off sets of 12 'step' coordinates, calls the routines to draw the connecting lines, erases the previous output and displays the new one.

If we define $R(\phi, z)$ to mean the rotation transformation about the z -axis, and similarly $R(\theta, y)$ to mean the rotation about the y -axis, then not only is it not true in general that their conjugate, namely $R(\phi, z) * R(\theta, y)$, is not commutative, but much more importantly for us, neither is it true that after the *second* transformation, viz. $R(\phi, z)$, a given resultant vector (x', y', z') will have maintained, with respect to the xz -plane, the angular relation that was specified in the *first* transformation, viz. $R(\theta, y)$, although of course the relation to the xy -plane will be the desired one. For our purposes this will clearly not do; d.f.s are *independent* from each other; and so a displacement directive about a given axis should not be 'contaminated' upon execution by its coupled directive, which, happening to be computed in second place, will therefore be satisfied. This introduces an unwarranted dissymmetry in the system which undermines the true independence of d.f.s. There is however a way out.

By modifying ϕ slightly, one can arrange for this bias to disappear; a simple trigonometric function of ϕ and θ , namely $\arctan[\cos \theta * \tan \phi]$, provides a new angle γ , which, when substituted for ϕ in the compound transformation above will correct for the error and produce the desired independence: γ effectively 'anticipates' upon the disturbing effect of the second transformation and nullifies it. This function was integrated within `init_tr()`.

However an inescapable dissymmetry, which arises as a consequence of the layout of the d.f.s on the effector, does enter the picture. The axes of rotation of the hand, for example, are both orthogonal to the axis of its reference sff, the forearm, but this breaks down in the latter's case: the forearm also revolves about its own axis. This again happens at the shoulder joint. A closer look at the layout reveals that for handling couples of d.f.s only two different conjugate matrices are needed to do the job, viz. $R(\gamma, z) * R(\theta, y)$ and $R(\theta, x) * R(\phi, z)$. The first is used for movements at the two hand d.f.s and the two 'orthogonal' shoulder d.f.s, which need the correction, while the second is used for the two intermediate pairs, which do not need the correction; for the last of this intermediate

pair, as mentioned above, $\phi=0$. Note that ϕ and θ are reversed in the two expressions; this is because rotations about a sff axis are always effected first; since "position 0" (i.e. the internal initial position) is used, the first rotation occurs about the x-axis for these cases, while in the others the first rotation is about the z-axis instead.

If one writes down the proper equations and develops the products, an interesting fact stands out: both matrices are identical except for the disposition of their elements; this enabled us to shorten somewhat the programme. Here are the transformation matrices:

$$R(\theta, x) = \begin{pmatrix} 1 & 0 & 0 \\ 0 & \cos \theta & \sin \theta \\ 0 & -\sin \theta & \cos \theta \end{pmatrix}, \quad R(\theta, z) = \begin{pmatrix} \cos \theta & \sin \theta & 0 \\ -\sin \theta & \cos \theta & 0 \\ 0 & 0 & 1 \end{pmatrix},$$

$$R(\theta, y) = \begin{pmatrix} \cos \theta & 0 & -\sin \theta \\ 0 & 1 & 0 \\ \sin \theta & 0 & \cos \theta \end{pmatrix};$$

wherefore

$$R(\gamma, z) * R(\theta, y) = \begin{pmatrix} \cos \gamma * \cos \theta & \sin \gamma & -\cos \gamma * \sin \theta \\ -\sin \gamma * \cos \theta & \cos \gamma & \sin \gamma * \sin \theta \\ \sin \theta & 0 & \cos \theta \end{pmatrix}$$

etc.

This is the logic behind the implementation of the arm. It was first programmed in the way described in order to make certain that the simulated arm indeed behaved the way it was supposed to. Once this was ascertained, the code segments that accepted and prepared the input parameters and fed them into the four subsequent steps of computation were replaced by the code segments that implemented the (second) control model described and discussed in sub-section 2.2.1.1. These new code segments (especially the "fetch()" function) appear in the listing provided below. Their function is equivalent to the one described above in which all the parameter values were individually provided upon calling the programme; this time however they were specified according to the control model, and then degrouped into the local values for each d.f.; this occurs in the "main()" loop.

The virtual humanoid arm described here was originally developed in C on a CromemcoTM microcomputer, from which it was ported to a MacintoshTM, on which the control model was then simulated.

