

A deterministic mathematical model of the transmission of information I

By E. GESZTELYI (Debrecen)

Dedicated to Professor Zoltán Daróczy on his 50th birthday

Introduction

According to the basic intuitive background of information theory, a communication system can be illustrated by Shannon's famous blockdiagram shown on Fig. 1.

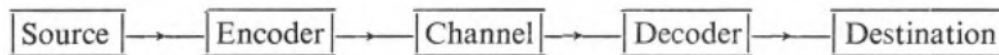


Fig. 1

The source and destination can be separated in space or time and the channel is destined for bridging over this separation. The terminology of the model shown on Fig. 1 reflects the point of view of communication between terminals separated in space. But, according to a widespread view, these models are equally suitable for describing data storage systems appropriately interchanging the roles of time and space ([1]).

The theoretical treatment of the transmission of information dates back to HARTLEY (1928). The basic concepts of information theory are due to SHANNON (1948). Since then the following technical developments came into being affecting both the telecommunication and the data processing systems a) the realization of astronautics, b) the appearance of the electronic digital computers, c) the widespreading of microprocessors and other integrated circuit (IC) chips such as Random Access Memories (RAM), Read Only Memories (ROM), programable communication interfaces, etc. (obtainable all in specialist shops¹⁾).

Now, we can speak not only theoretically but also practically about the case when the source and destination are separated both in space and time. Indeed, this situation occurs by the telecommunication with a satellite which runs along an orbit far from the Earth. This shows that *the separatedness of the source and destination in space or time is irrelevant from the point of view, whether a communication system is a device of data storage or else it is a device of transmission of information.*

¹⁾ See e.g. Intel Component Data Catalogs.

Several natural or artificial systems are able to emit information. Every known natural source emits continuous signals.²⁾

The whole mathematical discipline of information theory refers to the artificial case where the source emits discrete signals from a set, the source alphabet. Similarly, the channel is supposed to be capable of transmitting successively symbols from a given set, the input channel alphabet. The output alphabet of the channel may differ from the input alphabet of the channel. In the ideal case of a noiseless channel the output is identical to the input. The information theory is a statistical theory of communications studying the problem of reliable transmission of information at a possibly low cost for given source, channel and fidelity criterion ([1]). Thus, the noiseless channel is the trivial case from the point of view of information theory.

The telegraph was the first practical realization of the discrete channel. The importance of the discrete channel is emphasized by the fact that there are used discrete channels also in the modern digital computers.

But, it is an essential difference between the two kinds of transmission of information. Namely, if the channel is not noiseless in a computer then the computer is of no use since, as it is known from the programming of computers, the alteration of a single bit can ruin the whole program. Thus, *the information theory is unapplicable in our model since the noiselessness of the channel must be a starting requirement.*

Therefore, we try develop deterministic models of communication systems. Presently, we deal with the problem of transmission of information. ([4] deals with memory expansions.)

Preliminaries, denotations

Automata. By an automaton $A=(S, X, \delta)$ we mean a finite automaton with state set S , input set X and transition function $\delta: S \times X \rightarrow S$. The symbol δ as well as the extended transition function δ^* will be used here mainly in case of danger of misunderstandings. We shall write sp instead of $\delta^*(s, p)$ ($s \in S, p \in X^+$).

By a product of automata we shall mean always a generalized product introduced by F. GÉCSEG ([5]).

Definition A. By the product of automata $A_i=(S_i, X_i, \delta_i)$ ($i=1, \dots, n$) with respect to an alphabet X and a mapping

$$(1) \quad \psi: S_1 \times \dots \times S_n \times X \rightarrow X_1^* \times \dots \times X_n^*,$$

we mean the automaton $A=(S, X, \delta)$ where $S=S_1 \times \dots \times S_n$ and $\delta[(s_1, \dots, s_n), x] = (s_1 p_1, \dots, s_n p_n)$, where $x \in X$,

$$(s_1, \dots, s_n) \in S \quad \text{and} \quad (p_1, \dots, p_n) = \varphi(s_1, \dots, s_n, x).$$

²⁾ Except for the following neutrino event: "On February 23rd. 1987, a handful of neutrino particles were detected from Supernova 1987 A in the Large Magellanic Cloud... they brought us direct evidence of the violent conditions at the very heart of the tremendous catalysm some 160 000 light-years away. The observations represent the first neutrino detections from a known extre terrestrial source: a point source in another galaxi — the supernova". See May, 1987, Sky & Telescope.

