



An Information Theoretic Condition for Perfect Reconstruction

Idris Delsol, Olivier Rioul, Julien Béguinot, Victor Rabiet, Antoine Souloumiac

► To cite this version:

Idris Delsol, Olivier Rioul, Julien Béguinot, Victor Rabiet, Antoine Souloumiac. An Information Theoretic Condition for Perfect Reconstruction. *Entropy*, 2024, 26 (1), pp.86. 10.3390/e26010086 . hal-04416099

HAL Id: hal-04416099

<https://telecom-paris.hal.science/hal-04416099>

Submitted on 25 Jan 2024

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

An Information Theoretic Condition for Perfect Reconstruction

Idris Delsol ¹, Olivier Rioul ^{1,*} , Julien Béguinot ¹ , Victor Rabiet ^{1,2} and Antoine Souloumiac ³

¹ Laboratoire de Traitement et Communication de l'Information, Télécom Paris, Institut Polytechnique de Paris, 91120 Palaiseau, France; idris.delsol@telecom-paris.fr (I.D.); julien.beguino@telecom-paris.fr (J.B.); victor.rabiet@telecom-paris.fr or victor.rabiet@ens.fr (V.R.)

² Département de Mathématiques et Applications, École Nationale Supérieure, 75005 Paris, France

³ CEA-List, Université Paris-Saclay, 91120 Palaiseau, France; antoine.souloumiac@cea.fr

* Correspondence: olivier.rioul@telecom-paris.fr

Abstract: A new information theoretic condition is presented for reconstructing a discrete random variable X based on the knowledge of a set of discrete functions of X . The reconstruction condition is derived from Shannon's 1953 lattice theory with two entropic metrics of Shannon and Rajska. Because such a theoretical material is relatively unknown and appears quite dispersed in different references, we first provide a synthetic description (with complete proofs) of its concepts, such as total, common, and complementary information. The definitions and properties of the two entropic metrics are also fully detailed and shown to be compatible with the lattice structure. A new geometric interpretation of such a lattice structure is then investigated, which leads to a necessary (and sometimes sufficient) condition for reconstructing the discrete random variable X given a set $\{X_1, \dots, X_n\}$ of elements in the lattice generated by X . Intuitively, the components $X_1, \dots, X_n$ of the original source of information X should not be globally "too far away" from X in the entropic distance in order that X is reconstructable. In other words, these components should not overall have too low of a dependence on X ; otherwise, reconstruction is impossible. These geometric considerations constitute a starting point for a possible novel "perfect reconstruction theory", which needs to be further investigated and improved along these lines. Finally, this condition is illustrated in five specific examples of perfect reconstruction problems: the reconstruction of a symmetric random variable from the knowledge of its sign and absolute value, the reconstruction of a word from a set of linear combinations, the reconstruction of an integer from its prime signature (fundamental theorem of arithmetic) and from its remainders modulo a set of coprime integers (Chinese remainder theorem), and the reconstruction of the sorting permutation of a list from a minimal set of pairwise comparisons.

Keywords: information lattice; common information; complementary information; Rajska distance; Shannon distance; dependency coefficient; relative redundancy; convex envelope; perfect reconstruction

A movement is accomplished in six stages
And the seventh brings return.
The seven is the number of the young light
It forms when darkness is increased by one.

Change returns success
Going and coming without error.
Action brings good fortune.
Sunset, sunrise.

Syd Barrett, Chapter 24 (Pink Floyd).

1. Introduction

We consider the problem of perfectly reconstructing a discrete random variable X , based on the knowledge of a finite set $X_1, X_2, \dots, X_n$ of deterministic processings or transformations of X , denoted f_i , such that $X_i = f_i(X)$. Intuitively, the components X_i are

assumed to carry only a partial amount of the “information” present in X , and the perfect reconstruction of X would only be possible if the combination of the “information” in $X_1, X_2, \dots, X_n$ is enough to contain all the original “information” in X . Such intuitive considerations expressed in the language of information are very common in signal processing and in many other scientific fields; but, they were never mathematically formalized as far as the authors know. This article aims at formalizing precisely this trivial and vague intuition. Such a task implies, in particular, an accurate definition of “information”.

The Shannon’s 1948 classical information theory [1] cannot really answer this question as it is rather a theory of the measure of information rather than of the information itself. Fortunately, a “true information” theory was also developed by Claude Shannon in a relatively unknown 1953 article [2], which is *not* what is generally referred to as “Shannon’s information theory”. Said briefly, the information is defined there as an equivalence class of discrete random variables. A partial order on a set of classes allows one to build a lattice structure called the *information lattice*, which is made metric by the introduction of two related entropic distances.

“Claude [Shannon] did not like the term ‘information theory’” recalls Robert Fano, a colleague of Shannon’s working at MIT, who died almost a century old just seven years ago. In one of his last interviews [3], he said, “You see, the term ‘information theory’ suggests that it’s a theory about information, but it’s not. It’s about the transmission of information, not the information. Many people just didn’t understand that”. Fano is of course referring to Shannon’s famous theory in his 1948 seminal paper [1], which he entitled, “a mathematical theory of *communication*”—not information. But, very early on, it was the term “information” that prevailed. The entropy $H(X)$ of a discrete random variable X is presented as the measure of “*information contained in X* ”, and the notion of the *mutual information* $I(X; Y)$ between two variables X and Y , introduced precisely by the same Robert Fano in his course at MIT [4], quickly became central to the teaching of the theory. Moreover, the very first historical article on the theory, barely three years after its birth, is entitled “*A history of the theory of information*” [5].

This sudden craze for “information” in the early 1950s eventually became somewhat of a bore for Shannon, who in 1956, in his famous editorial, *The Bandwagon* [6] warned against the excesses of such popularity: “*It will be all too easy for our somewhat artificial prosperity to collapse overnight when we realize that the use of a few exciting words like information, entropy, redundancy, does not solve all our problems*”.

Under these conditions, it is understandable that Shannon wanted to go further: If several, unrelated, random variables can have the same *quantity* of information H , how can information itself be defined? Shannon presented a very brief summary of his findings (without proofs) at the International Congress of Mathematicians (ICM) in 1950 [7] and in a small, relatively unknown article [2] published in 1953 in the very first issue of what was to become the *IEEE Transactions on Information Theory*.

The remainder of this article is organized as follows. Section 2 presents in detail the Shannon theory of the lattice of information with complete proofs, and Section 3 does the same for the two entropic distances proposed, respectively, by Shannon and Rajski. The corresponding geometric point of view is further developed in Section 4. Two conditions of perfect reconstruction, a necessary one and a sufficient one, are then derived in Section 5. Finally, the condition is applied to five specific examples in Section 6.

Sections 2, 3, and 4.1 are a deepening of the article [8] previously published (in French) by four of the authors.

2. What Is Information? A Detailed Study of Shannon’s Information Lattice

For simplicity, we consider with Shannon discrete random variables X , which take a *finite* number of values in some alphabet $\mathcal{X}$. This amounts to considering all the random variables $X : \Omega \rightarrow \mathcal{X}$ defined on a given probability space $(\Omega, \mathcal{P}(\Omega), \mathbb{P})$, where the underlying universe Ω is finite and $\mathcal{P}(\Omega)$ is the power set of Ω .

2.1. Definition of the “True” Information

Quite arguably, the *information* contained in a discrete source or random variable X should not be confused with the “measure of quantity of information” such as the entropy $H(X)$. Shannon’s idea [2] is that this information contained in X should in fact be defined as X itself. Of course, any reversible encoding of X must be regarded as the *same* information, since one moves from one representation to another without loss of information. This amounts, in modern language, to the following definition:

Definition 1 (“True” information). *The information (contained in) X is the equivalence class of X for the equivalence relation:*

$$X \equiv Y \iff Y = f(X) \text{ and } X = g(Y) \text{ a.s. (almost surely)} \quad (1)$$

for two deterministic functions f and g .

Proof. Relation $\equiv$ is evidently reflexive (take f and g to be the identity function) and symmetric (by permuting the roles of f and g in the definition). It is also transitive by composition: if $X \equiv Y$ and $Y \equiv Z$, there exists f, g, h , and k such that $Y = f(X)$, $X = g(Y)$, and $Y = h(Z)$, $Z = k(Y)$ a.s.; then, $X = g(h(Z)) = g \circ h(Z)$ and $Z = k \circ f(X)$ a.s. $\square$

Proposition 1. $X \equiv Y$ if and only if (iff) there exists a bijective function h such that $Y = h(X)$ a.s.

Proof. If $X \equiv Y$, then there exist two deterministic functions f and g such that $X = f(Y)$ and $Y = g(X)$ a.s. Thus, $X = f(g(X))$ a.s. Then, for every value $X = x$ with non-zero probability, $f \circ g(x) = x$. Hence, $f \circ g$ coincides with the identity function a.s. Since the problem is symmetric in X and Y , $g \circ f$ also coincides with the identity function a.s. Thus, $h = g$ is bijective from the set of values that X can take with non-zero probability to the set of values that Y can take with non-zero probability, and we have $Y = g(X) = h(X)$ a.s.

Conversely, if $Y = h(X)$ a.s. with bijective h , then $X = h^{-1}(Y)$ a.s.; hence, $X \equiv Y$. $\square$

As suggested by Rajski [9], the equivalence between X and Y can be characterized by way of their joint probability matrix:

Proposition 2 (Matrix characterization). *If we restrain Ω to the elements of the non-zero probability measure, $X \equiv Y$ iff the matrix of joint probabilities $\mathbb{P}(X = x, Y = y)$ is a permutation matrix.*

Proof. By Proposition 1, $X \equiv Y$ iff there exists a bijective function h such that $Y = h(X)$ a.s. Thus, to each outcome of X corresponds exactly one outcome of Y and vice versa, which is equivalent to saying that the matrix of joint probabilities is a permutation matrix. $\square$

In the following, we shall denote (without possible confusion) X the equivalence class of the variable X , and thus, $X = Y$, the equality between the two classes X and Y (rather than $X \equiv Y$).

With this definition, it is clear that the equivalence relation is compatible with any functional relation $Y = f(X)$. If f is not bijective, it is tempting to say that there is *less* information in Y than in X , hence the following partial order.

Definition 2 (Partial order).

$$X \geq Y \iff Y = f(X) \text{ a.s.} \quad (2)$$

for some deterministic function f .

We also write $Y \leq X$. We are not necessarily considering real-valued variables, so the order $X \geq Y$ has nothing to do with the order in $\mathbb{R}$.

Proposition 3. *The relation $\geq$ is indeed a partial order on the set of equivalence classes of the relation $\equiv$ defined above.*

Proof. We first show that the relation $\equiv$ is compatible with the relation $\geq$. Let X_1, X_2 , and Y_1, Y_2 be such that $X_1 \equiv X_2$ and $Y_1 \equiv Y_2$. Then, if $X_1 \geq Y_1$, there exists a deterministic function f such that $Y_1 = f(X_1)$ a.s. Since $X_1 \equiv X_2$, there exists a bijective h such that $X_1 = h(X_2)$ a.s.; hence, $Y_1 = f \circ h(X_2)$ a.s. and $X_2 \geq Y_1$. Likewise, since $Y_1 \equiv Y_2$, there exists a bijective g such that $Y_2 = g(Y_1)$ a.s., so $Y_2 = g \circ f \circ h(X_2)$ a.s.; hence, $X_2 \geq Y_2$. This shows that the relation $\geq$ is well defined on the set of equivalence classes of the relation $\equiv$.