II: The Programme



```
#include <quickdraw.h>
#include <event.h>
#include <osutil.h>
#include <math.h>
#include <memory.h>
#include <toolutil.h>
#include <init.h>

#define GrRPtr &(thePort->portRect) /* GrafPort Rectangle Pointer */
#define PI 3.14159265
#define RAD PI/180.0
#define DEG 180.0/PI
#define PIby2 PI/2.0
#define PIby3 PI/3.0
#define HOR 341.0 /* horizontal adjustment */
#define VER 171.0 /* vertical adjustment */
#define L1 44.0 /* segment 3 length */
#define L2 2.0*L1 /* segments 1 & 2 length */
#define WD 5.0 /* width of segments */
#define JP L2 /* push for transformations */
#define NoMOMENTS 4
#define PLANE 700.0 /* projection plane distance */
#define PDIST 220.0 /* distance of arm from PLANE */

/* global variables */

int scrIP[12][2] = {402,143,402,136,400,123,400,129,407,174,408,167,
                   402,136,402,143,341,174,341,167,408,167,407,174};
int k3[9] = {0,0,0,1,1,1,2,2,2}, k3[9] = {0,1,2,0,1,2,0,1,2};
int k4[9] = {1,0,2,1,0,2,1,0,2}, k4[9] = {1,1,1,0,0,0,2,2,2};
float segMm[12][3] = {0,0,-WD,0,0,WD,L1,0,WD,L1,0,-WD,0,0,-WD,0,0,WD,
                     L2,0,WD,L2,0,-WD,0,0,-WD,0,0,WD,L2,0,WD,L2,0,-WD};
/* 3-D coordinates of initial position of three segments */

main()
{
    int ii,jj,p;
    Rect armrect,ctrlrect,ex0; /* arm's area & control area */
    long parms[101]; /* input quintuplets: global values */
    float armparms[15]; /* arm parameters: local values */
    long Gx1,Gy1,Gx2,Gy2,oGx,oGy;
    float Rot0,DirG; /* current rotation, direction of G-vector */
    long Xp,Yp; /* commands for sffl */
    float circle0,circle1; /* current extension, new extension */
    float RotInc,ExtInc; /* rotation & extension increments */

    toolboxinit(); /*set up toolbox & get new GrafPort */
    SetRect(&ctrlrect,0,36,270,306); /* control area - 2/3 of mpile's */
    SetRect(&armrect,271,0,512,342); /* prepare arm's display region */
    while(1){
```

```

p=0; /* screen & data initializations */
BackPat(black);
EraseRect(GrRPtr);
BackPat(white);
EraseRect(&ctrlrect);
EraseRect(&armrect);
PenPat(gray);
MoveTo(135,36); LineTo(135,305); MoveTo(0,171); LineTo(269,171);
PenPat(black); SetRect(&ex0,133,169,138,174); FrameOval(&ex0);
PenSize(3,3); MoveTo(168,229); LineTo(168,229);
PenNormal(); MoveTo(169,230); LineTo(135,171);
for(ii=0; ii<12; ii+=4){
    MoveTo(scrIP[ii+3][0],scrIP[ii+3][1]);
    for(jj=0; jj<4; ++jj)
        LineTo(scrIP[ii+jj][0],scrIP[ii+jj][1]);
    circle0=PIby2; /* initial 'circle' */
    Rot0=-PIby3; /* initial rotation position */
    oGx=0; oGy=0;
    armparms[0]=NoMOMENTS; /* temporary default */
    for(ii=1; ii<15; ++ii) armparms[ii]=0.0; /* initializing */
    armparms[7]=PIby2; /* adjustment on forearm rotation */
    armparms[9]=Rot0; /* adjustment on upper arm rotation */
    fetch(parms,&ctrlrect); /* get some parameters */
    if(*parms==1) {BackPat(white); EraseRect(GrRPtr); break;} /* QUIT */
    if(*parms==2) continue; /* RESTART */
    ++*parms; /* in case of zero - for the while loop */
    while(--*parms){
        /* degrouping loop */
        Gx1=parms[++p]; /* transferring values */
        Gyl=parms[++p];
        Gx2=parms[++p];
        Gy2=parms[++p];
        DirG=atan((double)(Gy2-Gyl)/(Gx2-Gx1));
        Xp=Gx1-oGx; Yp=Gyl-oGy;
        circle1=2.0*parms[++p]*RAD;
        ExtInc=circle0-circle1;
        RotInc=DirG-Rot0;
        oGx=Gx1; oGy=Gyl; circle0=circle1; Rot0=DirG; /* update */
        armparms[8]=ExtInc; /* initializations for arm() */
        armparms[10]=RotInc; /* already in radians */
        armparms[12]=Xp*RAD;
        armparms[14]=Yp*RAD;
        arm(parms,armparms);
    }
    /* end of degrouping loop */
}
ClosePort(thePort);
DisposPtr(thePort);
}

fetch(par,rp)
long * par; /* parameter list: first byte is number of triplets */
Rect * rp; /* pointer to control area rectangle */
{
    int i,j,finished,evx,evy,mevx,mevy;
    char * cd[2]; /* holds pointers to following 2 strings */
    char * okstring="\Pdegroup";
    char * rtstring="\Pcommand";
    char * rsstring="\Prestart";