The product A is denoted by $A = (A_1 \times \dots \times A_n) [X, \varphi]$. The product $(A_1 \times \dots \times A_n) [X, \varphi]$ will be called a Gluškov product of automata $A_1, \dots, A_n$ in the special case when

$$\varphi: S_1 \times \dots \times S_n \times X \rightarrow X_1 \times \dots \times X_n. \quad \square$$

Partitions. By a partition P of a nonempty set S , it is meant a set of nonempty, disjoint subsets of S of which union covers S . The set of all partitions of S will be denoted by $\text{PART}(S)$.

Every partition P of S can be generated by means of an appropriate surjective function $f_P: S \rightarrow H$ in the following sense. Any block of P is the inverse image of an $h \in H$: $P = \{B_h \subseteq S \mid B_h = f_P^{-1}(h), h \in H\}$ where f_P^{-1} is the inverse of f_P in the sense that $s \in f_P^{-1}(h) \Leftrightarrow f_P(s) = h$.

The mapping f_P is called a generator function of the partition P . If $v = |P|$ is finite, we choose $H = \Gamma_v = \{0, 1, \dots, v-1\}$.

Theorem P1. Let $f: S \rightarrow H_1$ and $g: S \rightarrow H_2$ be surjective functions. Then f and g generate the same partition P iff there exists a bijective map $\Psi: H_1 \rightarrow H_2$ such that $g = \Psi(f)$. ■ ([3])

Definition P1. The product $P \cdot Q$ of $P, Q \in \text{PART}(S)$ is defined by

$$(2) \quad P \cdot Q = \{B = K \cap L \mid (K, L) \in P \times Q \wedge K \cap L \neq \emptyset\}. \quad \square$$

Theorem P2. With respect to the product (2), $\text{PART}(S)$ forms an idempotent, commutative semigroup with $I = \{S\}$ as unit element. ■

We denote by $\text{FPART}(S)$ the set of all finite partitions of the set S . A finite subset of $\text{FPART}(S)$ is called a *core* on S .

Theorem P3. If $P_1, \dots, P_n \in \text{FPART}(S)$ then

$$\max \{|P_1|, \dots, |P_n|\} \leq |P_1 \cdot \dots \cdot P_n| \leq |P_1| \cdot \dots \cdot |P_n|. \quad \blacksquare$$

The set of all cores on S will be denoted by $\text{CORE}(S)$.

Definition P2.

- (i) The core $C = \{P_1, \dots, P_n\} \in \text{CORE}(S)$ is said to be *independent* if $|P_1 \dots P_n| = |P_1| \cdot \dots \cdot |P_n|$.
- (ii) By the *volume* of the core $\{P_1, \dots, P_n\} \in \text{CORE}(S)$ we mean the product $P_1 \cdot \dots \cdot P_n \in \text{PART}(S)$.
- (iii) By the *size* of the core $C = \{P_1, \dots, P_n\} \in \text{CORE}(S)$, we mean the cardinality of the volume of C . If $\|C\|$ denotes the size of C then we have $\|C\| = |P_1 \cdot \dots \cdot P_n|$.
- (iv) Let $C_1, C_2 \in \text{CORE}(S)$. C_1 is said to be *independent of* C_2 if $C_1 \cup C_2$ is independent and $C_1 \cap C_2 \subseteq \{I\}$. ■