We now show that $\geq$ is indeed a partial order:

- *Reflexivity:* $X = \text{Id}(X)$ so $X \geq X$.
- *Antisymmetry:* If $X \geq Y$ and $Y \geq X$, $X = f(Y)$ a.s., and $Y = g(X)$ a.s. for deterministic functions f and g , so $X \equiv Y$.
- *Transitivity:* If $X \geq Y$ and $Y \geq Z$, then there exist two deterministic functions f and g such that: $Z = g(Y)$ a.s. and $Y = f(X)$ a.s. Then, $Z = g(f(X))$ a.s.; hence, $X \geq Z$. $\square$

2.2. Structure of the Information Lattice: Joint Information; Common Information

Beyond the partial order, Shannon [2] established the natural mathematical structure of information: it is a *lattice*, i.e., two variables X, Y always admit a maximum $X \vee Y$ and a minimum $X \wedge Y$. Let us recall that these quantities (necessarily unique if they exist) are defined by the relations:

$$\begin{aligned} (X \leq Z \text{ and } Y \leq Z) &\iff X \vee Y \leq Z, \\ (X \geq Z \text{ and } Y \geq Z) &\iff X \wedge Y \geq Z. \end{aligned} \quad (3)$$

Shannon, in his paper [2], used Boolean notations instead, $X + Y$ for $X \vee Y$ and $X \cdot Y$ for $X \wedge Y$.

Proposition 4 (Joint information). *The joint information $X \vee Y$ of X and Y is the random pair $X \vee Y = (X, Y)$.*

Proof. If X and Y are functions of Z , then the pair (X, Y) is also a function of Z . Conversely, since X and Y are functions of (X, Y) , if (X, Y) is a function of Z , then so are X and Y . $\square$

The definition of $X \wedge Y$ (*common information*) is more difficult and was not made explicit by Shannon. Following Gács and Körner [10], let us adopt the following definition:

Definition 3. *We say that $x \in \mathcal{X}$ and $y \in \mathcal{Y}$ communicate, denoted by $x \sim y$, if there exists a path $xy_1x_1y_2 \cdots y_nx_ny$ in which all transitions are of non-zero probability: $\mathbb{P}(X = x, Y = y_1) > 0$, $\mathbb{P}(Y = y_1, X = x_1) > 0, \dots, \mathbb{P}(X = x_n, Y = y) > 0$.*

For convenience, we also write $y \sim x$ when x and y communicate. Strictly speaking, the relation $x \sim y$ is not an equivalence relation because x and y do not belong to the same set. However, it has similar properties:

Proposition 5. *The relation $\sim$ on the set of pairs (x, y) for which $\mathbb{P}(X = x) > 0$ and $\mathbb{P}(Y = y) > 0$ is transitive in the sense that $x_1 \sim y_1, y_1 \sim x_2$, and $x_2 \sim y_2$ implies $x_1 \sim y_2$.*

Proof. If $x_1 \sim y_1, y_1 \sim x_2$, and $x_2 \sim y_2$, then there exists a path from x_1 to y_1 , another from y_1 to x_2 , and a least one from x_2 to y_2 , whose transitions are of non-zero probability. The concatenated path from x_1 to y_2 has non-zero transition probabilities; hence, $x_1 \sim y_2$. $\square$

Definition 4 (Communication class). *If $x \sim y$, we define the communication class $C(x, y)$ as the set of all (x', y') such that $x' \sim y'$ and $x \sim y'$.*

Thus, by transitivity, $C(x, y) = C(x', y')$ for all (x', y') in the communication class, so that two classes are either equal or distinct. Therefore, the distinct communication classes partition the set of all values (x, y) for which $\mathbb{P}(X = x) > 0$ and $\mathbb{P}(Y = y) > 0$. We may identify any communication class C with its characteristic function $1_{(x,y) \in C}$ so that $C(X, Y)$ is a binary random variable.

Proposition 6 (Common information). *The common information $X \wedge Y$ of X and Y is $X \wedge Y = C(X, Y)$.*

Proof. If $Z = f(X) = g(Y)$ a.s., then Z is constant for each pair (x, y) such that $x \sim y$; in other words, Z is a function of the class $C(X, Y)$. $\square$

Remark 1. *In order to compute the common information between X and Y in practice, one has to fully determine the communication classes, which is only possible if there is a finite number of classes, each of which contains a finite number of elements. In other words, X and Y should take a finite number of values. This is the reason why we restrict ourselves to finitely valued variables in this paper.*

Remark 2. *As in any lattice, $X \leq Y$ is equivalent to saying that $X \vee Y = Y$ or that $X \wedge Y = X$.*

2.3. Computing Common Information

As shown in the previous section, the definition of common information is not a simple one, but one can compute it efficiently using the following algorithm. Given two variables X and Y , this algorithm turns the joint probability matrix of (X, Y) into a *block-diagonal* matrix, where each block corresponds to each communication class.

Let X and Y be two random variables taking values in $\mathcal{X}$ and $\mathcal{Y}$, respectively. Consider the graph $G = (V, E)$ whose vertices V are $\mathcal{X} \cup \mathcal{Y}$ and such that the vertices x and y of V are connected by an edge if and only if $\mathbb{P}(X = x, Y = y) > 0$. Hence, G is fully described by the joint probability matrix $\mathbb{P}_{X,Y}$. Furthermore, this is a bipartite graph (no edge connects two vertices x_1 and x_2 belonging to $\mathcal{X}$ or two vertices y_1 and y_2 belonging to $\mathcal{Y}$).

Then, the communication classes $C(X, Y)$ correspond to the *connected components* of G . Indeed, a connected component C is a subset of V such that each of its elements is accessible to all the others by a path in the subgraph (C, E) . So, for any two vertices x, y in the connected component C , there exists $y_1, x_1, \dots, y_k, x_k$ such that all the edges $(x, y_1), (y_1, x_1), \dots, (y_k, x_k), (x_k, y)$ belong to E , that is all the transition probabilities between these vertices are non-zero, which is equivalent to saying that they belong to the same communication class. Now, it is known that the connected components of G can be determined by a depth-first search.

We propose an algorithm, whose pseudo-code is given in Algorithm 1, that takes as the input the joint probability matrix $\mathbb{P}_{X,Y}$ and outputs a block-diagonal form of $\mathbb{P}_{X,Y}$ representing the common information $X \wedge Y$, an array storing the permutation of the columns of $\mathbb{P}_{X,Y}$, and an array storing the permutation of the rows $\mathbb{P}_{X,Y}$. Since the matrix $\mathbb{P}_{X,Y}$ is sufficient to fully describe G , we adapt the depth-first search algorithm to browse the rows and columns of the matrix $\mathbb{P}_{X,Y}$ to find which of its rows and columns must be swapped in order to write this matrix in a block-diagonal form. In this algorithm (Algorithm 1), the i th row of $\mathbb{P}_{X,Y}$ will be represented by the pair (r, i) and the j th column by the pair (c, j) .

Algorithm 1 Algorithm to compute the common information.

```

1: input  $\mathbb{P}_{X,Y}$ :  $n_R \times n_C$  matrix ▷ Joint probability matrix
2:  $\sigma_R \leftarrow$  array of integers of length  $n_R$  ▷ Rows' permutation vector
3:  $\sigma_C \leftarrow$  array of integers of length  $n_C$  ▷ Columns' permutation vector
4:  $S \leftarrow$  empty stack ▷ Stack contains row indices  $(r, i)$  or column indices  $(c, j)$ 
5: push  $(r, 0)$  into stack  $S$  ▷ First row put into stack
6:  $bottom \leftarrow 1$  ▷ Bottommost row index not yet assigned
7:  $up \leftarrow n_R - 1$  ▷ Uppermost row index that may have non-zero entries
8:  $left \leftarrow 0$  ▷ Leftmost column index not yet assigned
9:  $right \leftarrow n_C - 1$  ▷ Rightmost column index that may have non-zero entries
10: while There is an unmarked row or column do
11:   while  $S$  is not empty do
12:      $(s, i) \leftarrow S.pop()$  ▷ The  $pop()$  operation removes the top stack element and returns it.
13:     if  $(s, i)$  is not marked then
14:       mark  $(s, i)$ 
15:       if  $s = r$  then ▷ Current index  $i$  is a row index
16:         for  $left \leq j \leq right$  do ▷ Scan all columns
17:           if  $\mathbb{P}_{X,Y}(i, j) > 0$  then
18:             push  $(c, j)$  into stack  $S$ 
19:              $\sigma_C[j] \leftarrow left$ ; swap columns  $left$  and  $j$  in  $\mathbb{P}_{X,Y}$ 
20:              $left \leftarrow left + 1$ 
21:           end if
22:         end for
23:         if all entries on  $i$ th row are zeros then
24:            $\sigma_R[i] \leftarrow up$ ; swap rows  $i$  and  $up$  in  $\mathbb{P}_{X,Y}$ 
25:            $up \leftarrow up - 1$ 
26:         end if
27:       else ▷ Current index  $i$  is a column index
28:         for  $bottom \leq j \leq up$  do ▷ Scan all rows
29:           if  $\mathbb{P}_{X,Y}(j, i) > 0$  then
30:             push  $(r, j)$  into stack  $S$ 
31:              $\sigma_R[j] \leftarrow bottom$ ; swap rows  $bottom$  and  $j$  in  $\mathbb{P}_{X,Y}$ 
32:              $bottom \leftarrow bottom + 1$ 
33:           end if
34:         end for
35:         if all entries on  $i$ th column are zeros then
36:            $\sigma_C[i] \leftarrow right$ ; swap columns  $i$  and  $right$  in  $\mathbb{P}_{X,Y}$ 
37:            $right \leftarrow right - 1$ 
38:         end if
39:       end if
40:     end if
41:   end while ▷ Empty Stack
42:   if there is an unmarked  $i$ th row  $(r, i)$  then
43:     push  $(r, i)$  into stack  $S$ 
44:   else if there is an unmarked  $j$ th column  $(c, j)$  then
45:     push  $(c, j)$  into stack  $S$ 
46:   end if
47: end while ▷ All rows and columns marked
48: return  $\mathbb{P}_{X,Y}, \sigma_R, \sigma_C$ 

```

The complexity of this algorithm can be determined as follows. Let $n = \text{Card}(\mathcal{X}) + \text{Card}(\mathcal{Y})$ be the sum of the alphabet sizes on which X and Y take their values, i.e., the sum of the number of rows of $\mathbb{P}_{X,Y}$ and the number of columns of $\mathbb{P}_{X,Y}$. The algorithm passes through each row and column at most once. Indeed, for the index of a row or column to enter the stack, it must be *unmarked*, but as soon as we put it on the stack, we mark it. Then,

each time the index of a row or a column is unstacked, we look at each coefficient of the corresponding row or column. Therefore, our algorithm looks at each coefficient of the joint probability matrix $\mathbb{P}_{X,Y}$ exactly once. Four elementary operations are performed each time we cross a non-zero coefficient. Thus, the algorithm's complexity is quadratic in n .

Notice that the output of our algorithm gives a visualization of the common information: The stochastic matrix $\mathbb{P}(X = x, Y = y)$ is written, after the permutation of the rows/columns, in the “block-diagonal” form

$$\mathbb{P}_{X,Y} = \begin{pmatrix} C_1 & & & & & \\ & C_2 & & & & \\ & & \ddots & & & \\ & & & C_k & & \\ \mathbf{0} & & & & 0 & \\ & & & & & \ddots \\ & & & & & & 0 \end{pmatrix} \quad (4)$$

where k , the number of blocks, is maximal. The k rectangular matrices then represent the k different equivalence classes, the probability $\mathbb{P}(C(X, Y) = i)$ being the sum of all entries in block C_i .

2.4. Boundedness and Complementedness: Null, Total, and Complementary Information

Proposition 7 (Null information; total information). *The information lattice is bounded, i.e., it admits a minimum of 0 and a maximum of 1, such that, for any X , $0 \leq X \leq 1$:*

- The minimal element 0 (“null information”) is the equivalence class of all deterministic variables. Thus, $X = 0$ means that X is a deterministic variable.
- The maximal element 1 (“total information”) of the lattice is the equivalence class of the identity function Id on Ω .