```

```

char * qtstring="\Pquit";
long x,y,tx,ty,clr;
EventRecord ctrl;
Point mous; /* mouse position */
Rect crect; /* rectangle in which to draw selected circle */
Rect okr,rst,qut; /* command & degroup, restart and quit rects */
CursHandle curshdl;

/* initializations */
cd[0]=okstring;
cd[1]=rtstring;
FlushEvents(everyEvent,0);
curshdl=GetCursor(crossCursor); /* handle to cross cursor */
messrect(&rst,189,5,266,31,rsstring); /* restart button */
messrect(&qut,5,5,68,31,qtstring); /* quit button */
messrect(&okr,5,311,68,337,cd[0]);
finished=0;
j=0;
while(!finished) /* passes maximum 20 commands */
  for(i=0; i<2; ++i) { /* get two parameters */
    messrect(&okr,5,311,68,337,cd[i]);
    while(!GetNextEvent(mDownMask,&ctrl)) { /* wait for event */
      GetMouse(&mous);
      if(PtInRect(pass(mous),rp)) SetCursor(*curshdl);
      else InitCursor();
      if(PtInRect(pass(ctrl.where),&okr) && !i) {
        BackPat(black); EraseRoundRect(&okr,12,12); BackPat(white);
        finished=1; break; /* flag the while & quit for-loop */
      }
      else if(PtInRect(pass(ctrl.where),&qut)) { finished=-1; break; }
      else if(PtInRect(pass(ctrl.where),&rst)) { finished=-2; break; }
      else if(PtInRect(pass(ctrl.where),rp)) {
        evx=ctrl.where.h;
        evy=ctrl.where.v;
        x=evx-135; /* transform coordinates */
        y=171-evy;
        if(i==0) { SysBeep(2); /* signal OK */
          mevx=evx; mevy=evy; /* remember */
          SetRect(&crect, evx-2, evy-2, evx+2, evy+2);
          FrameOval(&crect);
          ++par=x*0.666666; ++par=y*0.666666; /* first parameter */
        }
        else { tx=x*(par-1)/0.666666;
          ty=y-par/0.666666;
          if(((cir=0.666666*sqrt((double)(tx*tx+ty*ty)))>=
              90 || evx<mevx)
            { i=0; continue; } /* get another value */
          else { SysBeep(2); /* signal OK */
            MoveTo(evx-1, evy-1);
            PenSize(3,3); LineTo(evx-1, evy-1); PenNormal();
            MoveTo(evx, evy);
            LineTo(mevx, mevy);
            ++par=x*0.666666; ++par=y*0.666666;
            ++par=cir; /* second parameter */
            if(++j>=20) { InitCursor(); finished=1; break; }
            /* quintuplet counter */
          }
        }
      }
    }
    else i=-1;
  }
if(finished!=1) *(par-j*5-i*3)=finished; /* flag central loop */

