

2 Neural Dynamics of Category Learning and Recognition: Structural Invariants, Reinforcement, and Evoked Potentials

Gail A. Carpenter

Northeastern University and Boston University

Stephen Grossberg

Center for Adaptive Systems Boston University

LEARNING OF COGNITIVE RECOGNITION CODES

This chapter describes how cognitive recognition codes can be learned in response to a temporal stream of input patterns. This self-organizing learning process automatically buffers, or self-stabilizes, its learning against recoding by the “blooming buzzing confusion” of irrelevant experience, no matter how many input patterns are processed, and no matter how complex these input patterns may be, until it utilizes its full memory capacity. These results are part of an extensive development of Grossberg’s *adaptive resonance theory*, or ART (Carpenter & Grossberg, 1987a,b,c, 1988). The ART model discussed herein is called ART 1 (Carpenter & Grossberg, 1987a, 1988). It was developed to explain how a world of binary input patterns can be learned and recognized. A more recent model, called ART 2, was developed to deal with an analog or binary input world (Carpenter & Grossberg, 1987b,c). An account of the historical factors leading to the introduction of adaptive resonance theory and its relationship to other neural network models, such as competitive learning, interactive activation, and back propagation, is found in Grossberg (1987a).

Recognition codes are discovered and learned by a real-time neural network that is sensitive both to the global structure of input patterns and to their local features. In particular, the network’s long-term memory (LTM) traces learn *critical feature patterns*, rather than individual features. These critical feature patterns are sensitive to global structural properties of the inputs but are not governed by “structural rules.” They capture structural properties such as those that Mumford (this volume) has discussed, but they are not determined a priori

of the network. The network actively regulates its learning process by matching its top-down templates against bottom-up input patterns and using the degree of match or mismatch to determine processes of learning or memory search, respectively. Thus the present theory analyses the fundamental role of novelty-sensitive, or matching, processes in the regulation of learning.

Using Context to Distinguish Noise from Features

The matching process takes place at level F_1 in Fig. 2.1a, where it compares whole activity patterns across a field of feature-selective cells, rather than activations of individual feature detectors. The constraints on this matching process that enable it to stabilize network learning also automatically rescale the matching criterion in response to input and template patterns of variable complexity. Thus the network can both differentiate finer details of simple input patterns and tolerate larger mismatches of complex input patterns. This rescaling property also defines the difference between irrelevant features and significant pattern mismatches in the following way.

If a mismatch within the attentional subsystem does not activate the orienting subsystem (viz., an "approximate match" occurs), then a search for a different code is not initiated. Consequently, mismatched features are treated as noise in the sense that they are rapidly suppressed in short-term memory (STM) at F_1 and are eliminated from the critical feature pattern learned by the $F_2 \rightarrow F_1$ template that is active at that time. If the mismatch is great enough to generate a search, then the mismatched features may be included in the critical feature pattern of the category to which the search leads. Due to the self-scaling nature of the matching process, if a template mismatches a simple input pattern at just a few features, a search may be elicited, thereby enabling the network to learn fine discriminations among patterns composed of few features. On the other hand, if a template mismatches the same number of features within a complex input pattern, a search may not be elicited; then the mismatched features will be suppressed as noise.

Thus the pattern-matching process of the model automatically exhibits properties that are akin to attentional focusing, or "zooming in." No separate attentional "zooming" mechanism is needed to generate these properties. On the other hand, attentional processing is of critical importance in regulating the network's ability to learn a recognition code in a stable way.

Refining Recognition Codes: Influence of Reinforcement on Vigilance

The present theory distinguishes three types of attentional mechanism; attentional priming, gain control, and vigilance. A single network parameter, the attentional vigilance parameter, determines how fine the learned categories will be. This vigilance mechanism allows the network to adjust for the fact that, no

matter how well feature detectors are made, the environment may impose variable demands on recognition performance. The present system is capable of refining its discriminative capacity, without any change of the feature detectors within the attentional subsystem, simply by changing its vigilance level. Such a vigilance change is distinct from the automatic-scaling property of the network, which holds at each fixed vigilance level.

Suppose that disconfirming environmental feedback, or predictive failure, causes an increase in the vigilance parameter. This change increases the sensitivity of the orienting subsystem to mismatches within the attentional subsystem (Fig. 2.1a). As a result, the network automatically searches for and learns finer categorical distinctions.

Disconfirming environmental feedback of this type constitutes a particular type of "teaching" function. Such a teaching function acts *only* to make the system more sensitive to its environment. The architecture of the network as a whole, through its interactions with a unique input environment, translates this sensitivity change into a refinement of the categorical invariants of the learned code.

A punishing event is one of the types of disconfirming environmental feedback capable of causing an increase in vigilance. A punishing signal is assumed to have this effect via a network subsystem that is called a drive representation, as in Fig. 2.1a.

Multiple Effects of Reinforcement: Transfer and Attention Shifts

A network of drive representations models part of the midbrain neural circuitry within which nonhomeostatic processes, such as reinforcement, interact with homeostatic processes, such as appetitive drives and circadian pacemaker signals. In addition to the reinforcement-contingent vigilance signals from a drive representation to the orienting subsystem, the drive representations also interact with the attentional subsystem. These interactions illustrate how a classification rule can be transferred to unfamiliar exemplars of a category (Herrnstein, this volume). In particular, conditionable pathways exist from the unitized representations of the attentional system—in the present example, level F_2 —to the several drive representations (Fig. 2.1b). When a category in F_2 is consistently paired with activation of a drive representation by a reinforcer, the LTM traces in the pathways from the category to the drive representation can grow. Such a pairing can, for example, occur when a conditioned stimulus (CS) is paired with an unconditioned stimulus (US) during classical conditioning. The CS activates a recognition category in F_2 . The US activates a drive representation; for example, when a shock US activates its drive representation, a fear reaction is elicited. The LTM changes in the pathways from F_2 to the active drive representation endow the CS with properties of a conditioned reinforcer. When the US is a shock, the

CS can then elicit a conditioned emotional response of fear. More complex contingencies between cues and reinforcing events can also alter the LTM traces in these pathways (Grossberg, 1982a,b,c, 1984).

To understand how transfer of a classification rule can occur, consider the special case in which a particular exemplar of an F_2 category v_1 is paired with an active drive representation D_1 . As a result of this pairing, the exemplar becomes a conditioned reinforcer by increasing the LTM trace in the pathway from its category v_1 to the drive representation D_1 . Subsequently, without any additional learning, every exemplar of the category v_1 can activate the same drive representation D_1 . Thus the classification of every exemplar in category v_1 as a conditioned reinforcer to D_1 need not be learned through pairing each exemplar of v_1 with activation of D_1 by a reinforcer. The learning process that controls transfer of the classification rule is neither classical conditioning nor reinforcer dependence. Instead, it is due to the LTM tuning of the adaptive filter $F_1 \rightarrow F_2$, which subserves category learning. The conjoint action of the two types of learning enables the reinforcing properties of one exemplar to transfer automatically to all other exemplars that share its categorical recognition invariants. Similarly, if a single exemplar is paired with a motor representation, then every exemplar of its category can activate that motor representation. Again, category learning subserves the transfer property of the sensory-motor map.