Proof. If X is any random variable and $Z = c$ a.s. is any deterministic variable, then it is clear that $Z = f(X)$, where f is the constant function c . Letting 0 be the equivalence class of constant variables, one has $0 \leq X$ for all X .

Also, for any random variable X , $X = X \circ \text{Id}$; hence, $X \leq \text{Id}$. Letting 1 be the equivalence class of the identity function on Ω , one has $X \leq 1$ for all X . $\square$

Proposition 8 (Complementary information). *The information lattice is complemented, i.e., any $X \leq Y$ admits a complement Z (“complementary information”) such that $X \vee Z = Y$ and $X \wedge Z = 0$.*

This Z is the information missing from X to obtain Y : It allows Y to be reconstructed from X without requiring more information than necessary. Shannon in [2] did not say how to determine it. The following proof gives an explicit construction:

Proof. Since $X \leq Y$, we simply have $X = X \wedge Y = C(X, Y)$. Thus, a given class $C(X, Y) = x$ has only one value $X = x$ per class, corresponding in general to several values of Y , say $y_1^x, y_2^x, \dots, y_{k_x}^x$. Now, let $Z \in \{1, \dots, k_X\}$ be the unique index such that $Y = Y_Z^X$.

By construction, $Z \leq X \vee Y = Y$, and since $X \leq Y$, one also has $X \vee Z \leq Y$. But, the formula $Y = Y_Z^X$ shows that $Y \leq X \vee Z$; hence, the equality $X \vee Z = Y$ holds.

Finally, the value $Z = 1$ connects each pair (x, z) , so there is only one class according to (X, Z) , i.e., $X \wedge Z = 0$. $\square$

This construction can be visualized on the stochastic tensor of (X, Y, Z) described in Figure 1.

Remark 3. The complementary information Z is not uniquely determined by X and Y . In the above construction, it depends on how the values of Y are indexed by the class $X = x$.

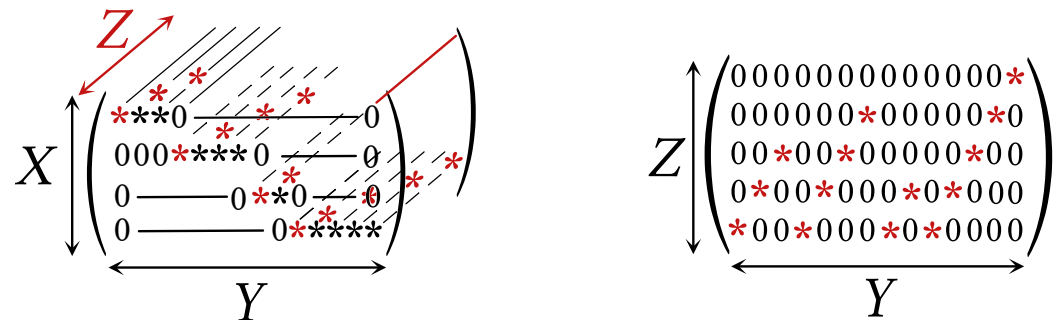


Figure 1. Construction of the complementary information Z allowing passing from X to Y . The stochastic tensor of (X, Y, Z) representing $\mathbb{P}_{X,Y,Z}$ has non-zero entries marked in red. The distribution $\mathbb{P}_Z$ of Z is obtained by marginalizing the tensor on the Z -axis.

2.5. Computing the Complementary Information

Given $X \leq Y$, Algorithm 2 determines a random variable Z corresponding to the complementary information from X to Y . This algorithm takes as the input the joint probability matrix $\mathbb{P}_{X,Y}$ in its block-diagonal form and outputs the tensor of the joint probability $\mathbb{P}_{X,Y,Z}$, where $X \vee Z = Y$ and $X \wedge Z = 0$. The tensor is built by spreading the non-zero coefficients of the joint probability matrix $\mathbb{P}_{X,Y}$ on the Z -axis as shown in Figure 1.

Algorithm 2 Algorithm for computing the complementary information.

```

1: input  $\mathbb{P}_{X,Y}$ :  $n_R \times n_C$  matrix                                ▷ Joint probability matrix
2:  $k \leftarrow 0$                                                     ▷ Z index
3: for  $0 \leq i < nR$  do
4:   for  $0 \leq j < nC$  do
5:     if  $\mathbb{P}_{X,Y}(i, j) > 0$  then
6:        $\mathbb{P}_{X,Y,Z}(i, j, k) \leftarrow \mathbb{P}_{X,Y}(i, j)$ 
7:        $k \leftarrow k + 1$ 
8:     end if
9:   end for
10:   $k \leftarrow 0$ 
11: end for
12: return  $\mathbb{P}_{X,Y,Z}$ 

```

The algorithm looks at each coefficient of the joint probability matrix $\mathbb{P}_{X,Y}$ exactly once and performs at most two elementary operations for each coefficient it processes. Therefore, it is quadratic in $n = \text{Card}(\mathcal{X}) + \text{Card}(\mathcal{Y})$ (since the number of coefficients in the matrix $\mathbb{P}_{X,Y}$ is quadratic in n).

2.6. Relationship Between Complementary Information and Functional Representation

There is a striking resemblance between Proposition 8 and the “functional representation lemma”, which has been used in recent years in various applications of information theory for network coding (see Appendix B, pp. 626–627, of [11]).

For the convenience of the notations, we write $X \perp Y$ if $X \wedge Y = 0$ (null common information) and write $X \perp\!\!\!\perp Y$ iff X and Y are independent. It is easily seen (see Remark 4 below) that $X \perp\!\!\!\perp Y \implies X \perp Y$. Now, Proposition 8 and the “functional representation lemma” can be rewritten as follows.

Lemma 1 (Complementary information lemma (Proposition 8)).

$$\forall X \leq Y, \quad \exists Z \perp X \quad \text{s.t.} \quad Y = X \vee Z. \quad (5)$$

Lemma 2 (Functional representation lemma ([11])).

$$\forall X, Y, \quad \exists Z \perp\!\!\!\perp X \quad \text{s.t.} \quad Y \leq X \vee Z. \quad (6)$$

Thus, compared to the “complementary information lemma”, the “functional representation lemma” (i) has a general assumption of X and Y (X need not be a function of Y), but (ii) requires a stronger condition on Z ($Z \perp\!\!\!\perp X$ instead of $Z \perp X$) and (iii) has a weaker conclusion (Y is only a function of X and Z). It would be interesting to further investigate the relationship between these two lemmas since it is apparent that one lemma cannot be deduced from the other.

2.7. Is the Information Lattice a Boolean Algebra?

Interestingly, it was Shannon who, as early as 1938 in his master’s thesis, used the *Boolean algebra* to study relay-based circuits—“the most important master’s thesis of the century” for which Shannon received the Alfred Noble prize (not to be confused with the Alfred Nobel Prize) in 1940. But, alas, as Shannon noted, his information lattice is *not* a Boolean algebra. It would have been one if it were *distributive* ($\wedge$ distributive with respect to $\vee$ or vice versa), since a Boolean algebra is, by definition, a distributive complemented bounded lattice. However:

Proposition 9. *The information lattice is not distributive.*

Indirect proof. In any Boolean algebra, the complement is unique. As seen above, this is not the case for the information lattice. $\square$

Direct proof. As a direct second proof, we provide an explicit counterexample to distributivity. Consider the probability space $(\Omega, \mathcal{P}(\Omega), \mathbb{P})$, where $\Omega = \{0, 1, 2, 3\}$ and $\mathbb{P}$ is the uniform probability measure, and define $X(\omega) = 0$ if ω is even, $X(\omega) = 1$ otherwise. Now, let Z_1, Z_2 be given as in Table 1 below. As we read in the table, $(X \wedge Z_1) \vee (X \wedge Z_2) = 0$ is constant, while $X \wedge (Z_1 \vee Z_2)$ is not. Therefore, $(X \wedge Z_1) \vee (X \wedge Z_2) \neq X \wedge (Z_1 \vee Z_2)$, and the information lattice is not distributive. $\square$

Table 1. Computation of $X \wedge (Z_1 \vee Z_2)$ and of $(X \wedge Z_1) \vee (X \wedge Z_2)$.

ω	0	1	2	3
X	0	1	0	1
Z_1	1	1	2	2
Z_2	2	1	1	2
$Z_1 \vee Z_2$	(1, 2)	(1, 1)	(2, 1)	(2, 2)
$X \wedge (Z_1 \vee Z_2)$	0	1	0	1
$X \wedge Z_1$	0	0	0	0
$X \wedge Z_2$	0	0	0	0
$(X \wedge Z_1) \vee (X \wedge Z_2)$	(0, 0)	(0, 0)	(0, 0)	(0, 0)

3. Metric Properties of the Information Lattice

3.1. Information and Information Measures

First of all, it is immediate to check the compatibility of the information lattice with respect to the entropy or the mutual information as logarithmic measures of information.

We use the following standard notations. The entropy of X is denoted $H(X)$. If X takes values in $\mathcal{X}$ of size N , then the entropy of X satisfies $H(X) \leq \log N$ with equality iff X is uniformly distributed. Here, “log” refers to the logarithm taken to *any* base. The conditional entropy of X given Y is denoted $H(X|Y)$, and the mutual information between X and Y is denoted $I(X; Y)$.

Proposition 10. Entropy, conditional entropy, and mutual information are compatible with the definition of information as an equivalence class:

Proof.

- *Entropy:* If $X \equiv Y$, there exist functions f and g such that $Y = f(X)$ a.s. (hence, $H(Y) \leq H(X)$) and $X = g(Y)$ a.s. (hence, $H(X) \leq H(Y)$). Thus, $H(X) = H(Y)$.
- *Conditional entropy:* Let $X_1 \equiv X_2$ with f and g be two functions such that $X_1 = f(X_2)$ and $X_2 = g(X_1)$ a.s. Then, $H(X_1|Y) = H(f(X_2)|Y) \leq H(X_2|Y)$. Similarly, $H(X_2|Y) = H(g(X_1)|Y) \leq H(X_1|Y)$. Therefore, $H(X_1|Y) = H(X_2|Y)$. Finally, if $Y_1 \equiv Y_2$ with two functions h and k such that $Y_1 = h(Y_2)$ and $Y_2 = k(Y_1)$ a.s., then $H(X|Y_1) = H(X|h(Y_2)) \geq H(X|Y_2, h(Y_2)) = H(X|Y_2)$ and likewise $H(X|Y_2) = H(X|k(Y_1)) \geq H(X|Y_1, k(Y_1)) = H(X|Y_1)$. Therefore, $H(X|Y_1) = H(X|Y_2)$.
- *Mutual information:* Since $I(X; Y) = H(X) - H(X|Y)$, compatibility follows from the two previous cases. $\square$

We then have some obvious connections:

Proposition 11 (Partial order and conditional entropy).

$$X \leq Y \iff H(X|Y) = 0 \quad (7)$$

In particular, H is “order-preserving” (greater information implies higher entropy):

$$X \leq Y \implies H(X) \leq H(Y). \quad (8)$$

Also, $X \leq Y$ with $H(X) = H(Y)$ implies $X = Y$.

Finally, $H(X) \geq 0$ for all X , with equality $H(X) = 0$ iff $X = 0$.

Proof. $H(X|Y) = 0$ means that $H(X|Y = y) = 0$ for all $y \in \mathcal{Y}$, which amounts to saying that X is deterministic equal to $f(y)$ given $Y = y$. In other words, $X = f(Y)$ a.s. We then have $H(X) = H(X) - H(X|Y) = I(X; Y) = H(Y) - H(Y|X) \leq H(Y)$.