```

```

else *(par-j*5)=j; /* store length */
}

messrect(rtp,l,t,r,b,mss)
Rect * rtp;
int l,t,r,b;
char * mss;
{
SetRect(rtp,l,t,r,b);
EraseRoundRect(rtp,12,12);
MoveTo(rtp->left+(((unsigned)rtp->right-rtp->left-6*strlen(mss))>>1),rtp->top+18);
DrawString(mss);
}

arm(prmtrs,thphd)
long * prmtrs;
float * thphd;
{
float tr[3][3]; /* composite rotation matrix */
float segM[12][3]; /* moved arm - segment positions in 3-D */
EventRecord stev;
Point pt; /* pen location */
Rect lrect,rstrect,qstrect; /* looping & stop, restart and quit rectangles */
char jstr[4]; /* ascii-converted j */
int i,j,m,n,pair,r,num;
float z;
register int (*PM)[2]; /* heaped projection coordinates */
char *PMp; /* new pointer */
char * lpstring="\Ploop. = ";
char * ctstring="\Pcontinue";
long endt,offset=12;

num=thphd[0];
SetRect(&rstrect,189,5,266,31);
SetRect(&qstrect,5,5,68,31);
BackPat(black); EraseRoundRect(&qstrect,12,12); BackPat(white);
for(i=2; i<=14; i+=2) thphd[i]=thphd[i]/num; /* step sizes */
if(!(PM=PMp=NewPtr((long)48*(num+1)))) /* space for positions */
{SysBeep(5); exit();}
jstr[0]=2; /* string j in pascal format */
messrect(&lrect,189,311,266,337,lpstring); /* message area */
GetPen(&pt); /* current pen location */
pt.h-=14; /* backup to print 2 numbers */
TextMode(srcXor); /* trace-erase loop counter */

for(j=0; j<=num; ++j){ /* head of positions generation loop */
FlushEvents(everyEvent,0);

/* TRANSFORMATION */

ftoa((double)j,jstr+1,0,1); /* convert j to ascii */
MoveTo(pt.h,pt.v); /* relocate the pen */
DrawString(jstr); /* print it */
for(m=0; m<12; ++m) for(n=0; n<3; ++n) segM[m][n]=segMm[m][n];
/* restart at position 0 */

if(j) /* compute IP for display */

```

```

(thphd[7]--PIby2; /* check for limit of motion in d.f.s */
thphd[9]++PIby3;
for(i=1; i<=13; i+=2){
    thphd[i]=thphd[i+1];
    if(thphd[i]>PIby2) thphd[i]=PIby2-0.0174533;
    if(thphd[i]<--PIby2) thphd[i]=--PIby2+0.0174533;
thphd[7]++PIby2;
thphd[9]--PIby3;
    }
    /* transformation calls */
    transf(thphd[1],thphd[3],k3,13,tr); mult(3,1,tr,segM); /* hand */
    transf(thphd[5],thphd[7],k4,14,tr); mult(7,1,tr,segM); /* forearm rot & ext */
    transf(thphd[9],0,0,k4,14,tr); mult(11,0,tr,segM); /* upper arm rot */
    transf(thphd[11],thphd[13],k3,13,tr); mult(11,0,tr,segM); /* upper arm */

MoveTo(pt.h,pt.v); /* erase procedure */
if(j!=num) DrawString(jstr);

/* PROJECTION & STORAGE */

for(m=0; m<12; ++m){
    z--segM[m][2]+PLANE+PDIST;
    PM[m][0]=segM[m][0]*PLANE/z+HOR;
    PM[m][1]=--segM[m][1]*PLANE/z+VER;
    PM+=offset; /* point to next coordinates block */
    } /* end of positions generation loop */

/* OUTPUT */

messrect(&lrect,189,311,266,337,ctstring); /* stop button */
offset=12*num;
FlushEvents(everyEvent,0);
while(1){
    PM=PMp; /* reset PM to point to beginning */
    if(GetNextEvent(mDownMask,&stev))
        if(PtInRect(pass(stev.where),&lrect)){
            if(*prmtrs==1) i=1;
            else i=-1;
            BackPat(black); EraseRoundRect(&lrect,12,12);
            BackPat(white); break;}
        else if(PtInRect(pass(stev.where),&rstrect))
            (*prmtrs=1; i=-1; break;)
    for(i=0; i<=offset; i+=12){
        if(i) PenMode(patXor);
        if(i==offset && *prmtrs==1) PenMode(patCopy);
        for(r=0; r<2; ++r){
            /* ON-OFF except for first position and last of last */
            for(m=0; m<12; m+=4){
                MoveTo(PM[~+3][0],PM[m+3][1]);
                for(pair=0; pair<4; ++pair)
                    LineTo(PM[m+pair][0],PM[m+pair][1]);
                if(!r) Delay((long)4,&endt);} /* do not erase first position */
            PenMode(patCopy);
            MoveTo(PM[2][0],PM[2][1]); /* retrace hand extremity */
            LineTo(PM[3][0],PM[3][1]);
            PM+=12;}
    }
}