Reinforcing cues can thus have multiple effects on the model's circuitry. Regulation of vigilance level and of the learning process whereby categories become conditioned reinforcers have already been mentioned. In addition, each drive representation gives rise to conditionable feedback pathways to all the categories in F_2 . The LTM traces in such a feedback pathway encode incentive motivational properties of a drive representation (Fig. 2.1b). Activation of a drive representation can hereby subliminally prime a set of F_2 categories that are motivationally compatible with the drive representation. When activation of F_2 by a sensory input occurs conjointly with incentive motivational signals from a drive representation, the choice of F_2 category is influenced by both sensory and motivational factors.

Suppose, for example, that an F_2 category v_1 is chosen in response to a sensory input, and v_1 then activates a drive representation D_1 . Let D_1 deliver larger incentive motivational feedback signals to the F_2 category v_2 than to v_1 . Suppose that v_2 then generates large conditioned reinforcer signals to D_1 . In this situation, activation of v_1 can cause a rapid shift of attention to v_2 . Short-term memory storage of v_2 is then sustained by reverberation within the learned conditioned reinforcer-incentive motivational feedback loop between v_2 and D_1 (Fig. 2.1b). In other words, such an attention shift is caused and maintained by the greater motivational salience of v_2 , despite the larger sensory input to v_1 . In models wherein F_2 can store a spatially distributed representation in STM, such as a masking field (Cohen & Grossberg, 1986, 1987, this volume), interactions between these distributed sensory representations and the drive representations

can account for attentional effects such as blocking, unblocking, latent inhibition, and conditioned inhibition (Grossberg, 1982a,b, 1987b; Grossberg & Levine, 1987; Grossberg & Schmajuk, 1987). In summary, learning of recognition codes, as illustrated by $F_1 \rightarrow F_2$ category learning, can enable transfer of a classification rule to occur. Learning of incentive motivation can shift the set of exemplars that share the transfer property.

Why Categories? Noise Suppression and Search

The previous sections summarize some properties of category learning in the theory but do not comment on why categories form in the first place. The subsequent sections review some properties of category learning that clarify the following assertions.

The tendency to form categories follows from the fact that F_2 transforms the input pattern that it receives from F_1 via the adaptive filter $F_1 \rightarrow F_2$ in a way that enhances its spatial contrast, or compresses it. This contrast-enhanced pattern, not the more broadly distributed input pattern, is stored in STM across F_2 . The case under consideration in which F_2 chooses for STM storage the node that receives the largest input is an extreme case of this contrast-enhancement process.

Why is contrast enhancement necessary? Two related answers can be given. The first answer concerns the ability of F_2 to suppress noise: Level F_2 is designed as a competitive feedback network to enable it to store activity patterns in STM. To store STM patterns without enhancing noise, such a network automatically contrast-enhances its input patterns (Grossberg, 1982a, 1983).

The second answer concerns the ability of the network as a whole to search for appropriate recognition codes. Such a search tends to inhibit previously active codes that have led to predictive failure. If these codes were distributed coextensively with their activating input patterns, shutting them off would prevent the subsequent search from being able to activate any new codes whatsoever (Grossberg, 1982b, 1984).

In summary, the computational requirements that lead to categorical learning are based on considerations of reliable signal processing and the possibility of search. An outline of these processing steps is given next.

Bottom-Up Adaptive Filtering and Contrast-Enhancement in Short-Term Memory

Consider the typical network reactions to a single input pattern I within a temporal stream of input patterns. Each input pattern may be the output pattern of a preprocessing stage. The input pattern I is received at the stage F_1 of an attentional subsystem. Pattern I is transformed into a pattern X of activation across the nodes of F_1 (Fig. 2.2). The transformed pattern X represents a pattern in short-

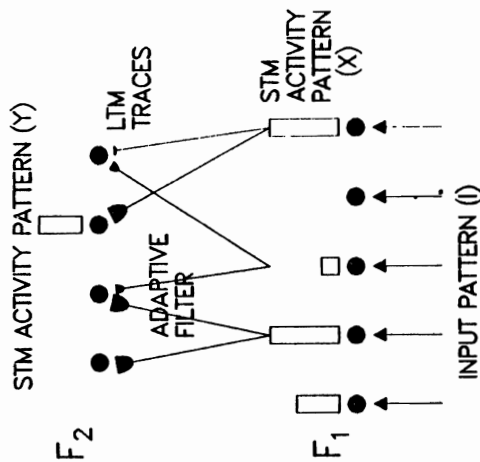


FIG. 2.2. Stages of bottom-up activation: The input pattern I generates a pattern of STM activation X across F_1 . Sufficiently active F_1 nodes emit bottom-up signals to F_2 . This signal pattern S is gated by long-term memory (LTM) traces within the $F_1 \rightarrow F_2$ pathways. The LTM-gated signals are summed before they activate their target nodes in F_2 . This LTM-gated and summed signal pattern T generates a pattern of activation Y across F_2 .

term memory (STM). In F_1 each node whose activity is sufficiently large generates excitatory signals along pathways to target nodes at the next processing stage F_2 . A pattern X of STM activities across F_1 thereby elicits a pattern S of output signals from F_1 . When a signal from a node in F_1 is carried along a pathway to F_2 , the signal is multiplied, or *gated*, by the pathway's long-term memory (LTM) trace. The LTM-gated signal (i.e., signal times LTM trace), not the signal alone, reaches the target node. Each target node sums up all its LTM-gated signals. In this way, pattern S generates a pattern T of LTM-gated and summed input signals to F_2 (Fig. 2.3a). The transformation from S to T is called an *adaptive filter*.

The input pattern T to F_2 is quickly transformed by interactions among the nodes of F_2 . These interactions contrast-enhance the input pattern T . The resulting pattern of activation across F_2 is a new pattern Y . The contrast-enhanced pattern Y , rather than the input pattern T , is stored in STM by F_2 .

A special case of this contrast-enhancement process is one in which F_2 chooses the node that receives the largest input. The chosen node is the only one that can store activity in STM. In general, the contrast-enhancing transformation from T to Y enables more than one node at a time to be active in STM. Such transformations are designed to represent in STM many subsets, or groupings, of an input pattern simultaneously (Cohen & Grossberg, 1986, 1987, this volume; Grossberg, 1986). When F_2 is designed to make a choice in STM, it selects the global grouping of the input pattern that is preferred by the adaptive filter. Feldman (this volume) calls such a choice "winner take all." This process automatically enables the network to partition all the input patterns that are received by F_1 into disjoint sets of recognition categories, each corresponding to a particular node (or "pointer," or "index") in F_2 . This example of a categorical mechanism is both interesting in itself and a necessary prelude to the analysis of recognition codes in which multiple groupings of X are simultaneously represented by Y .