Next, suppose that $X \leq Y$ and $H(X) = H(Y)$. By the chain rule, $H(X, Y) = H(X) + H(Y|X) = H(Y) + H(X|Y)$. Therefore, it follows from the equality $H(X) = H(Y)$ that $H(Y|X) = H(X|Y)$. But, since $X \leq Y$, $H(X|Y) = 0$; hence, $H(Y|X) = 0$ also, that is $Y \leq X$. This shows equivalence $Y = X$.

Finally, since $X \geq 0$ for all X , $H(X) \geq H(0) = 0$, and it is well known that the entropy $H(X)$ is zero if and only if the variable X is deterministic, that is $X = 0$. $\square$

An example of an information lattice with associated entropies is given in Figure 2.

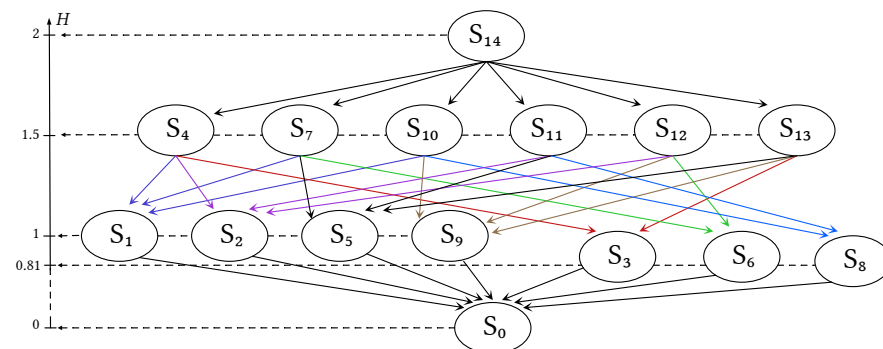


Figure 2. Hasse diagram of the information lattice defined on a universe Ω of size 4 with uniform probability. There are 15 different random variables corresponding to the 15 different ways to partition Ω . The corresponding entropies are, in descending order: 2; 1.5; 1; ≈ 0.81 , and 0 bits.

3.2. Common Information vs. Mutual Information

Proposition 12. The entropy of the joint information is the joint entropy, i.e., $H(X \vee Y) = H(X, Y)$.

Proof. Obvious, since $X \vee Y = (X, Y)$. $\square$

One may wonder by analogy with the usual Venn diagram in information theory (Figure 3) if the entropy of joint information is equal to the mutual information: Is it true that $H(X \wedge Y) = I(X; Y)$? The answer is *no*, as shown next. Proposition 13 is implicit in [10] and made explicit by Wyner in [12], who credits a private communication from Kaplan.

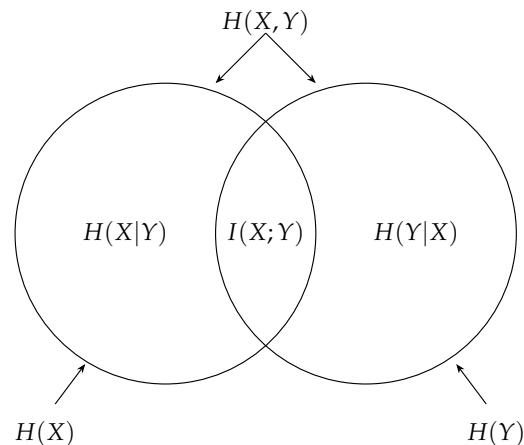


Figure 3. Usual Venn diagram in information theory.

Proposition 13. $H(X \wedge Y) \leq I(X; Y)$ always, with equality $H(X \wedge Y) = I(X; Y)$ iff one can write $X = (U, W)$ and $Y = (V, W)$, where U and V are conditionally independent given W .

Proof. Let $W = X \wedge Y$. Since $W \leq X$ and $W \leq Y$, by complementarity, we can write $X = W \vee U = (U, W)$ and $Y = W \vee V = (V, W)$. By the chain rule for mutual information, $I(X; Y) = I(U, W; V, W) = I(W; V, W) + I(U; V, W|W) = H(W) + I(U; V|W) \geq H(W)$ with equality iff U and V are conditionally independent given W . $\square$

Remark 4. In particular, if X and Y are independent, they have null common information $X \wedge Y = 0$. However, common information $H(X \wedge Y)$ can be far less [10] than mutual information $I(X; Y)$.

Remark 5. Notice that the case of equality corresponds to the case where the matrix blocks C_i in (4) are stochastic matrices of two independent variables X, Y knowing $W = i$, i.e., matrices of rank one.

Remark 6. Shannon's notion of common information should not be confused with the well-known Wyner's accepting of "common information", which is defined as the maximum of $I(X, Y; W)$ when X and Y are conditionally independent knowing W . This quantity is not less, but greater than the mutual information $I(X; Y)$ [12].

3.3. Submodularity of Entropy on the Information Lattice

From the results in [13], we can show that entropy is *submodular* on the information lattice:

Proposition 14 (Submodularity of entropy). $H(X \vee Y) + H(X \wedge Y) \leq H(X) + H(Y)$.

Proof. Since $X \wedge Y \leq Y$, $H(Y) = H(Y, X \wedge Y) = H(X \wedge Y) + H(Y|X \wedge Y)$. But, since $X \wedge Y \leq X$, $H(Y|X \wedge Y) \geq H(Y|X \wedge Y, X) = H(Y|X) = H(X \vee Y) - H(X)$. Combining gives the announced inequality. $\square$

Remark 7. Submodularity is in fact equivalent to the inequality $H(X \wedge Y) \leq I(X; Y)$ of Proposition 13, since $H(X) + H(Y) - H(X \vee Y) = H(X) + H(Y) - H(X, Y) = I(X; Y)$.

Remark 8. The submodularity property of entropy that is generally studied in the information theory literature is with respect to the set lattice (or algebra), where the entropy is that of a collection of random variables indexed by some index set (thus considered as a set function). Such considerations have been greatly developed in recent years; see, e.g., [14]. By contrast, it is the information lattice that is considered here. It can be easily shown using Proposition 13 that the two notions of submodularity coincide for collections of independent random variables.

3.4. Two Entropic Metrics: Shannon Distance; Rajski Distance

Since $X = Y \iff (X \leq Y \text{ and } X \geq Y)$, according to Proposition 11, it suffices that $H(X|Y) + H(Y|X) = 0$ in order for X and Y to be equivalent: $X = Y$. Shannon [2] noted that this defines a *distance*, which makes the information lattice a *metric space*:

Proposition 15 (Shannon's entropic distance). $D(X, Y) = H(X|Y) + H(Y|X)$ is a distance over the information lattice:

Proof.

- *Positivity:* As just noted above, $D(X, Y) \geq 0$ vanishes only when $X = Y$.
- *Symmetry:* $D(X, Y) = D(Y, X)$ is obvious by the commutativity of addition.
- *Triangular inequality:* First note that $H(X|Z) \leq H(X, Y|Z) = H(X|Y, Z) + H(Y|Z) \leq H(X|Y) + H(Y|Z)$. By permuting X and Z , we also obtain that $H(Z|X) \leq H(Z|Y) + H(Y|X)$. Summing up the two inequalities, we obtain the triangular inequality $D(X, Z) = H(X|Z) + H(Z|X) \leq H(X|Y) + H(Y|X) + H(Y|Z) + H(Z|Y) = D(X, Y) + D(Y, Z)$. $\square$

It is interesting to note that this is not the only distance (nor the only topology). By normalizing $D(X, Y)$ by the joint entropy $H(X, Y)$, we obtain another distance metric:

Proposition 16 (Rajski's entropic distance [9]). $d(X, Y) = \frac{D(X, Y)}{H(X, Y)}$ (with the convention $d(0, 0) = 0$) is a distance taking values in $[0, 1]$.

Notice that normalization by $H(X, Y)$ is valid when X and Y are non-deterministic since $X \neq 0$ and $Y \neq 0$ implies $H(X, Y) > 0$.

Proof. First of all, symmetry $d(X, Y) = d(Y, X)$ is obvious, and positivity follows from that of D . We follow Horibe [15] to prove the triangular inequality. One may always assume non-deterministic random variables. Observe that:

$$\frac{H(X|Y)}{H(X, Y)} = \frac{H(X|Y)}{H(X|Y) + H(Y)} \geq \frac{H(X|Y)}{H(X|Y) + H(Y, Z)} = \frac{H(X|Y)}{H(X|Y) + H(Y|Z) + H(Z)} \quad (9)$$

and

$$\frac{H(Y|Z)}{H(Y, Z)} = \frac{H(Y|Z)}{H(Y|Z) + H(Z)} \geq \frac{H(Y|Z)}{H(X|Y) + H(Y|Z) + H(Z)}. \quad (10)$$

Summing (9) and (10) yields

$$\frac{H(X|Y)}{H(X, Y)} + \frac{H(Y|Z)}{H(Y, Z)} \geq \frac{H(X|Y) + H(Y|Z)}{H(X|Y) + H(Y|Z) + H(Z)}. \quad (11)$$

Now, from the above proof of the triangular inequality of D , one has $H(X|Y) + H(Y|Z) \geq H(X|Z)$. Noting that $a \geq b > 0$ and $c \geq 0$ imply $\frac{a}{a+c} \geq \frac{b}{b+c}$, we obtain

$$\frac{H(X|Y) + H(Y|Z)}{H(X|Y) + H(Y|Z) + H(Z)} \geq \frac{H(X|Z)}{H(X|Z) + H(Z)} = \frac{H(X|Z)}{H(X, Z)}. \quad (12)$$

Therefore,

$$\frac{H(X|Y)}{H(X,Y)} + \frac{H(Y|Z)}{H(Y,Z)} \geq \frac{H(X|Z)}{H(X,Z)}. \quad (13)$$

Permuting the roles of X and Z gives

$$\frac{H(Y|X)}{H(X,Y)} + \frac{H(Z|Y)}{H(Y,Z)} \geq \frac{H(Z|X)}{H(X,Z)}. \quad (14)$$

Summing (13) and (14), we conclude that $d(X,Y) + d(Y,Z) \geq d(X,Z)$. $\square$

Remark 9. Rajski's distance between two variables X and Y can be visualized as the Jaccard distance between the region corresponding to X and the region corresponding to Y in the Venn diagram of Figure 3. The Jaccard (or Jaccard–Tanimoto) distance [16] between two sets A and B is defined by $d_J(A,B) = \frac{|A \Delta B|}{|A \cup B|}$, where Δ is the symmetric difference between A and B . Thus, if A and B are, respectively, the regions corresponding to X and to Y in the Venn diagram, we have: $H(X,Y) = |A \cup B|$, $H(X|Y) = |A \setminus B|$ and $H(Y|X) = |B \setminus A|$. Thus, $\frac{H(X|Y) + H(Y|X)}{H(X,Y)} = \frac{|(A \setminus B) \cup (B \setminus A)|}{|A \cup B|} = \frac{|A \Delta B|}{|A \cup B|}$.

3.5. Dependency Coefficient

From the Rajski distance, we can define a quantity that measures the dependence between two non-deterministic (i.e., non-zero) random variables X and Y .

Definition 5 (Dependency coefficient). For all non-zero elements X and Y of the information lattice, their dependency coefficient is $\rho(X,Y) = 1 - d(X,Y) \in [0,1]$.

Proposition 17. The dependency coefficient can be seen as normalized mutual information: $\rho(X,Y) = \frac{I(X;Y)}{H(X,Y)}$.