```

```

/* END */

if(i!=-1)
    while(1){ /* wait for restart signal */
        while(!GetNextEvent(mDownMask,&steve));
        if(PtInRect(pass(steve.where),&rstrect)) break;}
        PenMode(patCopy);
        DisposPtr(PMp); /* free the block */
    }

transf(th,ph,k,l,Tr)
float th,ph;
int *k,*l;
float Tr[3][3];
{
    float thcos,thsin,phcos,phsin,g;
    thcos=cos((double)th); thsin=sin((double)th);
    phcos=cos((double)ph); phsin=sin((double)ph);
    if(!*k){
        g=atan((double)thcos*phsin/phcos); /* compute gamma */
        phcos=cos((double)g); /* & recompute values */
        phsin=sin((double)g);}
    Tr[k[0]][l[0]]=thcos*phcos;
    Tr[0][l[1]]=phsin;
    Tr[k[2]][l[2]]=-phcos*thsin;
    Tr[l[1]][0]=-phsin*thcos;
    Tr[k[4]][l[4]]=phcos;
    Tr[k[5]][l[5]]=thsin*phsin;
    Tr[k[6]][l[6]]=thsin;
    Tr[k[7]][l[7]]=0.0;
    Tr[2][2]=thcos;
}

mult(id,shft,mx,sgM)
int id,shft;
float mx[3][3],sgM[12][3];
{
    int m,n;
    float row[3],adj; /* row product of (sgM=)segM[ ] with (mx=)tr[ ] */
    for(m=0; m<=id; ++m){
        for(n=0; n<3; ++n){
            adj=(!n && shft)*JP;
            row[n]=sgM[m][0]*mx[0][n]+sgM[m][1]*mx[1][n]+sgM[m][2]*mx[2][n]+adj;}
        for(n=0; n<3; ++n) sgM[m][n]=row[n]; /* copy row back for next pass */
    }
}

toolbxinit()
{
    /* cf II:27 */
    GrafPtr gp;

    InitGraf(&thePort);
    InitFonts();
    InitWindows();
    InitMenus();
    TEInit();
    InitDialogs(0L);
}

```

```
InitCursor();  
FlushEvents(everyEvent,0);  
gp=NewPtr((long)sizeof(GrafPort)); /* room for new GrafPort */  
OpenPort(gp);  
TextFont(4); /* Monaco font */ }
```

References

- Adams, J.A. (1971) A Closed-Loop Theory of Motor Learning. *Journal of Motor Behavior*, 3, 111-149.
- Albus, J.S. (1975) A New Approach to Manipulator Control: The Cerebellar Model Articulation Controller (CMAC). *Journal of Dynamic Systems, Measurement, and Control*, 97, 220-227.
- Arbib, M.A. (1961) Turing Machines, Finite Automata and Neural Nets. *Journal of the Association of Computing Machinery*, 8, 467-475.
- Asanuma, H. (1981) Functional Role of Sensory Inputs to the Motor Cortex. *Progress in Neurobiology*, 16, 241-262.
- Bartlett, F.C. (1932) *Remembering*. Cambridge University Press.
- Brady, M. (1982) Trajectory Planning. In Brady et al. (1982).
- Brady, M., Hollerbach, J.M., Johnson, T.L., Lozano-Perez, T., Mason, M.T. (eds.) (1982) *Robot Motion: Planning and Control*. MIT Press.
- Buchthal, F., Schmalbruch, H. (1980) Motor Unit of Mammalian Muscle. *Physiological Reviews*, 60, 90-142.
- Burke, D. (1981) The Activity of Human Muscle Spindle Endings in Normal Motor Behavior. *International Review of Physiology*, 25, 91-126.
- Cajal, S.R. (1894) La fine structure des centres nerveux. *Proceedings of the Royal Society of London*, 55, 444-468.
- Cajal, S.R. (1911) *Histologie du système nerveux de l'homme et des vertébrés*. Volume 2, Paris, Maloine.
- Craik, K.J.W. (1943) *The Nature of Explanation*. Cambridge University Press.
- Craik, K.J.W. (1947) Theory of the Human Operator in Control Systems. I. The operator as an engineering system. *British Journal of Psychology*, 38(2), 56-61.
- Craik, K.J.W. (1948) Theory of the Human Operator in Control Systems. II. Man as an element in a control system. *British Journal of Psychology*, 38(3), 142-148.
- Craik, K.J.W. (1966) *The Nature of Psychology*. Edited by S.L. Sherwood, Cambridge University Press.
- DeLong, M.R. (1974) Motor Functions of the Basal Ganglia: Single-Unit Activity During Movement. In F.O. Schmitt, F.G. Worden (eds.) *The Neurosciences, Fourth Study Program*. MIT Press.
- Desmedt, J.E. (ed.) (1978) *Cerebral Motor Control in Man: Long Loop Mechanisms*. *Progress in Clinical Neurophysiology*, 4.