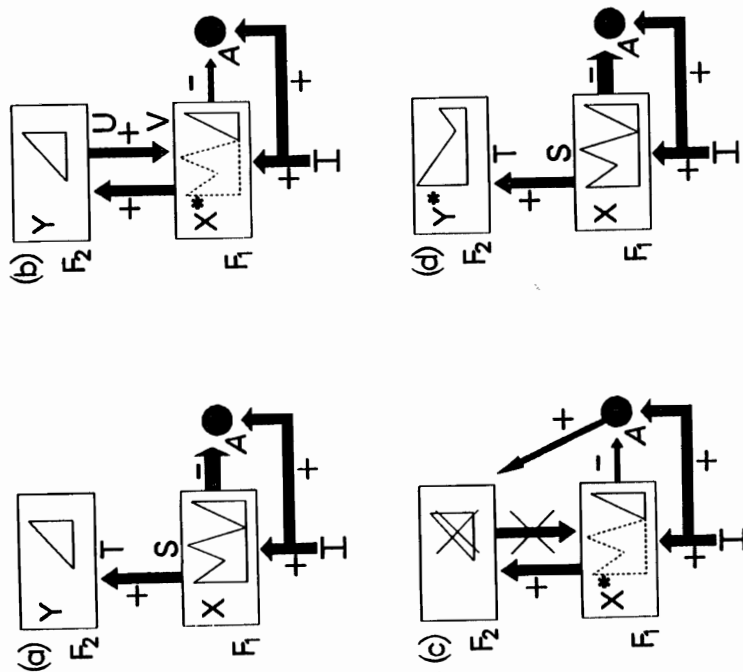


FIG. 2.3. Search for a correct F_2 code: (a) The input pattern I generates the specific STM activity pattern X at F_1 as it nonspecifically activates A . Pattern X both inhibits A and generates the output signal pattern S . Signal pattern S is transformed into the input pattern T , which activates the STM pattern Y across F_2 ; (b) pattern Y generates the top-down signal pattern U that is transformed into the template pattern V . If V mismatches I at F_1 , then a new STM activity pattern X^* is generated at F_1 . The reduction in total STM activity that occurs when X is transformed into X^* causes a decrease in the total inhibition from F_1 to A . (c) Then the input-driven activation of A releases a nonspecific arousal wave to F_2 , which resets the STM pattern Y at F_2 ; (d) after Y is inhibited, its top-down template is eliminated, and X is reinstated at F_1 . Now X once again generates input pattern T to F_2 , but because Y remains inhibited T can activate a different STM pattern Y^* at F_2 . If the top-down template due to Y^* also mismatches I at F_1 , then the rapid search for an appropriate F_2 code continues.

All the LTM traces in the adaptive filter, and thus all learned past experiences of the network, are used to determine the recognition code Y via the transformation $I \rightarrow X \rightarrow S \rightarrow T \rightarrow Y$. However, only those nodes of F_2 that maintain stored activity in the STM pattern Y can elicit new learning. Such learning can occur only within the contiguous LTM traces whose $F_1 \rightarrow F_2$ pathways terminate at the active F_2 nodes. Because the recognition code Y is, in general, a more contrast-

enhanced pattern than T , many F_2 nodes that receive positive inputs ($I \rightarrow X \rightarrow S \rightarrow T$) may not store any STM activity ($T \rightarrow Y$). The LTM traces adjacent to these nodes thus influence the recognition event but are not altered by the recognition event. Similarly, some memories that influence the focus of attention are not themselves attended.

Top-Down Template Matching and Stabilization of Code Learning

Top-down template matching stabilizes the code-learning process. To do so, top-down template matching at F_1 must prevent learning at bottom-up LTM traces whose contiguous F_2 nodes are only momentarily activated in STM. This ability depends on the different rates at which STM activities and LTM traces change. The STM transformation $I \rightarrow X \rightarrow S \rightarrow T \rightarrow Y$ takes place very quickly. "Very quickly" means much more quickly than the rate at which the LTM traces in the adaptive filter $S \rightarrow T$ change. As soon as the bottom-up STM transformation $X \rightarrow Y$ takes place, the STM activities Y in F_2 elicit a top-down excitatory signal pattern U to F_1 . Only sufficiently large STM activities in Y elicit signals in U along the feedback pathways $F_2 \rightarrow F_1$.

As in the bottom-up adaptive filter, the top-down signals U are also gated by LTM traces and the LTM-gated signals are summed at F_1 nodes. The pattern U of output signals from F_2 thereby generates a pattern V of LTM-gated and summed input signals to F_1 . The transformation from U to V is thus also an adaptive filter. The pattern V is called a *top-down template*, or *learned expectation* (Fig. 2.3b).

Two sources of input now perturb F_1 : the bottom-up input pattern I that gave rise to the original activity pattern X , and the top-down template pattern V that resulted from activating X . The activity pattern X^* across F_1 that is induced by I and V taken together is typically different from the activity pattern X that was previously induced by I alone. In particular, F_1 acts to match V against I . The result of this matching process determines the future course of learning and recognition by the network.

The entire activation sequence

$$I \rightarrow X \rightarrow S \rightarrow T \rightarrow Y \rightarrow U \rightarrow V \rightarrow X^* \quad (1)$$

takes place very quickly compared to the rate at which the LTM traces in either the bottom-up adaptive filter $S \rightarrow T$ or the top-down adaptive filter $U \rightarrow V$ can change. Although none of the LTM traces changes during such a short time, their prior learning strongly influences the STM patterns Y and X^* that evolve within the network. The next section contains a discussion of how a match or mismatch of I and V at F_1 regulates the course of learning in response to the pattern I .

INTERACTIONS BETWEEN ATTENTIONAL AND ORIENTING SUBSYSTEMS: STM RESET AND SEARCH

This section outlines how a mismatch at F_1 regulates the learning process. In Fig. 2.3a, an input pattern I instates an STM activity pattern X across F_1 . The input pattern I also excites the orienting population A , but pattern X at F_1 inhibits A before it can generate an output signal. Activity pattern X also generates an output pattern S that, via the bottom-up adaptive filter, instates an STM activity pattern Y across F_2 . In Fig. 2.3b, pattern Y reads a top-down template pattern V into F_1 . Template V mismatches input I , thereby significantly inhibiting STM activity across F_1 . The amount by which activity in X is attenuated to generate X^* depends on how much of the input pattern I is encoded within the template pattern V .

When the mismatch attenuates STM activity across F_1 , this activity no longer prevents the arousal source A from firing. Figure 2.3c depicts how disinhibition of A releases a nonspecific arousal burst to F_2 . This arousal burst, in turn, selectively inhibits the active population in F_2 . This inhibition is longlasting. One physiological design for F_2 processing that has these properties is a *dipole field* (Grossberg, 1980, 1984). A dipole field consists of opponent processing channels that are gated by habituating chemical transmitters. A nonspecific arousal burst induces selective and enduring inhibition within a dipole field. In Fig. 2.3c, inhibition of Y leads to inhibition of the top-down template V and thereby terminates the mismatch between I and V . Input pattern I can thus reinstate the activity pattern X across F_1 , which again generates the output pattern S from F_1 and the input pattern T to F_2 . Due to the enduring inhibition at F_2 , the input pattern T can no longer activate the same pattern Y at F_2 . A new pattern Y^* is thus generated at F_2 by I (Fig. 2.3d). Despite the fact that some F_2 nodes may remain inhibited by the STM reset property, the new pattern Y^* may encode large STM activities. This is because level F_2 is designed so that its total suprathreshold activity remains approximately constant, or normalized, despite the fact that some of its nodes may remain inhibited by the STM reset mechanism. This property is related to the limited capacity of STM. A physiological process capable of achieving the STM normalization property is based on on-center off-surround, or cooperative-competitive, interactions among cells obeying membrane, or shunting, equations (Grossberg, 1980, 1983).