Theorem P4. For all $C_1, C_2 \in \text{CORE}(S)$ the inequality

$$\|C_1 \cup C_2\| \leq \|C_1\| \cdot \|C_2\|$$

holds moreover, if C_1 is independent of C_2 then the equality

$$\|C_1 \cup C_2\| = \|C_1\| \cdot \|C_2\|$$

is valid. ■

Remark P1. It is easy to prove the following statement:

Let $f_1: S \rightarrow H_1, \dots, f_n: S \rightarrow H_n$ be some generator functions of the partitions $P_1, \dots, P_n \in \text{PART}(S)$, respectively. Then a generator function of the volume $P_1 \cdot \dots \cdot P_n$ of the core $\{P_1, \dots, P_n\}$ is the vector valued function $\mathbf{f}: S \rightarrow H$ for which $\forall s \in S$:

$$\mathbf{f}(s) = (f_1(s), \dots, f_n(s)) \in H \subseteq H_1 \times \dots \times H_n. \quad \blacksquare$$

A partition Q of a set S is known as a *refinement* of the partition P of S if every block of Q is a subset of some block of P .

Proposition P1. *If $Q \in \text{FPART}(S)$ is a refinement of $P \in \text{FPART}(S)$ then $|P| \equiv |Q|$ where the equality holds exactly in case $P = Q$.*

PROOF. Since $P, Q \in \text{PART}(S)$, if Q is a refinement of P then any block of P covers at least a block of Q and so, $|P| \equiv |Q|$ is clear. Now, if here $|P| = |Q|$ then any block of P agrees with a block of Q . Thus, $P = Q$. The inference $P = Q \Rightarrow |P| = |Q|$ is trivial. $\blacksquare$

Memories. As mentioned above, according to the basic concept of the information theory, the unique distinction between transmission and storage of information is the separatedness of the source and destination in time or space. But, we showed that the separatedness of the source and destination in time or space, practically seems to be an unessential distinction. Moreover, we will not introduce the notion of time and space in our theory. Therefore it is meaningless to speak about separatedness in time or space in the framework of the present model.

Definition M1 ([3]). By a memory $\mathbf{M}$, we mean a triplet $\mathbf{M} = (S, \mathbf{C}, s)$ where S is a nonvoid set, $\mathbf{C}$ is a set of some partitions of S and s is a variable on S . S is called the set of possible states, the current value of s is the instant state of $\mathbf{M}$. $\mathbf{C}$ is called the core of $\mathbf{M}$ and any partition $P \in \mathbf{C}$ is said to be a memory cell of $\mathbf{M}$. It is supposed, that the blocks of the partition P are indexed with the elements of some index-set Γ_P . (If P is finite, we can choose $\Gamma_P = \{0, 1, \dots, |P| - 1\}$.) Thus, we can write $P = \{B_v | v \in \Gamma_P\}$. The content of the memory cell P depends on the instant state s : The content of P at state s is the index of that block of P which contains s . The memory $\mathbf{M}$ is finite if $\mathbf{C}$ is a finite set of finite partitions of S . $\square$

We showed in [3] that the above notion is a faithful abstraction of the memories used in the practical computer technics.

It is well known that the information cannot spread in itself. In practice, the information is carried always by some material or energy. If the information is carried by discrete signals, then a memory, which is capable to store these signals, may be regarded also as a carrier of information.

Thus, the memory and the channel are two aspects of one and the same thing and we can describe the channels by means of the same mathematical tools as the memories are described.

Definition M2. The set of all memories over the set S will be denoted by $\text{MEM}(S)$:

$$(3) \quad \text{MEM}(S) = \{\mathbf{M} = (S, \mathbf{C}, s) | \emptyset \neq \mathbf{C} \subset \text{PART}(S)\}.$$

Although a memory is not a set, we define the set theoretical operations in MEM (S). For $M_i = (S, C_i, s)$ ($i=1, 2$) let

$$(4) \quad M_1 \cup M_2 = (S, C_1 \cup C_2, s), \quad \text{union of } M_1 \text{ and } M_2,$$

$$(5) \quad M_1 \cap M_2 = (S, C_1 \cap C_2, s), \quad \text{intersection of } M_1 \text{ and } M_2,$$

$$(6) \quad M_1 \subset M_2 \Leftrightarrow C_1 \subset C_2, \quad M_2 \text{ is an extension of } M_1. \quad \square$$

We shall suppose tacitly that any core contains the improper partition $I = \{S\}$ even if I is not indicated between the elements of C in one enumeration. Thus, the intersection of memories exists always since in this way, the intersections of the cores will be never empty.