Proof. One has $1 - d(X,Y) = 1 - \frac{H(X|Y) + H(Y|X)}{H(X,Y)} = \frac{H(X,Y) - H(X|Y) - H(Y|X)}{H(X,Y)}$, where the numerator $= H(X) + H(Y|X) - H(X|Y) - H(Y|X) = I(X;Y)$. $\square$

Proposition 18. For non-deterministic X and Y , one has $0 \leq \rho(X,Y) \leq 1$, where $\rho(X,Y) = 0$ vanishes (equivalently, $d(X,Y) = 1$) iff X and Y are independent and $\rho(X,Y) = 1$ (equivalently, $d(X,Y) = 0$) iff $X = Y$ are equivalent.

Proof. If X and Y are independent, then $I(X;Y) = 0$; hence, $\rho(X,Y) = 0$. If X and Y are equivalent, then $d(X,Y) = 0$ and $\rho(X,Y) = 1 - d(X,Y) = 1$. Since $0 \leq I(X;Y) \leq H(X) \leq H(X,Y)$, $0 \leq \rho(X,Y) = \frac{I(X;Y)}{H(X,Y)} \leq 1$. $\square$

Remark 10. The property of ρ in Proposition 18 is similar to the usual property of the linear correlation coefficient. However, while two independent random variables have zero correlation (but not conversely), the corresponding converse property holds for the dependence coefficient since two random variables are independent if and only if $\rho(X,Y) = 0$.

3.6. Discontinuity and Continuity Properties

Perhaps the biggest flaw in Shannon's lattice information theory [2] is that the different constructions of elements in the lattice (e.g., common and complementary information) do not actually depend on the values of the probabilities involved, but only on whether they are equal to or different from zero. Thus, a small perturbation on the probabilities can greatly influence the results. As an illustration, we have the following.

Proposition 19 (Discontinuity of common information). *The application $(X, Y) \mapsto X \wedge Y$ is discontinuous in the metric lattice with distance D (or d).*

Proof. Let $(X_\varepsilon, Y_\varepsilon)$ be defined by the stochastic matrix:

$$\mathbb{P}_{X,Y} = \begin{pmatrix} \frac{1-\varepsilon}{N} & \frac{\varepsilon}{N} & 0 & \cdots & 0 \\ 0 & \frac{1-\varepsilon}{N} & \frac{\varepsilon}{N} & \cdots & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ \frac{\varepsilon}{N} & 0 & \cdots & 0 & \frac{1-\varepsilon}{N} \end{pmatrix}. \quad (15)$$

Since there is a single class of communication, common information $X_\varepsilon \wedge Y_\varepsilon = 0$ is zero for every $\varepsilon > 0$. By contrast, when $\varepsilon = 0$, $X_\varepsilon \wedge Y_\varepsilon$ is uniformly distributed among N communication classes. Consequently, $D(X_\varepsilon \wedge Y_\varepsilon, 0) = 0$ for any $\varepsilon > 0$, whereas $D(X_0 \wedge Y_0, 0) = H(X_0 \wedge Y_0) = \log N$ is arbitrarily large for $\varepsilon = 0$. $\square$

However, it should be noted that the joint information $\vee$ is continuous with respect to Shannon's distance. In fact, we have the following.

Proposition 20. *For any X, X', Y , and Y' ,*

$$D(X \vee Y, X' \vee Y') \leq D(X, X') + D(Y, Y'). \quad (16)$$

Proof. One has

$$\begin{aligned} H(X \vee Y | X' \vee Y') &= H(X, Y | X', Y') \\ &\stackrel{(a)}{=} H(X | X', Y') + H(Y | X', Y', X) \\ &\stackrel{(b)}{\leq} H(X | X') + H(Y | Y'). \end{aligned} \quad (17)$$

where (a) is the consequence of the chain rule and (b) is due to the fact that conditioning reduces entropy. Since X, X' and Y, Y' play a symmetrical role in (b), we can permute the roles of X, X' and Y, Y' , which gives $H(X' \vee Y' | X \vee Y) \leq H(X' | X) + H(Y' | Y)$. Summing both inequalities yields the result. $\square$

Remark 11. *In particular, for $X = X'$, for any X, Y, Z ,*

$$D(X \vee Y, X \vee Z) \leq D(Y, Z). \quad (18)$$

In other words, joining the same X can only reduce the Shannon distance: in this respect, the joining operator $Y \mapsto X \vee Y$ is a contraction operator.

Furthermore, the entropy, the conditional entropy, and the mutual information are continuous with respect to the entropic distance of Shannon. Indeed, we have the following inequalities (see Problem 3.5 in [17]):

Proposition 21. *For all X, Y, X' , and Y' :*

- (i) $|H(X) - H(Y)| \leq D(X, Y)$.
- (ii) $|H(X, Y) - H(X', Y')| \leq D(X, X') + D(Y, Y')$.
- (iii) $|H(X | Y) - H(X' | Y')| \leq D(X, X') + 2D(Y, Y')$.
- (iv) $|I(X; Y) - I(X'; Y')| \leq 2(D(X, X') + D(Y, Y'))$.

Proof.

- (i) By the chain rule: $H(X) + H(Y | X) = H(X, Y) = H(Y) + H(X | Y)$, hence $|H(X) - H(Y)| = |H(X | Y) - H(Y | X)| \leq H(X | Y) + H(Y | X) = D(X, Y)$.

- (ii) Applying the inequality (i) to the variables (X, Y) and (X', Y') , we obtain $|H(X, Y) - H(X', Y')| \leq D((X, Y), (X', Y'))$. From the continuity of joint information (Proposition 20), one can further bound $D((X, Y), (X', Y')) \leq D(X, X') + D(Y, Y')$.
- (iii) By the chain rule, $|H(X|Y) - H(X'|Y')| = |H(X, Y) - H(Y) - (H(X', Y') - H(Y'))| \leq |H(X, Y) - H(X', Y')| + |H(Y') - H(Y)|$. The conclusion now follows from (i) and (ii).
- (iv) By the chain rule, $|I(X; Y) - I(X'; Y')| = |H(X) - H(X') + H(Y) - H(Y') + H(X', Y') - H(X, Y)| \leq |H(X) - H(X')| + |H(Y) - H(Y')| + |H(X', Y') - H(X, Y)|$. The conclusion follows from bounding each of the three terms in the sum using (i) and (ii). $\square$

In the remainder of this paper, we only consider quantities that are *continuous* with respect to the entropic metrics (Shannon and Rajsiki distance). As a result, the discontinuity of the $\wedge$ operator will not hinder our derivations in the sequel.

4. Geometric Properties of the Information Lattice

4.1. Alignments of Random Variables

Definition 6 (Alignment). Let δ be any distance on the information lattice. The random variables X, Y , and Z are said to be aligned with respect to δ if the triangular inequality is met with equality:

$$\delta(X, Y) + \delta(Y, Z) = \delta(X, Z). \quad (19)$$

Proposition 22 (Alignment with respect to the Shannon distance D). The random variables X, Y , and Z are aligned with respect to D if and only if $X - Y - Z$ is a Markov chain and $Y \leq X \vee Z$.

This alignment condition is illustrated in Figure 4.

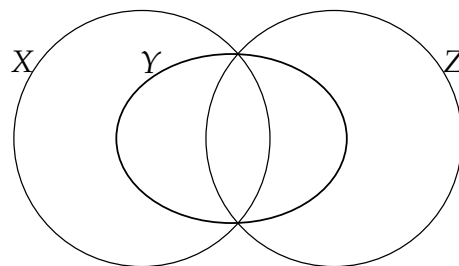


Figure 4. Venn diagram illustrating the alignment condition for the Shannon distance.

Proof. From the proof of the triangular inequality for D (Proposition 15), the equality holds iff equality holds in both inequalities $H(X|Z) \leq H(X, Y|Z) = H(X|Y, Z) + H(Y|Z) \leq H(X|Y) + H(Y|Z)$ and those inequalities obtained by permuting the roles of X and Z . Since $H(X, Y|Z) - H(X|Z) = H(Y|X, Z)$ and $H(X|Y) - H(X|Y, Z) = I(X; Z|Y)$, the equality holds iff $H(Y|X, Z) = 0$ and $I(X; Z|Y) = 0$, both conditions being symmetric in (X, Z) . Now, $H(Y|X, Z) = 0$ means that Y is a function of (X, Z) , i.e., $Y \leq X \vee Z$. Also, $I(X; Z|Y) = 0$ means that X and Z are conditionally independent given Y , which characterizes the fact that $X - Y - Z$ forms a Markov chain. $\square$

Proposition 23 (Alignment with respect to Rajsiki's distance d). The random variables X, Y , and Z are aligned with respect to d if and only if either $X = Y$, $Y = Z$, or $Y = X \vee Z$.

This alignment condition $Y = X \vee Z$ is illustrated in Figure 5.

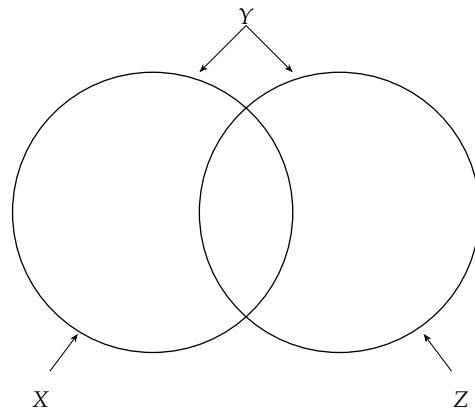


Figure 5. Venn diagram illustrating the alignment condition for the Rajski distance.

Proof. Alignment trivially holds when $X = Y$ or $Y = Z$. More generally, from the proof of the triangular inequality for d (Proposition 16), alignment holds iff the equality holds in all inequalities (9), (10), and (12) and those inequalities obtained by permuting the roles of X and Z .

Now, a close inspection of (9) shows that it achieves equality iff $H(X|Y) = 0$ or $H(Z|Y) = 0$, that is $X \leq Y$ or $Z \leq Y$. This condition is written as $X \vee Z \leq Y$ and is already symmetric in X, Z .

Similarly, (10) achieves equality iff $H(Y|Z) = 0$ or $H(X|Y) = 0$, that is $Y \leq Z$ or $X \leq Y$. Permuting the roles of X, Z , we also have the condition $Y \geq Z$ or $X \geq Y$. Thus, leaving aside the trivial solutions $X = Y$ or $Y = Z$, one necessarily has either $X \leq Y$ or $Z \leq Y$, which is the same as the equality condition in (9), or the opposite inequalities, $X \geq Y$ and $Z \geq Y$. In this latter case, combining with the equality condition in (9), we again end up with the trivial solutions $X = Y$ or $Y = Z$.

Thus, leaving aside the trivial solutions $X = Y$ or $Y = Z$, both conditions are written as $X \vee Z \leq Y$, which is symmetric in (X, Z) . Finally, (12) achieves equality iff $H(X|Y) + H(Y|Z) = H(X|Z)$ and the corresponding equality obtained by permuting the roles of X and Z . This means that X, Y , and Z are aligned with respect to D , that is $X - Y - Z$ is a Markov chain and $Y \leq X \vee Z$. Overall, $Y = X \vee Z$, which already implies that X and Z are conditionally independent given $Y = (X, Z)$, i.e., $X - Y - Z$ is a Markov chain. $\square$

Remark 12. Note that if X, Y , and Z are aligned in the sense of Rajski's distance, then they are also aligned in the sense of Shannon's entropic distance since $Y = (X, Z)$ implies that $X - Y - Z$ is a Markov chain. Thus, the alignment condition is stronger in the case of the Rajski distance.

Remark 13. The alignment condition with respect to the Rajski distance is simpler and expressed by using only the operators of the information lattice, whereas that with respect to the Shannon distance requires the additional notion of the Markov chain. Therefore, in the sequel, we develop some geometrical aspects of the information lattice based essentially on the Rajski distance.