The new activity pattern Y^* reads out a new top-down template pattern V^* . If a mismatch again occurs at F_1 , the orienting subsystem is again engaged, which thereby leads to another arousal-mediated reset of STM at F_2 . In this way, a rapid series of STM matching and reset events may occur. Such an STM matching and reset series controls the system's search of LTM by sequentially engaging the novelty-sensitive orienting subsystem. Such a mismatch-mediated, self-terminating search provides a mechanistic theory of the probabilistic operations described by Heinemann and Chase (this volume). Although STM is reset se-

quantitatively in time, the mechanisms that control the LTM search are all parallel network interactions, rather than serial algorithms. Such a parallel-search scheme is necessary in a system the LTM codes of which do not exist a priori. In general, the spatial configuration of codes in such a system depends on both the system's initial configuration and its unique learning history. Consequently, no prewired serial algorithm could possibly anticipate an efficient order of search.

The mismatch-mediated search of LTM ends when an STM pattern across F_2 reads out a top-down template that matches I , to the degree of accuracy required by the level of attentional vigilance, or that has not yet undergone any prior learning. In the former case, the learned code that is accessed by the search is refined using any additional information that the input I may contain. In the latter case, a new recognition category is established as a bottom-up code and top-down template are learned.

Attentional Gain Control and Attentional Priming

The top-down template-matching process that stabilizes learning is also a mechanism of attentional priming. Consider, for example, a situation in which F_2 is activated by a level other than F_1 before F_1 is itself activated (Fig. 2.4a). In such a situation, F_2 can generate a top-down template V to F_1 . The level F_1 is then primed, or ready, to receive a bottom-up input that may or may not match the active expectancy. Level F_1 can be primed to receive a bottom-up input without necessarily eliciting suprathreshold output signals in response to the priming expectancy.

On the other hand, an input pattern I must be able to generate a suprathreshold activity pattern X even if no top-down expectancy is active across F_1 (Fig. 2.3a and 2.4b). How does F_1 know that it should generate a suprathreshold reaction to a bottom-up input pattern, but not to a top-down input pattern? In both cases, input patterns stimulate F_1 cells. Some auxiliary mechanism must exist to distinguish between bottom-up and top-down inputs. This auxiliary mechanism is called *attentional gain control* to distinguish it from *attentional priming* by the top-down template itself. The attentional priming mechanism delivers *specific* template patterns to F_1 . The attentional gain control mechanism has a *nonspecific* effect on the sensitivity with which F_1 responds to the template pattern, as well as to other patterns received by F_1 . The combination of attentional gain control with patterned input allows F_1 to tell the difference between bottom-up and top-down signals.

Matching: The 2/3 Rule

The 2/3 Rule follows naturally from the distinction between attentional gain control and attentional priming. It states that two out of three signal sources must activate an F_1 node for that node to generate suprathreshold output signals. In

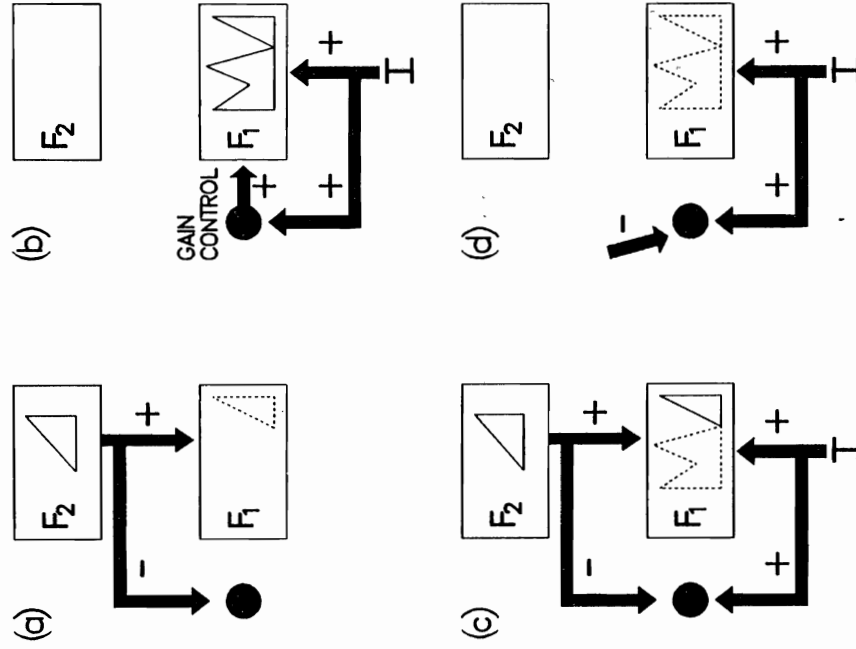


FIG. 2.4. Matching by 2/3 Rule: (a) A top-down template from F_2 inhibits the attentional gain control source as it subliminally primes target F_1 cells; (b) nonspecific attentional gain control signals are activated by the bottom-up input. The bottom-up input arouses two different nonspecific channels: the attentional gain control channel and the orienting subsystem. Only F_1 cells that receive bottom-up inputs and gain control signals become supraliminally active; (c) when a bottom-up input pattern and a top-down template are simultaneously active, only those F_1 cells that receive inputs from both sources become supraliminally active, because the gain control source is inhibited; (d) intermodality inhibition shuts off the gain control source and thereby prevents a bottom-up input from supraliminally activating F_1 .

Fig. 2.4a, during top-down processing, or priming, the nodes of F_1 receive inputs from at most one of their three possible input sources. Hence, no cells in F_1 are supraliminally activated by the top-down template. In Fig. 2.4b during bottom-up processing, a suprathreshold node in F_1 is one that receives a specific

Consequently, the template of v_2 is recoded from pattern C to its superset pattern A.

Code Stability Example

In Fig. 2.5b, the 2/3 Rule does hold due to stronger top-down attentional gain control. Thus, the network experiences a sequence of recodings that ultimately stabilizes. In particular, on trial 4, node v_2 reads out the template C, which mismatches the input pattern A. Here, a search is initiated, as indicated by the numbers beneath the template symbols in row 4. First v_2 's template C mismatches A. Then v_1 's template B mismatches A. Finally A activates the uncommitted node v_3 , which resonates with F_1 as it learns the template A.