If $P \in C$ is a memory cell of a memory $M = (S, C, s)$ then the memory $P = (S, \{P\}, s)$ will be called a *register* of M . The content of the register P at state s will be denoted by $(P)_s$ and per definition, it agrees with the content of the cell P at s . Thus, e.g. the improper register $I = (S, \{I\}, s)$ contains always the zero: $(I)_s = 0$ for all $s \in S$ (or $(I)_s$ is equal to another fixed element according to the chosen index set Γ_I which in consequence of $I = \{S\}$, can be a singleton only). So, the improper register is unsuited for storage or transmission of information.

The notion of content can be defined not only for registers but also for the general memories.

Definition M3. Let $M = (S, C, s)$ be a finite memory on S .

- (i) The volume of the core C will be called the *vector cell* of M .
- (ii) Let a generator function $f: S \rightarrow H$ of the vector cell of M be fixed. We shall say that f is the generator function of the memory M .
- (iii) The content of the memory M at state s will be denoted by $(M)_s$ and defined by the equation

$$(7) \quad (M)_s = f(s) \quad \forall s \in S. \quad \square$$

Remark M1. If we know the content of the memory M then we know the content of every single register of M too. Namely, according to Remark P1 and Definition M3, the generator function of M can be written in the form

$$(8) \quad f(s) = (f_1(s), \dots, f_n(s))$$

where $f_i: S \rightarrow H_i$ is the generator function of P_i ($i=1, \dots, n$) and $C = \{P_1, \dots, P_n\}$. Thus, we have $\forall s \in S$:

$$(9) \quad (P_i)_s = f_i(s) \in H_i \quad (i = 1, \dots, n). \quad \blacksquare$$

Theorem M1. Let M_1 and M_2 be finite memories on the set S . Any content of M_1 can be exchanged for any possible content of M_1 in such a way that the content of M_2 remains unchanged, iff the core of M_1 is independent of the core of M_2 . $\blacksquare$

All the results of this introductory chapter are proved in [3] or they are immediate consequences of the results given in [3].

1. The refinement of memories

If $M_1 \subset M_2$ holds for some memories $M_1, M_2 \in \text{MEM}(S)$ then we say that M_2 is an extension of the memory M_1 . In this chapter, we shall deal with another kind of increase of memories. Namely, we can obtain a memory M_2 with larger storage capability as the memory M_1 by means of the refinement of the memory cells of M_1 . We shall speak in this case about the *refinement* of the memory in the sense of the following definition.

Definition 1.1. The memory $M_2 = (S, C_2, s)$ is said to be a *refinement* of the memory $M_1 = (S, C_1, s)$ if there exists a bijection $\varphi: C_1 \rightarrow C_2$ such that for all $P \in C_1$, $\varphi(P)$ is a refinement of P . Moreover, if there exists a $P \in C_1$ such that $|P| < |\varphi(P)|$ then we say that M_2 is a *proper refinement* of M_1 .

Definition 1.2. We shall denote by $\text{FMEM}(S)$ the set of all finite memories on S . We say that the memory M is independent (in itself) if the core of M is independent.