4.2. Convex Sets of Random Variables in the Information Lattice

Definition 7 (Convexity). Given two random variables X and Y , we define the segment $[X, Y]$ of endpoints X and Y as the set of all random variables Z such that X, Z , and Y are aligned with respect to the Rajski distance, i.e., such that $d(X, Z) + d(Z, Y) = d(X, Y)$.

We say that a set $\mathcal{C}$ of points (random variables) in the information lattice is convex if, for all points $X, Y \in \mathcal{C}$, the segment $[X, Y] \subseteq \mathcal{C}$. If $\mathcal{S}$ is any set of points of the information lattice, its convex envelope is the smallest convex set containing $\mathcal{S}$.

By its very definition, the convex envelope of the two-element set $\{X, Y\}$ is the segment $[X, Y]$. We have the following simple characterization.

Proposition 24 (Segment characterization). *For any two elements X and Y of the information lattice, the segment $[X, Y]$ is the three-element set $[X, Y] = \{X, (X, Y), Y\}$, with the respective distances to the endpoints given by $d(X, (X, Y)) = \frac{H(Y|X)}{H(X, Y)}$ and $d(Y, (X, Y)) = \frac{H(X|Y)}{H(X, Y)}$.*

Proof. X and Y do belong to the segment $[X, Y]$ since $d(X, X) + d(X, Y) = d(X, Y)$ and $d(X, Y) + d(Y, Y) = d(X, Y)$. Moreover, if $Z \in [X, Y]$, then X, Z , and Y are aligned with respect to the Rajski distance so that, necessarily, $Z = (X, Y)$. One calculates $d(X, (X, Y)) = \frac{H(X|X, Y) + H(X, Y|X)}{H(X, Y)} = \frac{H(Y|X)}{H(X, Y)}$ and similarly for $d(Y, (X, Y))$ by permuting the roles of X and Y . $\square$

Remark 14. By the above Proposition, segments in the information lattice are intrinsically discrete objects. In the case where $X \leq Y$ or $Y \leq X$, then the segment $[X, Y]$ contains only two distinct points, X and Y . Obviously, if $X = Y$, then $[X, Y]$ is a singleton. This gives three possible cases as illustrated in Figure 6.

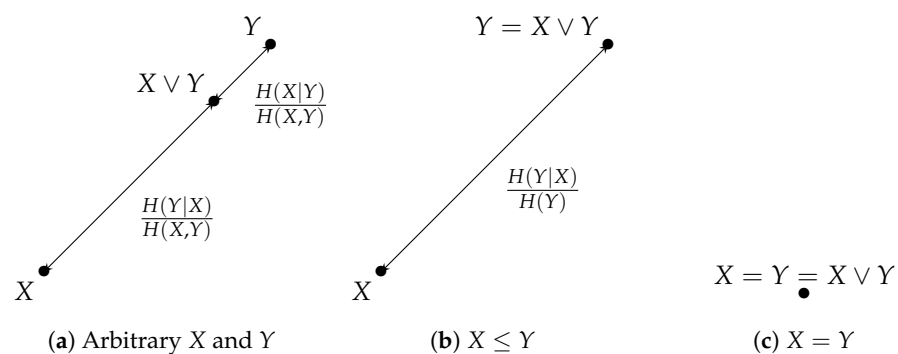


Figure 6. Visualization of the segment $[X, Y]$ for three possible cases.

Remark 15. As a result of this characterization, four or more distinct points cannot be aligned with respect to the Rajski distance, because a segment cannot contain more than three distinct points.

Proposition 25. $\mathcal{C}$ is convex iff it is closed under the $\vee$ operator.

Proof. $\mathcal{C}$ is convex iff, for all $X, Y \in \mathcal{C}$, $[X, Y] \subseteq \mathcal{C}$, that is X, Y , and $(X, Y) = X \vee Y \in \mathcal{C}$. $\square$

Beyond the case of a two-element set, we now characterize the convex envelope of any n -element set in the information lattice, that is the convex envelope of n random variables $X_1, X_2, \dots, X_n$. We adopt the following usual convention. For any n -tuple of indices $I = (i_1, i_2, \dots, i_n)$, the random vector $(X_{i_1}, X_{i_2}, \dots, X_{i_n}) = X_{i_1} \vee X_{i_2} \vee \dots \vee X_{i_n}$ is denoted by X_I . Again, by convention, for the empty set, $X_\emptyset = 0$, so that one always has $X_{I \cup J} = X_I \vee X_J$ for any two finite sets of indices I and J .

Proposition 26. Let I be a finite index set and $(X_i)_{i \in I}$ be random variables. The convex envelope of $(X_i)_{i \in I}$ is $\{\vee_{j \in J} X_j \mid \emptyset \neq J \subseteq I\} = \{X_J \mid \emptyset \neq J \subseteq I\}$, that is the set of all sub-tuples of the X_i .

Proof. With every X_i ($i \in I$), the convex envelope in question should be closed by the $\vee$ operator, hence contain any tuple $\vee_{j \in J} X_j$ for any nonempty $J \subseteq I$. Now, $\mathcal{C} = \{\vee_{j \in J} X_j = X_J \mid \emptyset \neq J \subseteq I\}$ contains all X_i for $i \in I$ and is already convex. Indeed, for all $X_J \in \mathcal{C}$ and $X_K \in \mathcal{C}$, $X_J \vee X_K = X_{J \cup K} \in \mathcal{C}$. $\square$

Remark 16. Given a finite set I of an index of cardinality $|I| = n$, the convex envelope of $(X_i)_{i \in I}$ contains at most $2^n - 1$ distinct elements, since there are $2^n - 1$ nonempty subsets of I . The number $2^n - 1$ is only an upper bound since it might happen that two different subsets J and K of I are such that $X_J = X_K$.

An example of the convex envelope of a family of three random variables is shown in Figure 7.

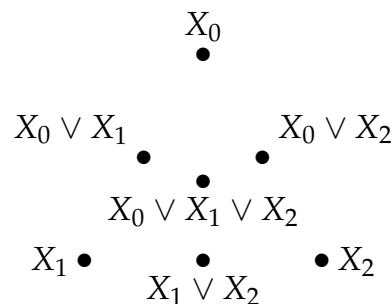


Figure 7. Seven-element convex envelope of three random variables X_0 , X_1 , and X_2 . These three random variables are represented as vertices of an (equilateral) triangle. The other points in the convex envelope are obtained as intersections of medians and edges, and the common information $X_0 \vee X_1 \vee X_2$ is the center of gravity (intersection of the three medians). Similarly, the 15-element convex envelope of four distinct points can be visualized as a tetrahedron, etc.

It is also interesting to note that any sublattice of the information lattice does have some convexity properties:

Proposition 27. *Any sublattice of the information lattice (including the information lattice itself) is convex in the sense of Definition 7. If a sublattice contains a subset of points $(X_i)_{i \in I}$, it also contains every point in the convex envelope of $(X_i)_{i \in I}$.*

Proof. With every two points X, Y , the sublattice should contain their maximum $X \vee Y$, hence the whole segment $[X, Y]$. It is, therefore, convex. Now, every convex set contains the convex envelope of any of its subsets. $\square$

4.3. The Lattice Generated by a Random Variable

In the sequel, we are interested in all possible deterministic functions of a given random variable X . In fact, their set constitutes a sublattice of the information lattice:

Proposition 28 (Sublattice generated by a random variable). *Let X be any random variable in the information lattice. The set of all random variables $\leq X$ is a sublattice, which we call lattice generated by X , denoted $\langle X \rangle$. It is a bounded lattice with maximum (total information) X and minimum 0.*

Proof. Let $Y \leq X$ and $Z \leq X$. There exists deterministic functions f and g such that $Y = f(X)$ a.s. and $Z = g(X)$ a.s. Clearly, $Y \wedge Z \leq Y \leq X$ and $Y \vee Z = (Y, Z) = (f(X), g(X)) \leq X$. Therefore, the set of random variables $\leq X$ forms a sublattice. Clearly, X is maximum, and 0 (deterministic random variable seen as a constant function of X) is minimum. $\square$

Remark 17. *One may also define the sublattice $\langle X_1, X_2, \dots, X_n \rangle$ generated by several random variables $X_1, X_2, \dots, X_n$ simply as the sublattice generated by the variable $X_1 \vee X_2 \vee \dots \vee X_n$. Therefore, it is enough to restrict ourselves to one random variable X as the lattice generator.*

Proposition 29. $\langle X \rangle$ is a complemented lattice.

Proof. Let $Y \leq Z \leq X$, so that both $Y, Z \in \langle X \rangle$. By Proposition 8, Y admits at least one complement information $\bar{Y}$ with respect to Z in the information lattice, such that $Y \wedge \bar{Y} = 0$ and $Y \vee \bar{Y} = Z$. Now, $\bar{Y} \leq Z \leq X$; hence, the complement $\bar{Y} \in \langle X \rangle$ belongs to the sublattice generated by X . $\square$

4.4. Properties of Rajski and Shannon Distances in the Lattice Generated by a Random Variable

We now investigate the metric properties of the sublattice $\langle X \rangle$ generated by a random variable X . To avoid the trivial case $\langle 0 \rangle = \{0\}$, we assume that X is *nondeterministic*. First of all, we observe that the entropy of an element of the sublattice increases as it is closer to X (in terms of either Shannon's or Rajski's distance):

Proposition 30. *For any $Y \in \langle X \rangle$, one has $D(X, Y) = H(X|Y) = H(X) - H(Y)$ and $d(X, Y) = \frac{H(X|Y)}{H(X)} = 1 - \frac{H(Y)}{H(X)}$. In particular, the maximum distance $d(X, Y) = 1$ is achieved iff $Y = 0$.*

Proof. One has

$$\begin{aligned} d(X, Y) &= \frac{D(X, Y)}{H(X, Y)} = \frac{H(X|Y) + H(Y|X)}{H(X)} \\ &\stackrel{(a)}{=} \frac{H(X|Y)}{H(X)} \stackrel{(b)}{=} \frac{H(X) - H(Y)}{H(X)} = 1 - \frac{H(Y)}{H(X)}. \end{aligned} \quad (20)$$

where (a) is because $Y \leq X$ and (b) is a consequence of the chain rule: $H(X) = H(X, Y) = H(Y) + H(X|Y)$. $\square$

Remark 18. *In the language of data compression, $d(X, Y) = \frac{H(X) - H(Y)}{H(X)} = 1 - \frac{H(Y)}{H(X)}$ can be seen as the relative entropic redundancy of X when it is represented ("encoded") by Y .*

Remark 19. *The maximum distance case in the proposition can be stated as follows: The only random variables Y that can be obtained as functions of X ($Y \in \langle X \rangle$) while being also independent of X ($d(X, Y) = 1$) are the constant (deterministic) random variables.*

4.5. Triangle Properties of the Shannon Distance

At least one attempt has been made previously by Donderi [18,19] to relate the entropic distance D to Euclidean geometry. Referring to Shannon's lattice of information, Donderi defined the distance between X and Y to be $\sqrt{D(X, Y)} = \sqrt{H(X|Y) + H(Y|X)}$, rather than $D(X, Y)$, and postulated that such a distance satisfies the usual properties of a Euclidean distance, such as the trigonometric law of cosines for a triangle (see Figure 1 in [18] and Figure 2 in [19]). This, in fact, is not the case, and the geometry of the triangle has to be re-thought in a non-Euclidean way as follows.