- Dodge, R., Bott, E.A. (1927) Antagonistic Muscle Action in Voluntary Flexion and Extension. *The Psychological Review*, 34, 241-272.
- Evarts, E.V. (1980) Brain Mechanisms in Voluntary Movement. In D. McFadden (ed.) *Neural Mechanisms in Behavior*. Springer-Verlag.
- Fitts, P.M. (1954) The Information Capacity of the Human Motor System in Controlling the Amplitude of Movement. *Journal of Experimental Psychology*, 47, 381-391.
- Fitts, P.M. (1964) Perceptual-Motor Skill Learning. In A.W. Melton (ed.) *Categories of Human Learning*. Academic Press.
- Flatau, C.R. (1973) Design Outline for Mini-Arms Based on Manipulator Technology. MIT Artificial Intelligence Laboratory, AIM-300, 45 pp.
- Foley, J.D., Van Dam A. (1982) *Fundamentals of Interactive Computer Graphics*. Addison-Wesley. (reprinted with corrections March 1983)
- Fritsch, G., Hitzig, E. (1870) Über die elektrische Erregbarkeit des Grosshirns. *Archiv f. Anat. u. Physiol. Wissenschaftliche Medizin*, 37, 300-332. (Excerpt translated by D. Harris as 'On the Electrical Excitability of the Cerebrum' in K.H. Pribram (ed.) *Brain and Behaviour 2 Perception and Action*, Selected Readings. Penguin Books 1969.)
- Hallet, M., Shahani, B.T., Young, R.R. (1975) EMG Analysis of Stereotyped Voluntary Movements in Man. *Journal of Neurology, Neurosurgery, and Psychiatry*, 38, 1154-1162.
- Hammond, H., Merton, P.A., Sutton, G.G. (1956) Nervous Gradation of Muscular Contraction. *British Medical Bulletin*, 12, 214-218.
- Harvey, R.J. (1980) Cerebellar Regulation in Movement Control. *Trends in NeuroSciences*, 3, 281-284.
- Hebb, D.O. (1949) *The Organization of Behavior*. New York, John Wiley & Sons.
- Henneman, E. (1981) Recruitment of Motoneurons: the Size Principle. In J.E. Desmedt (ed.) *Motor Unit Types, Recruitment and Plasticity in Health and Disease. Progress in Clinical Neurophysiology*, 9, 26-60.
- Hildreth, E.C., Hollerbach, J.M. (1985) The Computational Approach to Vision and Motor Control. MIT Artificial Intelligence Laboratory, AIM-846, 84pp.
- Hillis, W.D. (1985) *The Connection Machine*. MIT Press.
- Hollerbach, J.M. (1980) An Iterative Lagrangian Formulation of Manipulator Dynamics. MIT Artificial Intelligence Laboratory, AIM-533, 22 pp.
- Hollerbach, J.M. (1982) Computers, Brains, and the Control of Movement. *Trends in NeuroSciences*, 5, 189-192.
- von Holst, E. (1954) Relations Between the Central Nervous System and the Peripheral Organs. *British Journal of Animal Behaviour*, 2, 89-94.

- von Holst, E., Mittelstaedt, H. (1950) Das Reafferenzprinzip. Wechselwirkungen zwischen Zentralf-nervensystem und Peripherie. *Naturwissenschaften*, 37, 464-476. (Translated by R.D. Martin, ed., *The Behavioural Physiology of Animals and Man*. Collected Papers, vol. 1, University of Miami Press.)
- Hubel, D.H., Wiesel, T.N. (1962) Receptive Fields, Binocular Interaction, and Functional Architecture in the Cat's Visual Cortex. *Journal of Physiology*, 160, 106-154.
- Hubel, D.H., Wiesel, T.N. (1979) Brain Mechanisms of Vision. *Scientific American*, 241(3), 150-162.
- Jastrow, J. (1892) On the Judgment of Angles and Positions of Lines. *American Journal of Psychology*, 5, 214-223.
- Johnson, T.L. (1982) Feedback Control. In Brady et al. (1982).
- Jones, B. (1978) The Role of Efference in Motor Control: A Centralist Emphasis for Theories of Skilled Performance. In D.M. Landers, R.W. Christina (eds.) *Psychology of Motor Behavior and Sport*. Champaign, Illinois, Human Kinetics Publishers.
- Keele, S.W. (1968) Movement Control in Skilled Motor Performance. *Psychological Bulletin*, 70, 387-403.
- Kelso, J.A.S., Stelmach, G.E. (1976) Central and Peripheral Mechanisms in Motor Control. In G.E. Stelmach (ed.) *Motor Control: Issues and Trends*. Academic Press.
- Kleene, S.C. (1956) Representation of Events in Nerve Nets and Finite Automata. *Annals of Mathematical Studies*, 34, 3-41.
- Köhler, W. (1947) *Gestalt Psychology*. New American Library.
- Lamontagne, C. (1976) *Steps Towards a Computational Theory of Visual Motion Detection: Designing a Working System*. Doctoral dissertation, University of Edinburgh, Edinburgh, U.K.
- Lamontagne, C. (1986) Sensorimotor Emergence: Proposing a Computational Syntax. In W. Callebaut, R. Pinxten (eds.) *Evolutionary Epistemology: A Multiparadigm Program*. Reidel.
- Lamontagne, C., Beausoleil, J.R. (1982) Achieving Visual Spatiality: Towards a Psychologically Relevant, Physiologically Plausible, and Computationally Efficient Conjecture. *Cognition and Brain Theory*, 5, 343-365.
- Lamontagne, C., Beausoleil, J.R. (1984) Vers une représentation conceptuelle de la représentation préconceptuelle. *Communication Information*, 6(2/3), 177-201.
- Lamontagne, C., Howe, J.A.M. (1980) *Towards a Computational Theory of Visual Motion Perception*. Ghent, Communication & Cognition.
- Lashley, K.S. (1917) The Accuracy of Movement in the Absence of Excitation from the Moving Organ. *American Journal of Physiology*, 43, 169-194.
- Lashley, K.S. (1951) The Problem of Serial Order in Behavior. In L.A. Jeffress (ed.) *Cerebral Mechanisms in Behavior*. New York, Wiley.