Proposition 1.1. *The refinement of memories is an ordering relation on $\text{FMEM}(S)$.*

PROOF. Clearly, the refinement is reflexive, since any memory is a refinement of itself. Let $M_i = (S, C_i, s) \in \text{FMEM}(S)$ ($i=1, 2, 3$) be memories such that M_2 is a refinement of M_1 and M_3 is a refinement of M_2 . We have to show that M_3 is a refinement of M_1 . There exist bijective mappings $\varphi_{12}: C_1 \rightarrow C_2$ and $\varphi_{23}: C_2 \rightarrow C_3$ such that $\forall P \in C_1: \varphi_{12}(P)$ is a refinement of P and $\forall Q \in C_2: \varphi_{23}(Q)$ is a refinement of Q . Let the mapping $\varphi_{13}: C_1 \rightarrow C_3$ be defined by $\varphi_{13}(P) = \varphi_{23}[\varphi_{12}(P)]$. We see that any block of $\varphi_{13}(P)$ is covered by a block of $\varphi_{12}(P)$ which is a subset of some block of P . This shows that M_3 is a refinement of M_1 and the refinement is transitive. Let M_2 be a refinement of M_1 and M_1 be a refinement of M_2 . We shall show that $M_1 = M_2$. There exist bijective mappings $\varphi_{12}: C_1 \rightarrow C_2$ and $\varphi_{21}: C_2 \rightarrow C_1$ such that $\forall P \in C_1: \varphi_{12}(P)$ is a refinement of P and $\forall Q \in C_2: \varphi_{21}(Q)$ is a refinement of Q . Thus, by Proposition P1, we have

$$(1.1) \quad \forall P \in C_1: |P| \leq |\varphi_{12}(P)|$$

and

$$(1.2) \quad \forall Q \in C_2: |Q| \leq |\varphi_{21}(Q)|.$$

It follows from (1.1) and (1.2) for the mapping $\psi = \varphi_{21} \varphi_{12}: C_1 \rightarrow C_1$:

$$(1.3) \quad |P| \leq |\varphi_{12}(P)| \leq |\varphi_{21}[\varphi_{12}(P)]| = |\varphi_{21} \varphi_{12}(P)| = |\psi(P)|,$$

whence by induction, we get for all $k=1, 2, \dots$

$$(1.4) \quad |P| \leq |\psi^k(P)|$$

where $\psi^1(P) = \psi(P)$ and $\psi^{i+1}(P) = \psi[\psi^i(P)]$ if $i > 0$. If for some P_0 , the inequality $|P_0| < |\varphi_{12}(P_0)|$ holds, using (1.3), we obtain $|P_0| < |\psi(P_0)|$ and whence

$$|P_0| < |\psi(P_0)| \leq |\psi[\psi(P_0)]| = |\psi^2(P_0)|$$

and by induction, we get

$$(1.5) \quad |P_0| < |\psi^k(P_0)| \quad k = 1, 2, \dots$$

Now, we shall prove that M_2 cannot be a proper refinement of M_1 . Suppose, in contrary, that M_2 is a proper refinement of M_1 . Then, by Definition 1.1, there exists a $P_0 \in C_1$ such that $|P_0| < |\varphi_{12}(P_0)|$. Then the inequality (1.5) holds for all natural numbers k . On the other hand, $\psi: C_1 \rightarrow C_1$ is a permutation on C_1 and so, there exists a natural number m such that

$$(1.6) \quad \psi^m(P_0) = P_0.$$

thus, for $k=m$, (1.5) contradicts (1.6). This shows that M_2 cannot be a proper refinement of M_1 and so, $M_2 = M_1$. $\blacksquare$

Definition 1.3. If M_2 is a refinement of M_1 ($M_1, M_2 \in \text{FMEM}(S)$) then we shall write $M_1 \leq M_2$. If M_2 is a proper refinement of M_1 then we write $M_1 < M_2$. $\square$

Definition 1.4. By the size $\|M\|$ of the memory $M = (S, C, s) \in \text{FMEM}(S)$ we mean the size of its core. Thus, if $C_1 = \{P_1, \dots, P_n\}$ then

$$(1.7) \quad \|M\| = \|C\| = |P_1 \dots P_n|. \quad \square$$

Theorem 1.1. Let $M_1, M_2 \in \text{FMEM}(S)$ be some memories where M_2 is independent.