In Fig. 2.5b, pattern A is coded by v_1 on trial 1; by v_3 on trials 4 and 6; and by v_4 on trial 9. On all future trials, input pattern A is coded by v_4 . Moreover, all the input patterns A, B, C, and D have learned a stable code by trial 9. Thus the code self-stabilizes by the second run through the input list ABCAD. On trials 11 through 15, and on all future trials, each input pattern chooses a different node ($A \rightarrow v_4$; $B \rightarrow v_1$; $C \rightarrow v_3$; $D \rightarrow v_2$). Each pattern belongs to a separate category because the vigilance parameter was chosen to be large in this example. Moreover, after code learning stabilizes, each input pattern directly activates its node in F_2 without undergoing any additional search. Thus after trial 9, only the "RES" symbol appears under the top-down templates. The patterns shown in any row between 9 and 15 provide a complete description of the learned code. Examples of how a novel exemplar can activate a previously learned category are found on trials 2 and 5 in Fig. 2.5a and 2.5b. On trial 2, for example, pattern B is presented for the first time and directly accesses the category coded by v_1 , which previously learned from pattern A on trial 1. In terminology from artificial intelligence, B activates the same categorical "pointer," or "marker," or "index" as A. In so doing, B does not change the categorical "index," but it may change the categorical template, which determines which input patterns will also be coded by this index on future trials. The category does not change, but its invariants may change.

An example of how presentation of different input patterns can influence the category of a fixed input pattern is found through consideration of trials 1, 4, and 9 in Fig. 2.5b. These are the trials on which pattern A is recoded due to the intervening occurrence of other input patterns. On trial 1, pattern A is coded by v_1 . On trial 4, A is recoded by v_3 because pattern B has also been coded by v_1 and pattern C has been coded by v_2 in the interim. On trial 9, pattern A is recoded by v_4 both because pattern C has been recoded by v_3 and because pattern D has been coded by v_2 in the interim.

Two more conclusions follow from these considerations. First, the code-learning process is one of progressive refinement of distinctions. The distinctions that emerge are the cumulative result of all the input patterns that the network

ever experiences, rather than of some preassigned features. Second, the matching process compares whole patterns, not just separate features. Global measures of pattern match at F_1 —the learned critical feature patterns—determine whether STM reset or resonance will occur. For example, two different templates may overlap an input pattern to F_1 at the same set of feature detectors, yet the network could reset the F_2 node of one template and not reset the F_2 node of the other template. The degree of mismatch of template and input as a whole determines whether recoding will occur. Thus, the learning of categorical invariants resolves two opposing tendencies. As categories grow larger and hence code increasingly global invariants, the templates that define them become smaller and hence base the code on sets of critical feature groupings that drive the competition for STM activity. The simulations illustrate how these two opposing tendencies can be resolved, leading to dynamic equilibrium, or self-stabilization, of recognition categories in response to a prescribed input environment.

The next section describes how a sufficiently large mismatch between an input pattern and a template can lead to STM reset, whereas a sufficiently good match can terminate the search and enable learning to occur.

Vigilance, Orienting, and Reset

A summary is now given of how matching with the attentional subsystem at F_1 determines whether or not the orienting subsystem will be activated, thereby leading to reset of the attentional subsystem at F_2 .

Distinguishing Active Mismatch from Passive Inactivity. A severe mismatch at F_1 activates the orienting subsystem A. In the worst possible case of mismatch, no F_1 node satisfies the 2/3 Rule, and no supraliminal activation of F_1 occurs. In this case, when F_1 becomes totally inactive, the orienting subsystem must surely be engaged.

On the other hand, F_1 may be inactive simply because no inputs whatsoever are being processed. In this case, the orienting subsystem should not be activated. How does the network compute the difference between active mismatch and passive inactivity at F_1 ? This question led Grossberg (1980) to assume that the bottom-up input source activates two parallel channels (Fig. 2.3). The attentional subsystem receives a specific input pattern at F_1 . The orienting subsystem receives convergent inputs at A from all the active input pathways. Thus, the orienting subsystem can be activated only when F_1 is actively processing bottom-up inputs.

Competition between the Attentional and Orienting Subsystems. How, then, is a bottom-up input prevented from resetting its own F_2 code? What mechanism prevents the activation of A by the bottom-up input from always resetting the STM representation at F_2 ? Clearly, inhibitory pathways must exist from F_1 to A

(Fig. 2.3a). When F_1 is sufficiently active, it prevents the bottom-up input to A from generating a reset signal to F_2 . When activity at F_1 is attenuated due to mismatch, the orienting subsystem A resets F_2 (Fig. 2.3b,c,d). In this way, the orienting subsystem distinguishes between active mismatch and passive inactivity at F_1 .

Collapse of Bottom-Up Activation due to Template Mismatch. A finer analysis of network dynamics, with particular emphasis on the 2/3 Rule, leads to a vigilance mechanism capable of regulating how coarse the learned categories will be. Suppose that a bottom-up input pattern has activated F_1 and blocked activation of A (Fig. 2.3a). Suppose, moreover, that F_1 activates an F_2 node that reads out a template that badly mismatches the bottom-up input at F_1 (Fig. 2.3b). Due to the 2/3 Rule, many of the F_1 nodes that were activated by the bottom-up input alone are suppressed when the top-down template is read out. Because this mismatch event causes a large collapse in the total activity across F_1 , it leads to a large reduction in the total inhibition that F_1 delivers to A. If this reduction is sufficiently large, then the excitatory bottom-up input to A generates a non-specific reset signal from A to F_2 (Fig. 2.3c).

Reset Criterion. The reset criterion is now characterized precisely. Let $|I|$ denote the number of F_1 nodes activated by an input I when F_2 is inactive (Fig. 2.4a). Assume that each such input pathway sends an excitatory signal of fixed size γ to A whenever I is presented, so that the total excitatory input to A is $\gamma|I|$. Assume also that each active F_1 node generates an inhibitory signal of fixed size δ to A. Then, if $|X|$ is the number of active F_1 nodes in the pattern X, the total inhibitory input from F_1 to A is $\delta|X|$. When

$$\gamma|I| > \delta|X|, \quad (2)$$

A receives a net excitatory signal and generates a nonspecific reset signal to F_2 (Fig. 2.3c). The quantity

$$\rho = \frac{\gamma}{\delta} \quad (3)$$

is called the *vigilance parameter* of the orienting subsystem A. By equations (2) and (3), STM reset is initiated when

$$\rho > \frac{|X|}{|I|}. \quad (4)$$

STM reset is prevented when

$$\rho \leq \frac{|X|}{|I|}. \quad (5)$$

In other words, the proportion $|X|/|I|$ of the input pattern I that is matched by the top-down template to generate X must exceed ρ to present STM reset.

While F_2 is inactive (Fig. 2.4b), $|X| = |I|$. Activation of A is always forbidden in this case to prevent an input I from resetting its correct F_2 code. By (5), this constraint is achieved if

$$\rho \leq 1; \quad (6)$$

that is, if $\gamma \leq \delta$.