- Lozano-Perez, T. (1982) Robotics. *Artificial Intelligence*, 19, 137-143.
- Marsden, C.D. (1980) The Enigma of the Basal Ganglia and Movement. *Trends in NeuroSciences*, 3, 284-287.
- Mason, M.T. (1979) Compliance and Force Control for Computer Controlled Manipulators. MIT Artificial Intelligence Laboratory, AI-TR-515, 71 pp.
- Matthews, P.B.C. (1972) *Mammalian Muscle Receptors and their Central Actions*. London, Arnold.
- Max, L.W. (1934) An Experimental Study of the Motor Theory of Consciousness: I. Critique of Earlier Studies. *Journal of General Psychology*, 11, 112-125.
- McCulloch, W.S., Pitts, W.H. (1943) A Logical Calculus of the Ideas Immanent in Nervous Activity. *Bulletin of Mathematical Biophysics*, 5, 115-133.
- Merton, P.A. (1953) Speculations on the Servo-Control of Movement. In J.L. Malcolm, J.A.B. Gray, G.E.W. Wolstenholme, J.S. Freeman (eds.) *The Spinal Cord*. London, J & A Churchill.
- Miles, F.A., Evarts, E.V. (1979) Concepts of Motor Organization. *Annual Review of Psychology*, 30, 327-362.
- von Neumann, J. (1951) The General and Logical Theory of Automata. In L.A. Jeffress (ed.) *Cerebral Mechanisms in Behavior*. New York, Wiley.
- Paillard, J. (1960) The Patterning of Skilled Movements. *Handbook of Physiology*, (Part 1. Neurophysiology) 3, 1679-1708.
- Paillard, J. (1978) The Pyramidal Tract: Two Million Fibres in Search of a Function. *Journal de Physiologie*, 74, 155-162.
- Paillard, J. (1980) Nouveaux objectifs pour l'étude neurobiologique de la performance motrice intégrée: les niveaux de contrôle. In C.H. Nadeau, W.R. Halliwell, K.M. Narell, G.C. Roberts (eds.) *Psychology of Motor Behavior and Sport*. Champaign, Illinois, Human Kinetics Publishers.
- Paul, R.P. (1981) *Robot Manipulators: Mathematics, Programming, and Control*. MIT Press.
- Pellionisz, A. (1980) Tensorial Approach to the Geometry of Brain Function: Cerebellar Coordination via a Metric Tensor. *Neuroscience*, 5, 1125-1136.
- Penfield, W. (1954) Mechanisms of Voluntary Movement. *Brain*, 77, 1-17.
- Pew, R.W. (1974) Human Perceptual-Motor Performance. In B.H. Kantowitz (ed.) *Human Information-Processing: Tutorials in performance and cognition*. New York, John Wiley & Sons.
- Phillips, C.G. (1969) Motor Apparatus of the Baboon's Hand. *Proceedings of the Royal Society, B* 173, 141-174.
- Popper, K.R. (1963) *Conjectures and Refutations*. London, Routledge & Kegan Paul.