(i) If $M_1 \subset M_2$ then $\|M_1\| < \|M_2\|$.

(ii) If $M_1 < M_2$ then $\|M_1\| < \|M_2\|$.

PROOF. (i) Since M_2 is a proper extension of M_1 , the core of M_2 can be written in the form $C_2 = \{P_1, \dots, P_n, Q_1, \dots, Q_m\}$ where $C_1 = \{P_1, \dots, P_n\}$ is the core of M_1 . Here

$$(1.8) \quad |Q_i| > 1 \quad \text{for all } i = 1, \dots, m$$

since the partitions Q_i are those elements of C_2 which are not elements of C_1 and the improper partition $I = \{S\}$ is an element of C_1 according to our former agreement. Thus, on the basis of Theorem P3, by Definition 1.4, Definition P2 (i) using the independence of M_2 and (1.8), we get

$$\begin{aligned} \|M_1\| &= \|C_1\| = |P_1 \dots P_n| \leq |P_1| \dots |P_n| < \\ &< |P_1| \dots |P_n| \cdot |Q_1| \dots |Q_m| = |P_1 \dots P_n \cdot Q_1 \dots Q_m| = \|C_2\| = \|M_2\|. \end{aligned}$$

(ii) Since M_2 is a refinement of M_1 there is a bijection $\varphi: C_1 \rightarrow C_2$ such that for any partition $P \in C_1$, $\varphi(P)$ is a refinement of P . If $C_1 = \{P_1, \dots, P_n\}$ then let $Q_i = \varphi(P_i)$ ($i = 1, \dots, n$). Then $C_2 = \{Q_1, \dots, Q_n\}$ and according to Proposition P1, we have

$$(1.9) \quad |P_i| \leq |Q_i| \quad i = 1, \dots, n.$$

Since M_2 is a proper refinement of M_1 , on the basis of Definition 1.1, there exists at least a $j \in \{1, \dots, n\}$ for which

$$(1.10) \quad |P_j| < |Q_j|.$$

Thus, using that M_2 is independent, by Definition 1.4, Definition P2 (i) and (iii) applying Theorem P3, it follows from (1.9) and (1.10) that

$$\begin{aligned}\|M_1\| &= \|C_1\| = |P_1 \cdot \dots \cdot P_n| \leq |P_1| \cdot \dots \cdot |P_j| \cdot \dots \cdot |P_n| < \\ &< |Q_1| \cdot \dots \cdot |Q_j| \cdot \dots \cdot |Q_n| = \|C_2\| = \|M_2\|. \blacksquare\end{aligned}$$

2. The storing capacity of the memory and a characterization of that

We defined the storing capacity of finite memories previously as a mapping $\text{cap}: \text{FMEM}(S) \rightarrow \mathbf{R}_+$ such that

$$(2.1) \quad \text{cap}(M) = \log_2 \|M\|$$

where $\|M\|$ is the size of M . This definition coincides with the practical notion of the capacity if the memory is an independent union of one bit memory cells. However, not any memory is independent even if it is used in practice. For example, the memory of a microprocessor is not independent generally since the content of the Aagregister depends on the content of some other registers (see e.g. [6]). Thus, neither a software nor a hardware specialist can give a satisfactory exact answer for the question: how much is the storing capacity of a microprocessor? Using the formula (2.1), we can give an exact answer but, it will be unsatisfactory if we recommend (2.1) without any plausible theoretical justification. We gave the following characterization in [3]:

Theorem M2. *If $\text{cap}(M)$ is defined as an additive, normed function of M which is monotonic relative to the size and the size is monotonic relative to the capacity then (2.1) is valid for $\text{cap}(M)$. $\blacksquare$*

Professor Z. DARÓCZY remarked that a such characterization would be of more interest which does not rest to the notion of the size, just on the contrary, the size would be one of inferences of the characterization!