This line of argument can be intuitively recapitulated as follows. Due to the 2/3 Rule, a bad mismatch at F_1 causes a large collapse of total F_1 activity, which leads to activation of A. For this to happen, the system must maintain a measure of the prior level of total F_1 activity and compare this criterion level with the collapsed level of total F_1 activity. The criterion level is computed by summing bottom-up inputs at A. This sum provides a stable criterion because it is proportional to the initial activation of F_1 by the bottom-up input, and yet it remains unchanged as the matching process unfolds in real-time.

Distinguishing Signal from Noise in Patterns of Variable Complexity: Weighing the Evidence

Several properties follow from this conception of how the orienting system is engaged by mismatch within the attentional subsystem. The simulation in Fig. 2.6 illustrates, for example, how the network automatically rescales its matching criterion. On the first four trials, the patterns are presented in the order ABAB. By trial 2, coding is complete. Pattern A directly accesses node v_1 on trial 3, and pattern B directly accesses node v_2 on trial 4. Thus, patterns A and B are coded within different categories. On trials 5–8, patterns C and D are presented in the order CDCD. Patterns C and D are constructed from patterns A and B, respectively, by adding identical upper halves to A and B. Thus, pattern C differs from pattern D at the same locations where pattern A differs from pattern B. Because patterns C and D represent many more active features than patterns A and B, however, the difference between C and D is treated as noise, whereas the difference between A and B is considered significant. In particular, both patterns C and D are coded within the same category on trials 7 and 8.

VIGILANCE LEVEL TUNES CATEGORICAL COARSENESS: DISCONFIRMING FEEDBACK

The previous section showed how, given each fixed vigilance level, the network automatically rescales its sensitivity to patterns of variable complexity. This section shows that changes in the vigilance level can regulate the coarseness of

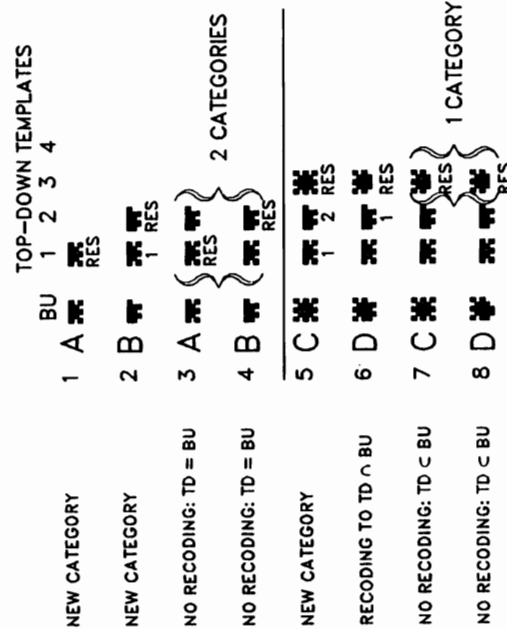


FIG. 2.6. Distinguishing noise from patterns in inputs of variable complexity: Input patterns A and B are coded by the distinct category nodes v_1 and v_2 , respectively. Input patterns C and D include A and B as subsets, but C and D also possess identical subpatterns of additional features. Due to this additional pattern complexity, C and D are coded by the same category node v_3 . At this vigilance level ($\rho = .8$), the network treats the difference between C and D as noise and suppresses the discordant elements in the v_3 template. By contrast, it treats the difference between A and B as informative and codes the difference in the v_1 and v_2 templates.

the categories that are learned in response to a fixed sequence of input patterns.

A low vigilance level leads to learning of coarse categories, whereas a high vigilance level leads to learning of fine categories. Suppose, for example, that a low vigilance level has led to a learned grouping of inputs that need to be distinguished for successful adaptation to a prescribed input environment, but that a punishing event occurs as a consequence of this erroneous grouping. Suppose that, in addition to its negative reinforcing effects, the punishing event also has a direct cognitive effect; namely, that it increases sensitivity to pattern mismatches. Such an increase in sensitivity is modeled within the network by an increase in the vigilance parameter, ρ (Fig. 2.1a). Increasing this single parameter enables the network to discriminate patterns that previously were lumped together. Once these patterns are coded by different categories in F_2 , the different categories can be associated with different behavioral responses. In this way, environmental feedback acts to reset a single nonspecific parameter that interacts with the internal organization of the network to parse more finely whatever input patterns happen to occur. The vigilance parameter is increased if all the signals

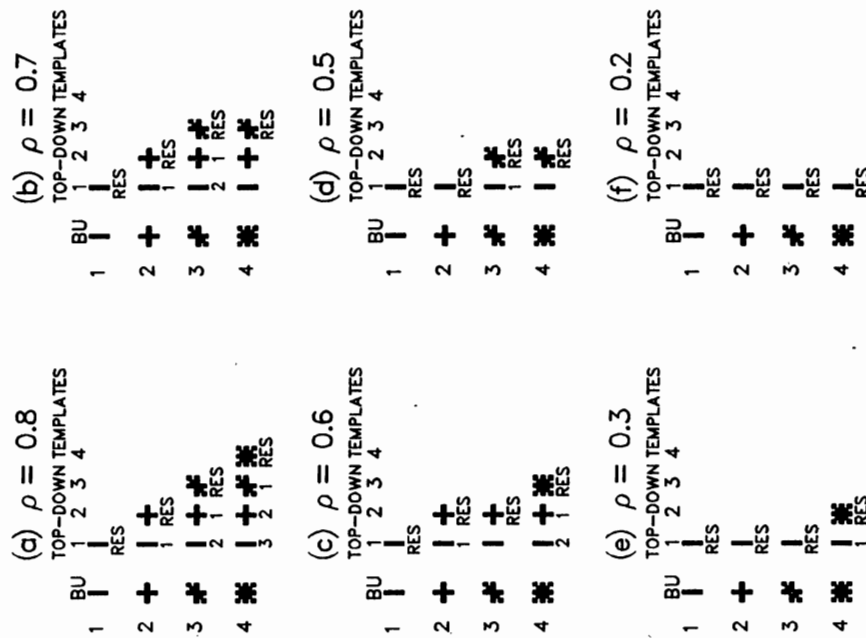


FIG. 2.7. Influence of vigilance level on categorical groupings: As the vigilance parameter ρ decreases, the number of categories progressively decreases.

from the input pattern to A are nonspecifically amplified, so that parameter γ increases. Alternatively, ρ may be increased either by a nonspecific decrease in the size δ of signals from F_1 to A; or by direct input signals to A. In a similar fashion, a decrease of the vigilance parameter enables coarser discriminations to be learned.

Figure 2.7 describes a series of simulations in which four input patterns—A, B, C, D—are coded. In these simulations, $A \subset B \subset C \subset D$. The different parts of the figure show how categorical learning changes with changes of ρ . When $\rho = .8$ (Fig. 2.7a), 4 categories are learned: (A)(B)(C)(D). When $\rho = .7$ (Fig. 2.7b), 3 categories are learned: (A) (B)(C,D). When $\rho = .6$ (Fig. 2.7c), 3 different categories are learned: (A)(B,C)(D). When $\rho = .5$ (Fig. 2.7d), 2

categories are learned: (A,B)(C,D). When $\rho = .3$ (Fig. 2.7e), 2 different categories are learned: (A,B,C)(D). When $\rho = .2$ (Fig. 2.7f), all the patterns are lumped together into a single category.