In Euclidean geometry, Apollonius's theorem allows one to calculate the length of the median of a triangle XYZ given the length of its other three sides. In the information lattice context, $Y \vee Z$ denotes the median of the segment $[Y, Z]$ (the only possible point in the segment that is not an endpoint). Thus, Apollonius's theorem gives a formula for the distance $D(X, Y \vee Z)$ in terms of $D(X, Y)$, $D(X, Z)$, and $D(Y, Z)$. The following Proposition is the analogue of Apollonius's theorem for the Shannon distance in the information lattice generated by X :

Lemma 3 (Apollonius's theorem in $\langle X \rangle$). *For any $Y, Z \in \langle X \rangle$,*

$$D(X, Y \vee Z) = \frac{D(X, Y) + D(X, Z) - D(Y, Z)}{2}. \quad (21)$$

This can also be written as

$$D(X, Y) + D(X, Z) = D(Y, Z) + 2D(X, Y \vee Z) \quad (22)$$

This is illustrated in Figure 8. Note that, when $X = Y \vee Z$, one recovers that Y, X, Z (in this order) are aligned.

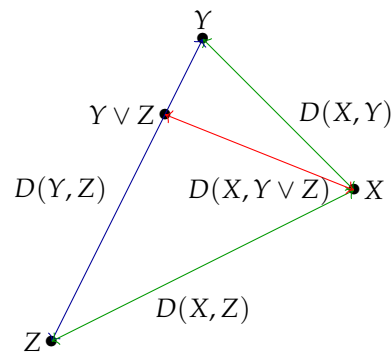


Figure 8. Graphical representation of Apollonius's theorem (Lemma 3).

Proof. From Proposition 30, $D(X, Y) = H(X) - H(Y)$ for any $Y \in \langle X \rangle$, in particular $D(X, Z) = H(X) - H(Z)$ and $D(X, Y \vee Z) = H(X) - H(Y, Z)$ also. Therefore, $D(X, Y) + D(X, Z) - 2D(X, Y \vee Z) = 2H(Y, Z) - H(Y) - H(Z) = H(Z|Y) + H(Y|Z) = D(Y, Z)$. $\square$

From Lemma 3, we derive the following,

Lemma 4. For any $Y, Z \in \langle X \rangle$,

$$d(X, Y) + d(X, Z) \leq d(X, Y \vee Z) + 1 \quad (23)$$

with equality if and only if Y and Z are independent.

Proof. Observe that $D(Y, Z) + D(X, Y \vee Z) = H(Y|Z) + H(Z|Y) + H(X|Y \vee Z) \leq H(Y) + H(Z|Y) + H(X|Y, Z) = H(Y, Z) + H(X|Y, Z) = H(X, Y, Z) = H(X)$ since $Y, Z \in \langle X \rangle$, with equality iff Y and Z are independent. Now, by Lemma 3, $D(X, Y) + D(X, Z) = D(Y, Z) + 2D(X, Y \vee Z) \leq D(X, Y \vee Z) + H(X)$. Dividing by $H(X) = H(X, Y) = H(X, Z) = H(X, Y, Z)$ yields the announced inequality. $\square$

In the other direction, we have the following.

Lemma 5. For any $Y, Z \in \langle X \rangle$,

$$d(X, Y \vee Z) \leq d(X, Y) + d(X, Z) \quad (24)$$

with equality if and only if $X = Y = Z$.

Proof. By the triangular inequality, $d(X, Y \vee Z) \leq d(X, Y) + d(Y, Y \vee Z)$ with equality iff $Y = X \vee Y \vee Z = X$ by the alignment condition. Similarly, $d(X, Y \vee Z) \leq d(X, Z) + d(Z, Y \vee Z)$ with equality iff $Z = X \vee Y \vee Z = X$. Summing the two inequalities, $2d(X, Y \vee Z) \leq d(X, Y) + d(X, Z) + d(Y, Y \vee Z) + d(Z, Y \vee Z)$, where $d(Y, Y \vee Z) + d(Z, Y \vee Z) = d(Y, Z) \leq d(X, Y) + d(X, Z)$ with equality iff $X = Y \vee Z$. Combining yields the announced inequality. $\square$

Remark 20. In the course of the proof, we proved the following stronger inequality: for any $Y, Z \in \langle X \rangle$,

$$d(X, Y \vee Z) \leq \frac{d(X, Y) + d(Y, Z) + d(Z, X)}{2} \quad (25)$$

with the same equality condition $X = Y = Z$.

Remark 21. By Lemmas 4 and 5, we see that, in terms of the Rajska distances to the generator X , $d(X, Y \vee Z)$ lies between $d(X, Y) + d(X, Z) - 1$ and $d(X, Y) + d(X, Z)$, where the lower and upper bounds differ by one and the minimum value is achieved in the case of independence. These two Lemmas are instrumental in the derivations of the next section.

5. The Perfect Reconstruction Problem

5.1. Problem Statement

Suppose one is faced with the following reconstruction problem. We are given a (discrete) source of information X (e.g., a digital signal, some text document, or any type of data), which is processed using deterministic functions into several “components”:

$$X_1 = f_1(X), \quad X_2 = f_2(X), \quad \dots, \quad X_n = f_n(X) \quad (26)$$

(e.g., different filtered versions of the signal at various frequencies, translated parts of the document, or some nonlinear transformations of the data). The natural question is: Did one *lose information* when processing X into its n components $X_1, X_2, \dots, X_n$, or else, can we *perfectly reconstruct* the original X from its n components using some (unknown) deterministic function $X = f(X_1, \dots, X_n)$?

We emphasize that all involved functions must be *deterministic* (no noise is involved), otherwise *perfect reconstruction* (without error) would not be possible. Yet, we do not require any precise form for the reconstruction function f , only that such a reconstruction exists. To our knowledge, the first occurrence of such a problem (for $n = 2$) is Exercise 6 of the textbook [20].

Stated in the information lattice language, the perfect reconstruction problem is as follows. Suppose we are given $X_1, X_2, \dots, X_n$ in $\langle X \rangle$, the sublattice generated by X . Is it true that $X \leq X_1 \vee X_2 \vee \dots \vee X_n$? Since the sublattice is convex (Proposition 27), i.e., stable by the $\vee$ operator (Proposition 25), one always has, by assumption, that $X_1 \vee X_2 \vee \dots \vee X_n \in \langle X \rangle$, i.e., $X_1 \vee X_2 \vee \dots \vee X_n \leq X$. Therefore, in the reconstruction problem, it is equivalent to determining whether $X = X_1 \vee X_2 \vee \dots \vee X_n$ or $X \neq X_1 \vee X_2 \vee \dots \vee X_n$.

Remark 22. Geometrically, by Proposition 26, determining whether $X \leq X_1 \vee X_2 \vee \dots \vee X_n$ or not is equivalent to determining whether X is in the convex envelope of $(X_i)_{i=1, \dots, n}$.

Thus, when $n = 2$, perfect reconstruction is possible iff X lies in the segment $[X_1, X_2]$. When $n = 3$, perfect reconstruction is possible iff, for every distinct index $i, j, k \in \{1, 2, 3\}$, X_i, X , and $X_{j,k}$ are aligned with respect to the Rajska distance, as illustrated in Figure 9.

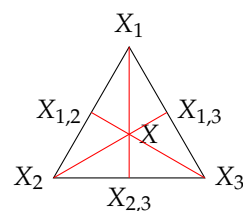


Figure 9. Geometric illustration of the three-component reconstruction problem.

Intuitively, the processed components X_i should not (on the whole) be too “far away” from the original source X in order that perfect reconstruction be possible. In other words, at least some of the distances $d(X, X_i)$ should not be too high. Such distances can be, in principle, evaluated when processing the source X into each of its components. In the following subsection, we give a simple necessary condition on the sum $d(X, X_1) + d(X, X_2) + \dots + d(X, X_n)$ to allow for perfect reconstruction.

5.2. A Necessary Condition for Perfect Reconstruction

The main result of this paper is the following.

Theorem 1 (Necessary condition for perfect reconstruction). *Let X be a random variable, and let $X_1, X_2, \dots, X_n \in \langle X \rangle$. If perfect reconstruction is possible: $X = X_1 \vee X_2 \vee \dots \vee X_n$, then*

$$\sum_{i=1}^n d(X, X_i) \leq n - 1 \quad (27)$$

with equality iff $X_1, X_2, \dots, X_n$ are independent.

Proof. By repeated use of Lemma 4, each joining operation of two components in the sum—e.g., passing from $d(X, X_i) + d(X, X_j)$ to $d(X, X_i \vee X_j)$ —decreases this sum by at most one. Thus,

$$\begin{aligned} \sum_{i=1}^n d(X, X_i) &\leq \sum_{i=1}^{n-2} d(X, X_i) + d(X, X_{n-1} \vee X_n) + 1 \\ &\leq \sum_{i=1}^{n-3} d(X, X_i) + d(X, X_{n-2} \vee X_{n-1} \vee X_n) + 2 \\ &\vdots \\ &\leq d(X, X_1 \vee X_2 \vee \dots \vee X_n) + n - 1 = n - 1. \end{aligned} \quad (28)$$

The equality holds iff all the above $n - 1$ inequalities are equalities. By the equality condition of Lemma 4, this means by induction that X_1 is independent of $X_2 \vee \dots \vee X_n$, where X_2 is independent of $X_3 \vee \dots \vee X_n$, and so on, until X_{n-1} is independent of X_n . Overall, this is equivalent to saying that all components $X_1, X_2, \dots, X_n$ are mutually independent. $\square$

Remark 23. To illustrate Theorem 1, consider a uniformly distributed two-bit random variable X (i.e., the result of two independent coin flips), and let X_1 be the result of the first coin toss and X_2 be that of the second coin toss. Clearly, reconstruction is possible since $X = (X_1, X_2)$. Now, a simple calculation gives $d(X, X_1) = \frac{H(X|X_1)}{H(X)} = \frac{\log 2}{\log 4} = \frac{1}{2}$, and similarly, $d(X, X_2) = \frac{1}{2}$, which shows that (27) is achieved with equality: $d(X, X_1) + d(X, X_2) = 2 - 1 = 1$. This is not surprising since X_1 and X_2 are independent, as can be checked directly.

Now, consider $X_3 = 0$ or 1 depending on whether $X_1 = X_2$ or not. Clearly, X can be also reconstructed from X_1, X_2, X_3 since it can already be reconstructed from X_1, X_2 . Again, one computes $d(X, X_3) = \frac{\log 2}{\log 4} = \frac{1}{2}$, so in this case, the sum of the distances to X is now $d(X, X_1) + d(X, X_2) + d(X, X_3) = \frac{3}{2} < 3 - 1 = 2$. This shows that (27) is still satisfied, but not with equality. In fact, it can easily be proven that, even though X_1, X_2, X_3 are pairwise independent, they are not mutually independent.

Remark 24. In practice, Theorem 1 gives an impossibility condition for the perfect reconstruction of the random variable X from components $X_1, X_2, \dots, X_n$. Indeed, if the latter are such that

$$\sum_{i=1}^n d(X, X_i) > n - 1 \quad (29)$$

then perfect reconstruction is impossible, however complex the reconstruction function f could have been. In other words, $X < X_1 \vee X_2 \vee \dots \vee X_n$, and information was lost by processing.

That perfect reconstruction is impossible does not mean that it would never be possible to deduce one particular value of X from some particular values of $X_1, X_2, \dots, X_n$. This means that such a deduction is not possible in general, for every possible value taken by $X_1, X_2, \dots, X_n$. In other words, there is at least one set of values $X_1 = x_1, X_2 = x_2, \dots, X_n = x_n$ for which X cannot be reconstructed unambiguously.

Remark 25. Another look at Theorem 1 can be made using the dependency coefficient $\rho = 1 - d$ in place of the Rajski distance. Then, the impossibility condition (29) is simply written as

$$\sum_{i=1}^n \rho(X, X_i) < 1. \quad (30)$$

In other words, perfect reconstruction can only occur if the components are (as a whole) sufficiently dependent on the original X . Otherwise, (30) precludes perfect reconstruction.