Our next result goes toward this direction with one step. Though, we cannot eliminate the size from the characterization, but as it will be shown, the condition "the capacity is monotonic relative to the size" may be omitted.

The meaning of the remaining condition concerning the size is that any enlargement of the storing capacity goes together with the enlargement of the size. In the light of Theorem 1.1, this condition seems to be plausible since, this theorem asserts that the enlargement of the capability of the storage in important cases goes together with the enlargement of the size.

We deened a definition and a lemma.

Definition 2.1. Let X be an arbitrary nonvoid set and Y_1, Y_2 be ordered sets. A function $g: X \rightarrow Y_2$ is said to be monoton increasing with respect to $f: X \rightarrow Y_1$ if

$$\forall x_1, x_2 \in X: f(x_1) < f(x_2) \Rightarrow g(x_1) < g(x_2)$$

(see [3]).

Lemma 2.1. *Let X be an arbitrary nonvoid set and Y_1, Y_2 be ordered sets such that Y_1 is totally ordered. If $g: X \rightarrow Y_2$ and $f: X \rightarrow Y_1$ are surjective functions such*

that g is monoton increasing with respect to f then there exists a monoton nondecreasing, surjective function $\lambda: Y_2 \rightarrow Y_1$ such that for all $x \in X$:

$$(2.1) \quad f(x) = \lambda[g(x)].$$

PROOF. Let $y \in Y_2$ be given arbitrarily. It follows from the surjectivity of g that there exists $x \in X$ such that $y = g(x)$. Then, let $\lambda(y) = f(x)$. We see that λ is defined in such a way that (2.1) holds. It is remained to prove only that also the other mentioned properties hold, but first at all that the definition of λ is correct in the sense that λ is a (one valued) function. Suppose that for some $y \in Y_2$ there exist $x_1, x_2 \in X$ such that $y = g(x_1) = g(x_2)$. We show that $f(x_1) = f(x_2)$. Indeed, g is monoton increasing with respect to f and so, the contrary case $f(x_1) < f(x_2)$ or $f(x_2) < f(x_1)$ yields $g(x_1) < g(x_2)$ or $g(x_2) < g(x_1)$ contradicting $g(x_1) = g(x_2)$. Now, we show the surjectivity of λ . Let $z \in Y_1$ be given arbitrarily. By the surjectivity of f , there exists $x \in X$ such that $z = f(x)$. Thus, according to the definition of λ , for $y = g(x) \in Y_2$, we have $\lambda(y) = f(x) = z$ showing the surjectivity of λ . Finally, we prove that λ is nondecreasing. Let $u, v \in Y_2$ be given so that $u < v$. Let $x_1, x_2 \in X$ be chosen so that $u = g(x_1)$ and $v = g(x_2)$. Then $f(x_1) \leq f(x_2)$ since in consequence of the condition that g is monoton increasing with respect to f , the contrary case $f(x_2) < f(x_1)$ yields the false inequality $v = g(x_2) < g(x_1) = u$. Thus $\lambda(u) = f(x_1) \leq f(x_2) = \lambda(v)$. ■

In order to guarantee the existence of arbitrary large memories, we shall suppose that the set S of states is infinite.

Definition 2.2. By a normed measure of memories, we mean a function $\mu: \text{FMEM}(S) \rightarrow \mathbf{R}$ which is additive in the sense that for all M_1, M_2

$$(2.2) \quad \mu(M_1 \cup M_2) = \mu(M_1) + \mu(M_2)$$

holds if $M_1 \in \text{FMEM}(S)$ is independent of $M_2 \in \text{FMEM}(S)$; moreover μ is normed i.e. there exists a partition $\{S_1, S_2\} \in \text{PART}(S)$ for which $\mu[(S, \{\{S_1, S_2\}, s\})] = 1$.

Theorem 2.1. If the storing capacity $\text{cap}(M)$ of memories is a normed measure such that any enlargement of the storing capacity goes together with the enlargement of the seize then $\text{cap}(M) = \log_2 \|M\|$.