Comparison with Data about Event-Related Potentials

An ideal experimental paradigm for testing the theory is one in which expectancies are learned, matching processes are manipulated, and the experimental measures are sensitive to state-dependent patterning of neuronal activities across large ensembles of cells. Experiments measuring event-related potentials (ERPs) are excellent probes of such processes. That is why ERP data have been analyzed and predicted by adaptive resonance theory since its inception (Banquet & Grossberg, 1987; Grossberg, 1975, 1978, 1984, 1987a).

Recent ERP data have provided increasing experimental support for the formal processes developed within the theory. An early success was the prediction of a learned top-down template-matching process whose formal properties anticipated discovery of the *processing negativity* (PN) evoked potential (Näätänen, 1982; Näätänen, Gaillard, & Mäntysalo, 1978). This discussion briefly compares recent ERP data reviewed by Näätänen (1985) with formal properties of the theory. Another comparison of recent ERP data and theory is provided by Banquet and Grossberg (1987). Together these comparisons indicate that theoretically predicted processing sequences have striking correlates in ERP data.

Näätänen (1985) noted that the largest and longest PN is generated on presentation of a target stimulus, that a somewhat smaller and shorter PN is generated on presentation of a different but similar stimulus, and that little or no PN is generated on presentation of a stimulus that is very different from the target stimulus. These properties are consistent with the hypothesis that a PN develops when a subliminal top-down template elicits supraliminal activation only to the extent to which it is matched by the bottom-up input pattern, as in the 2/3 Rule.

Näätänen (1985) remarked: "As time from the stimulus onset increased, the PN became increasingly more selective or 'tuned' to the target value. Thus PN appears to be a real-time physiological sign of a matching process which continues longer the more similar the stimulus is to the 'attended-channel' stimulus." This property supports the hypothesis that, as the matching process unfolds, it causes reverberating feedback to be exchanged between F_1 and F_2 . This reverberatory exchange can deform and amplify STM as it finds the best match, or fusion, of bottom-up input and top-down expectancy (Grossberg, 1978, 1980).

Using the term attentional trace (AT) for the process subserving the PN, Näätänen (1985) pointed out: "At the very instant when the subject notices the difference between the mental image maintained and the current stimulus, he stops this maintenance and the AT vanishes instantaneously." This property supports the hypothesis that mismatch of the top-down template with the bottom-

up input causes a rapid reset of the STM code that has been reading-out the top-down template (Fig. 2.3c). Näätänen (1985) commented further: "This involuntary sampling, or high distractibility of selective attention, is of vital biological significance." The theory clarifies this remark by showing how such "distractibility" can buffer old learning and initiate search for an appropriate STM representation of the current stimulus.

Näätänen (1985) alluded to this relationship by discussing a property of the mismatch negativity (MMN), which is one of the ERPs that are elicited by deviant stimuli: "When a mismatch process is vigorous (eliciting an early sharp MMN), the input almost invariably draws attention to it." In other words, the mismatch process, by resetting STM, sets the stage for instating a new STM representation that better represents the input (Fig. 2.3d). This new STM representation activates a top-down template that focuses attention on the salient groupings within the stimulus.

Näätänen (1985) also commented on the role of learning in establishing an AT and its PN: "I presume that the development of this sensory-perceptual mental image is paralleled by the emergence of AT in the sensory system." In other words, as the bottom-up LTM code is learned through tuning of the adaptive filter from F_1 to F_2 , learning of the top-down template from F_2 to F_1 also occurs. In addition, if "the frequency of occurrence of the infrequent target feature is increased it can enter, because of frequent sensory reinforcement, the AT which is then set up for recognition of the target directly." Thus, as the top-down template is progressively learned, it can more strongly resonate with an input pattern and can lead to a resonant recognition event without a search.

Näätänen (1985) observed: "Paying attention seems to have a self-inhibitory nature . . . , the pressure for switching increasing with length of time of continuously maintaining attention directed to a certain object . . . or when the change is repetitive." Within the adaptive resonance theory, an arousal burst from the orienting subsystem causes STM reset at F_2 by means of a *gated dipole field*. Using such a circuit design, an STM reverberation within F_2 becomes weaker through time due to the habituation of chemical transmitters that exist within the reverberating feedback pathways. Such habituation of STM within F_2 has been used to analyze a variety of data concerning spontaneous perceptual switches, as occur during perception of the Necker cube and during monocular or binocular rivalry (Grossberg, 1980). Within the present context, this tendency towards switching helps to explain the "pressure for switching" an attentional trace.

In summary, recent ERP data are consistent with theoretical predictions concerning the processes of top-down template learning and priming, matching with bottom-up input patterns, and mismatch-mediated reset of the STM code that controls the read out of the top-down template. This chapter thus illustrates how concepts and data concerning attention, recognition, learning, reinforcement, and event-related potentials can be unified within a single neurally based theory.

Mathematical Equations

Summarized next are the hypotheses and parameter constraints that define the ART 1 model. The model's STM processes are assumed to occur so much more quickly than its LTM processes that STM remains in approximate equilibrium with respect to LTM. Thus the STM processes are approximated by algebraic equations that define the equilibrium values of the STM differential equations. See Carpenter and Grossberg (1987a) for a definition of these STM differential equations.

This approximation obeys the following equations.

Binary Input Patterns

$$I_i = \begin{cases} 1 & \text{if } i \in I \\ 0 & \text{otherwise} \end{cases} \quad (7)$$

Automatic Bottom-Up Activation and 2/3 Rule

$$X = \begin{cases} 1 & \text{if } F_2 \text{ is inactive} \\ 1 \cap V^{(U)} & \text{if the } F_2 \text{ node } v_j \text{ is active} \end{cases} \quad (8)$$

Weber Law Rule and Bottom-Up Associative Decay Rule

$$\frac{d}{dt} z_{ij} = \begin{cases} K[(1 - z_{ij})L - z_{ij}(|X| - 1)] & \text{if } i \in X \text{ and } f(x_i) = 1 \\ -K|X| z_{ij} & \text{if } i \notin X \text{ and } f(x_i) = 1 \\ 0 & \text{if } f(x_i) = 0 \end{cases} \quad (9)$$

Template Learning Rule and Top-Down Associative Decay Rule

$$\frac{d}{dt} z_{ji} = \begin{cases} -z_{ji} + 1 & \text{if } i \in X \text{ and } f(x_j) = 1 \\ -z_{ji} & \text{if } i \notin X \text{ and } f(x_j) = 1 \\ 0 & \text{if } f(x_j) = 0 \end{cases} \quad (10)$$

Reset Rule

An active F_2 node v_j is reset if

$$\frac{|I \cap V^{(U)}|}{|I|} < \rho \equiv \frac{\gamma}{\delta} \quad (11)$$

Once a node is reset, it remains inactive for the duration of the trial.