- Popper, K.R. (1972) *Objective Knowledge*. Oxford, Clarendon Press.
- Prochazka, A. (1981) Muscle Spindle Function During Normal Movement. *International Review of Physiology*, 25, 47-90.
- Proske, U. (1981) The Golgi Tendon Organ. *International Review of Physiology*, 25, 127-171.
- Raibert, M.H. (1978) A Model for Sensorimotor Control and Learning. *Biological Cybernetics*, 29, 29-36.
- Roberts, T.D.M. (1978) *Neurophysiology of Postural Mechanisms*. (2nd edition) London, Butterworth & Co.
- Rooney, J. (1984) Kinematic and Geometric Structure in Robot Systems. In T. O'Shea, M. Eisenstadt (eds.) *Artificial Intelligence: Tools, Techniques, and Applications*. New York, Harper & Row.
- Schmidt, R.A. (1975) A Schema Theory of Discrete Motor Skill Learning. *Psychological Review*, 82, 225-260.
- Schmidt, R.A. (1980) On the Theoretical Status of Time in Motor Program Representations. In G.E. Stelmach & J. Requin (eds.) *Tutorials in Motor Behavior*. North-Holland Publishing Company.
- Schmidt, R.A., Zelaznik, H., Hawkins, B., Frank, J.S., Quinn, J.T. Jr. (1979) Motor-Output Variability: A Theory for the Accuracy of Rapid Motor Acts. *Psychological Review*, 86, 415-451.
- Schwartz, E.L. (1980) Computational Anatomy and Functional Architecture of Striate Cortex: A Spatial Mapping Approach to Perceptual Coding. *Vision Research*, 20, 645-669.
- Selfridge, O.G. (1959) Pandemonium: a Paradigm for Learning. In *Mechanisation of Thought Processes*. NPL Symposium no. 10, London, HMSO.
- Sherrington, C.S. (1905) *The Integrative Action of the Nervous System*. Yale University Press.
- Sherrington, C.S. (1925) Remarks on some Aspects of Reflex Inhibition. *Proceedings of the Royal Society, B*, 97, 519-545.
- Skoglund, S. (1973) Joint Receptors and Kinaesthesia. *Handbook of Sensory Physiology II: Somatosensory System*. (A. Iggo, ed.) Springer-Verlag.
- Sperry, R.W. (1950) Neural Basis of the Spontaneous Optokinetic Response Produced by Visual Inversion. *Journal of Comparative and Physiological Psychology*, 43, 482-489.
- Szentágothai, J. (1979) Local Neuron Circuits of the Neocortex. In F.O. Schmitt, F.G. Worden (eds.) *The Neurosciences, Fourth Study Program*. MIT Press.
- Thach, W.T. (1978) Correlation of Neural Discharge with Pattern and Force of Muscular Activity, Joint Position, and Direction of Intended Next Movement in Motor Cortex and Cerebellum. *Journal of Neurophysiology*, 41, 654-676.

- Turing, A.M. (1936) On Computable Numbers, with an Application to the Entscheidungsproblem. *Proceedings of the London Mathematical Society*, Series 2, 42, 230-265 & 43, 544-546.
- Turing, A.M. (1950) Computing Machinery and Intelligence. *Mind*, 59, 433-460.
- Waters, R.C. (1979) Mechanical Arm Control. MIT Artificial Intelligence Laboratory, AIM-549, 20 pp.
- Wertheimer, M. (1923) Untersuchungen zur Lehre von der Gestalt, II. *Psychologische Forschung*, 4, 301-350.
- Whitney, D.E. (1972) The Mathematics of Coordinated Control of Prosthetic Arms and Manipulators. *Journal of Dynamic Systems, Measurement, and Control*, 94, 303-309.
- Wilson, D.J. (1933) Antagonistic Muscle Action During the Initiatory Stages of Voluntary Effort. *Archives of Psychology*, No. 160, 48pp.
- Wise, S.P., Evarts, E.V. (1981) The Role of the Cerebral Cortex in Movement. *Trends in Neurosciences*, 4, 297-300.
- Woodworth, R.S. (1899) The Accuracy of Voluntary Movement. *Psychological Review Monograph Supplements*, 3(3), Whole No. 13, 114pp.

Note on the text

The foregoing thesis was produced on an Apple® Macintosh™ microcomputer using MacWrite 4.5; the figures were prepared with MacPaint 1.5, and MacDraw 1.9; final copy was printed on an Apple LaserWriter™. Typefaces are Times, Courier, & Palatino.

◊ ◊ ◊