PROOF. In order to apply the lemma for our case, let $X = \text{FMEM}(S)$, $Y_1 = \mathbf{R}_+$, $Y_2 = \mathbf{N}$ where $\mathbf{N}$ is the set of natural numbers. Moreover, let the functions $f: X \rightarrow Y_1$ and $g: X \rightarrow Y_2$ be the functions $f(M) = \text{cap}(M)$ and $g(M) = \|M\|$. Then the property, that the enlargement of the capacity goes together with the enlargement of the size, is another expression of the fact that the size is monoton increasing with respect to the storing capacity. Thus, by Lemma 2.1, there exists a monoton non-decreasing function $\lambda: \mathbf{N} \rightarrow \mathbf{C}_+$ such that

$$(2.3) \quad \text{cap}(M) = \lambda(\|M\|).$$

Now, we shall show that λ is a totally additive number theoretical function, i.e. for which $\forall n, m \in \mathbf{N}$:

$$(2.4) \quad \lambda(nm) = \lambda(n) + \lambda(m).$$

Let $n, m \in \mathbf{N}$ be given arbitrarily. Using the infiniteness of S , we can easy construct

partitions $P, Q \in \text{FPART}(S)$ such that $|P|=n$, $|Q|=m$ and the memory $\mathbf{M}_1 = (S, \{P\}, s)$ is independent of $\mathbf{M}_2 = (S, \{Q\}, s)$ (see [1]). Using (2.3) and (2.2) (which in consequence of the additivity of $\lambda(\mathbf{M}) = \text{cap}(\mathbf{M})$ holds), we get

$$\begin{aligned}\lambda(n) + \lambda(m) &= \lambda(\|\mathbf{M}_1\|) + \lambda(\|\mathbf{M}_2\|) = \text{cap}(\mathbf{M}_1) + \text{cap}(\mathbf{M}_2) = \text{cap}(\mathbf{M}_1 \cup \mathbf{M}_2) = \\ &= \text{cap}(S, \{P, Q\}, s) = \lambda[\|(S, \{P, Q\}, s)\|] = \lambda(|P \cdot Q|) = \lambda(|P| \cdot |Q|) = \lambda(nm).\end{aligned}$$

Thus, λ is a monoton nondecreasing totally additive number theoretical function, whence on the basis of a theorem of P. ERDŐS ([2]), we infer that $\lambda(n) = c \log n$. Thus (2.3) can be written in the form

$$(2.5) \quad \text{cap}(\mathbf{M}) = c \log \|\mathbf{M}\|.$$

Since the capacity is normed in the sense of Definition 2.2, there exists a partition $\{S_1, S_2\}$ of S such that $1 = \text{cap}[(S, \{\{S_1, S_2\}\}, s)]$. Thus, by (2.5), we get $1 = \text{cap}[(S, \{\{S_1, S_2\}\}, s)] = c \log |\{S_1, S_2\}| = c \log 2$. Whence $c = 1/\log 2$ and (2.5) can be written in the form

$$\text{cap}(\mathbf{M}) = (\log \|\mathbf{M}\|)/\log 2 = \log_2 \|\mathbf{M}\|. \quad \blacksquare$$

References

- [1] I. CSISZÁR and J. KÖRNER, Information Theory, *Akadémiai Kiadó, Budapest* (1981).
- [2] P. ERDŐS, On the distribution function of additive functions, *Ann. of Math.* 2. **47** (1946), 1—20.
- [3] E. GESZTELYI, A mathematical theory of processors I, *MTA Sztaki Közlemények* 37/1987, pp. 21—61.
- [4] E. GESZTELYI, On the memory expansion of processors. *To appear*.
- [5] F. GÉCSEG, Products of automata, *Springer-Verlag* 1986.
- [6] LANCE A. LEVENTHAL, Z80 assembly language programming, *Adam Osborne et Associates, Inc.* 1979.

(Received February 1, 1987)