F_2 Choice and Search

If J is the index set of F_2 nodes which have not yet been reset on the present learning trial, then

$$f(x_j) = \begin{cases} 1 & \text{if } T_j = \max\{T_k : k \in J\} \\ 0 & \text{otherwise} \end{cases} \quad (12)$$

where

$$T_j = D_2 \sum_{i \in X} z_{ij} \quad (13)$$

In addition, all STM activities x_i and x_j are reset to zero after each learning trial. The initial bottom-up LTM traces $z_{ij}(0)$ are chosen to satisfy the

Direct Access Inequality

$$0 < z_{ij}(0) < \frac{L}{L - 1 + M} \quad (14)$$

The initial top-down LTM traces are chosen to satisfy the

Template Learning Inequality

$$1 \geq z_{ji}(0) > \bar{z} \equiv \frac{B_1 - 1}{D_1} \quad (15)$$

TABLE 2.1
Parameter Constraints

$A_1 \geq 0$
$C_1 \geq 0$
$\max\{1, D_1\} < B_1 < 1 + D_1$
$0 < \epsilon < 1$
$K = O(1)$
$L > 1$
$0 < \rho \leq 1$
$0 < z_{ij}(0) < \frac{L}{L-1+M}$
$1 \geq z_{ji}(0) > \bar{z} \equiv \frac{B_1-1}{D_1}$
$0 \leq I_i, f_i, g_i, h_i \leq 1$

Fast Learning

It is assumed that fast learning occurs so that, when v_j in F_2 is active, all LTM traces approach the asymptotes,

$$z_{ij} \cong \begin{cases} \frac{L}{L-1+|X|} & \text{if } i \in X \\ 0 & \text{if } i \notin X \end{cases} \quad (16)$$

and

$$z_{ji} \cong \begin{cases} 1 & \text{if } i \in X \\ 0 & \text{if } i \notin X \end{cases} \quad (17)$$

on each learning trial. A complete listing of parameter constraints is provided in Table 2.1.

ACKNOWLEDGMENTS

This chapter was supported in part by a grant to the first author by the Army Research Office (ARO DAAG-29-85-K-0095 and the National Science Foundation (NSF-DMS-84-13119); also partial support by a grant to the second author by the Air Force Office of Scientific Research (AFOSR 85-0149) and the National Science Foundation (NSF IRI-84-17756).

We thank Cynthia Suchta and Carol Yanakakis for their valuable assistance in the preparation of the manuscript.

REFERENCES

- Banquet, J.-P., Grossberg, S. (1988). Probing cognitive processes through the structure of event-related potentials during learning: An experimental and theoretical analysis. *Applied Optics*, 26, 4931-4946.
- Carpenter, G. A., & Grossberg, S. (1987a). A massively parallel architecture for a self-organizing neural pattern recognition machine. *Computer Vision, Graphics, and Image Processing*, 37, 54-115.
- Carpenter, G. A., & Grossberg, S. (1987b). ART 2: Self-organization of stable category recognition codes for analog input patterns. In *Proceedings of the IEEE international conference on neural networks*, San Diego.
- Carpenter, G. A., & Grossberg, S. (1987c). ART 2: Stable self-organization of pattern recognition codes for analog input patterns. *Applied Optics*, 26, 4919-4930.
- Carpenter, G. A., & Grossberg, S. (1988). Neural dynamics of category learning and recognition: Attention, memory consolidation, and amnesia. In J. Davis, R. Newburgh, & E. Wegman (Eds.), *Brain structure, learning, and memory*. AAAS Selected Symposia Series, 105, 233-290.
- Cohen, M. A., & Grossberg, S. (1986). Neural dynamics of speech and language coding: Developmental programs, perceptual grouping, and competition for short-term memory. *Human Neurobiology*, 5, 1-22.
- Cohen, M. A., & Grossberg, S. (1987). Masking fields: A massively parallel architecture for

learning, recognizing, and predicting multiple groupings of patterned data. *Applied Optics*, 26, 1866-1891.

Grossberg, S. (1975). A neural model of attention, reinforcement, and discrimination learning. *International Review of Neurobiology*, 18, 263-327.

Grossberg, S. (1978). A theory of human memory: Self-organization and performance of sensory-motor codes, maps, and plans. In R. Rosen & F. Snell (Eds.), *Progress in theoretical biology* (Vol. 5, pp. 233-374). New York: Academic Press.

Grossberg, S. (1980). How does a brain build a cognitive code? *Psychological Review*, 87, 1-51.

Grossberg, S. (1982a). *Studies of mind and brain: Neural principles of learning, perception, development, cognition, and motor control*. Boston: Reidel Press.

Grossberg, S. (1982b). Processing of expected and unexpected events during conditioning and attention: A psychophysiological theory. *Psychological Review*, 89, 529-572.

Grossberg, S. (1982c). A psychophysiological theory of reinforcement, drive, motivation, and attention. *Journal of Theoretical Neurobiology*, 1, 286-369.

Grossberg, S. (1983). The quantized geometry of visual space: The coherent computation of depth, form, and lightness. *Behavioral and Brain Sciences*, 6, 625-692.

Grossberg, S. (1984). Some psychophysiological and pharmacological correlates of a developmental, cognitive, and motivational theory. In R. Karrer, J. Cohen, & P. Tuetting (Eds.), *Brain and information: Event related potentials*. New York: New York Academy of Sciences.

Grossberg, S. (1986). The adaptive self-organization of serial order in behavior: Speech, language, and motor control. In E. C. Schwab & H. C. Nusbaum (Eds.), *Pattern recognition by humans and machines, Vol. 1: Speech perception* (pp. 187-294). New York: Academic Press.

Grossberg, S. (1987a). Competitive learning: From interactive activation to adaptive resonance. *Cognitive Science*, 11, 23-63.

Grossberg, S. (1987b). The adaptive brain, I: Cognition, learning, reinforcement, and rhythm. Amsterdam: Elsevier/North-Holland.

Grossberg, S., & Levine, D. S. (1987). Neural dynamics of attentionally modulated Pavlovian conditioning: Blocking, interstimulus interval, and secondary reinforcement. *Applied Optics*, 26, 5015-5030.

Grossberg, S., & Schmajuk, N. A. (1987). A neural network architecture for attentionally modulated Pavlovian conditioning: Conditioned reinforcement, inhibition, and opponent processing. *Psychobiology*, 15, 195-240.

Näätänen, R. (1982). Processing negativity: An evoked-potential reflection of selective attention. *Psychological Bulletin*, 92, 605-640.

Näätänen, R. (1985). Theory of auditory selective attention based on event-related potentials in man. Preprint.

Näätänen, R., Gaillard, A., & Mäntysalo, S. (1978). The N1 effect of selective attention reinterpreted. *Acta Psychologica*, 42, 313-329.

Sokolov, E. N. (1958). Perception and the conditioned reflex. Moscow: Moscow University Press.

Sokolov, E. N. (1968). *Mechanisms of memory*. Moscow: Moscow University Press.

Tolman, E. C. (1932). *Purposive behavior in animals and men*. New York: Century Press.