Remark 26. Since the Rajski distance is always upper bounded by one, if the impossibility condition (29) is met, then the actual value of the sum $\sum_{i=1}^n d(X, X_i)$ necessarily lies in the interval $(n - 1, n]$.

In the worst situation $\sum_{i=1}^n d(X, X_i) = n$, all terms should equal one: $d(X, X_i) = 1$. This means that all components are independent of X . By Proposition 30, the components $X_i = 0$ are all constants: in this case, all information is lost.

Remark 27. By Theorem 1, for perfect reconstruction to be possible, the components X_i should be (at least slightly) tightened around X in the sense that (27) is satisfied. The example of Remark 23 shows that, under this condition (even when the inequality is strict), it may be actually possible to reconstruct X . However, proximity may not be enough: the necessary condition of Theorem 1 is not sufficient in general.

To see this, consider X uniformly distributed in the integer interval $\{0, 1, \dots, 11\}$, and define $X_1 = k$ if $X = 2k$ or $2k + 1$ and $X_2 = \ell$ if $X = 3\ell, 3\ell + 1$, or $3\ell + 2$. In other words X_1 is the integer division of X by 2, and X_2 is the integer division of X by 3. One easily computes

$$d(X, X_1) + d(X, X_2) = \frac{H(X|X_1) + H(X|X_2)}{H(X)} = \frac{\log 2 + \log 3}{\log 12} = \frac{\log 6}{\log 12} < 1. \quad (31)$$

While the necessary condition (27) of Theorem 1 is met, the value of X cannot be unambiguously determined from those of X_1 and X_2 . For example, $X_1 = X_2 = 0$ leaves two possibilities: $X = 0$ or 1. Therefore, perfect reconstruction is not possible.

Another way to see this is to observe that perfect reconstruction is equivalent to saying that X_1, X, X_2 are aligned, which in terms of the Shannon distance would be written as $D(X_1, X_2) = D(X, X_1) + D(X, X_2)$. But, while $D(X, X_1) + D(X, X_2) = \log 6$, one has

$$D(X_1, X_2) = H(X_1|X_2) + H(X_2|X_1) = \left(\frac{1}{3} \log 3 + \frac{2}{3} \log \frac{3}{2}\right) + \frac{2}{6} \log 2 = \log 3 - \frac{\log 2}{3} \quad (32)$$

which is clearly less than $\log 6$. Therefore, perfect reconstruction is impossible in our example, because X_1 and X_2 are too close together, i.e., there is too much redundant information between them.

A slight modification of the above example where X takes values in $\{0, 1, \dots, 12m - 1\}$ for arbitrarily large m shows that the sum $d(X, X_1) + d(X, X_2) = \frac{\log 6}{\log(12m)}$ can actually be as small as desired, while perfect reconstruction is still impossible. Therefore, there can be no condition of the form $\sum_{i=1}^n d(X, X_i) < c$ (or any condition based only on the value of this sum) to ensure perfect reconstruction. Such a sufficient condition cannot be established without assuming some other property of the components X_i , as seen in the next subsection.

5.3. A Sufficient Condition for Perfect Reconstruction

For independent components $X_1, X_2, \dots, X_n$ (with no redundant information between them), the necessary condition of Theorem 1 becomes also a sufficient condition:

Theorem 2 (Sufficient condition for perfect reconstruction). Let X be a random variable, and let $X_1, X_2, \dots, X_n \in \langle X \rangle$ be independent. If the inequality (27) holds, then it necessarily holds with equality:

$$\sum_{i=1}^n d(X, X_i) = n - 1 \quad (33)$$

and perfect reconstruction is possible: $X = X_1 \vee X_2 \vee \dots \vee X_n$.

Proof. A closer look at the proof of Theorem 1 shows that we have established (without the perfect reconstruction assumption) the general inequality:

$$\sum_{i=1}^n d(X, X_i) \leq d(X, X_1 \vee X_2 \vee \dots \vee X_n) + n - 1 \quad (34)$$

which holds with equality iff $X_1, X_2, \dots, X_n$ are independent. Therefore, by the independence assumption, (27) is written as $\sum_{i=1}^n d(X, X_i) = d(X, X_1 \vee X_2 \vee \dots \vee X_n) + n - 1 \leq n - 1$. Since the distance is nonnegative, this necessarily implies that the inequality holds with equality and that $d(X, X_1 \vee X_2 \vee \dots \vee X_n) = 0$, that is $X = X_1 \vee X_2 \vee \dots \vee X_n$. $\square$

Remark 28. Following Remark 26, we see that, for independent $X_1, X_2, \dots, X_n$, the sum of the distances to X : $\sum_{i=1}^n d(X, X_i)$ can only take values in the interval $[n - 1, n]$, with two possibilities:

- Either $\sum_{i=1}^n d(X, X_i) = n - 1$, and perfect reconstruction is possible;
- Or $\sum_{i=1}^n d(X, X_i) > n - 1$, and perfect reconstruction is impossible.

In other words, independent components cannot be arbitrarily tightly packed around X .

Following Remark 25, in terms of dependency coefficients, for independent $X_1, X_2, \dots, X_n$:

- Either $\sum_{i=1}^n \rho(X, X_i) = 1$, and perfect reconstruction is possible;
- Or $\sum_{i=1}^n \rho(X, X_i) < 1$, and perfect reconstruction is impossible.

Remark 29. Following Remark 22 and Figure 9 in the case of three independent components X_1, X_2, X_3 , one should have $d(X, X_1) + d(X, X_2) + d(X, X_3) = 2$ for perfect reconstruction to hold. Incidentally, the graphical Euclidean illustration of Figure 9 is faithful in this case, since, for an equilateral triangle $X_1 X_2 X_3$ with sides of length one, the sum of the Euclidean distances equals $d(X, X_1) + d(X, X_2) + d(X, X_3) = \frac{2}{3} + \frac{2}{3} + \frac{2}{3} = 2$.

Remark 30. By Proposition 30, $d(X, X_i) = 1 - \frac{H(X_i)}{H(X)}$, so that Theorems 1 and 2 can be rewritten using the standard assertions that $H(X) \leq \sum H(X_i)$ with equality when X_i are mutually independent. This, of course, does not require all the machinery developed earlier. We feel, however, that our geometric vision is still valuable because of its conceptual and pedagogical interest and also as a starting point for a “perfect reconstruction theory”, which, of course, needs to be improved and further investigated along these lines.

5.4. Approximate Reconstruction

Suppose we encode the information source X by n components $X_1, X_2, \dots, X_n$, but do not particularly insist that perfect reconstruction is possible. Rather, we assume that the encoding removes a fraction of redundancy in X equal to

$$d(X, X_1 \vee X_2 \vee \dots \vee X_n) = \delta \quad (35)$$

(see Remark 18). Since the case $\delta = 0$ corresponds to the previous case of perfect reconstruction ($X = X_1 \vee X_2 \vee \dots \vee X_n$), we assume that $\delta > 0$ in the sequel. Thus, in what follows, the reconstruction of X can only be approximate (up to a certain distance tolerance δ). We then have the following.

Theorem 3 (Approximate reconstruction). *Let X be a random variable, and let $X_1, X_2, \dots, X_n \in \langle X \rangle$ such that (35) holds with redundancy $= \delta > 0$. Then,*

$$\delta < \sum_{i=1}^n d(X, X_i) \leq n - 1 + \delta. \quad (36)$$

with equality in the second inequality iff the components $X_1, X_2, \dots, X_n$ are independent.

Proof. The rightmost inequality in (36) is just (34) (with the announced case of equality), which was established by repeated application of Lemma 4. Similarly, the repeated application of Lemma 5 gives

$$d(X, X_1 \vee X_2 \vee \dots \vee X_n) \leq \sum_{i=1}^n d(X, X_i) \quad (37)$$

with equality iff all $X_i = X$ ($i = 1, \dots, n$). But, such an equality condition would yield $\delta = d(X, X) = 0$, contrary to the assumption $\delta > 0$. This shows that the leftmost inequality in (36) is strict. $\square$

Remark 31. Similarly, as in the above two subsections, one can deduce from Theorem 3 that, for independent components $X_1, X_2, \dots, X_n$, one necessarily has $\sum_{i=1}^n d(X, X_i) = n - 1 + \delta$ and that, in general, approximate reconstruction within distance tolerance $\leq \delta$ will be impossible if $\sum_{i=1}^n d(X, X_i) > n - 1 + \delta$.

6. Examples and Applications

In this section, we develop five examples of the applications of the theorems of Section 5.

6.1. Reconstruction from Sign and Absolute Value

Consider a real-valued random variable X , and assume that it is symmetric (X is identically distributed as $-X$) and that $\mathbb{P}(X = 0) = 0$. Now, define $X_1 = |X|$ (absolute value) and $X_2 = \text{sgn}(X) \in \{-1, 1\}$ (sign of X). Clearly, if X follows probability distribution p , then X_1 has probability distribution $2p(x)$ for $x > 0$. Also, X_2 is Rademacher distributed (equiprobable ± 1).

One easily computes $H(X_1) = \sum_{x>0} 2p(x) \log \frac{1}{2p(x)} = H(X) - \log 2$ and $H(X_2) = \log 2$ (equiprobable ± 1); hence, $d(X, X_1) = 1 - \frac{H(X_1)}{H(X)} = \frac{\log 2}{H(X)}$ and $d(X, X_2) = 1 - \frac{H(X_2)}{H(X)} = 1 - \frac{\log 2}{H(X)}$. Therefore, $d(X, X_1) + d(X, X_2) = 1$: Inequality (27) is satisfied with equality.

Of course, in this trivial example, perfect reconstruction is possible since $X = |X|\text{sgn}(X) = X_1 X_2$. Then, by Theorem 1, we deduce that X_1 and X_2 are independent. This is easily checked directly since, by the symmetry assumption, $\mathbb{P}(X_1 = x_1 | X_2 = \pm 1) = \mathbb{P}(X_1 = x_1)$. Notice that, from this independence, by Theorem 2, we find anew that perfect reconstruction of X is possible from X_1 and X_2 .

This example can be easily generalized to the case of a “symmetric” complex-valued random variable X with modulus $X_1 = |X|$ and argument $X_2 = \arg(X)$, where X_1 is independent of X_2 and X_2 is uniformly distributed over M possible values. Then, $H(X_2) = \log M$, $H(X_1) = H(X) - \log M$ by symmetry, and similar conclusions hold.

Of course, perfect reconstruction $X = X_1 X_2$ is always possible even in the case where X is *not* symmetric, in which case X_1 and X_2 are *not* independent, and therefore, by the alignment condition, $d(X, X_1) + d(X, X_2) = d(X_1, X_2) < 1$.

6.2. Linear Transformation over a Finite Field

Consider X uniformly distributed over $\mathbb{F}_q^k$, where $\mathbb{F}_q$ is the field with q elements. Suppose X is linearly transformed using some matrix $\mathbf{G}$ to obtain

$$(X_1, X_2, \dots, X_n) = X \cdot \mathbf{G} \quad (38)$$

in row vector notation, where $\mathbf{G} \in \mathbb{F}_q^{k \times n}$ has k rows and n columns. For example, X may represent information symbols to be transmitted over a channel, and $(X_1, X_2, \dots, X_n)$ would be the associated codeword using an (n, k) linear code over $\mathbb{F}_q$ with generator matrix $\mathbf{G}$.

If the i th column of $\mathbf{G}$ is not the all-zero vector, then it is easily checked that, since X is uniformly distributed over $\mathbb{F}_q^k$, X_i is likewise uniformly distributed over $\mathbb{F}_q$. Therefore,

$$d(X, X_i) = 1 - \frac{H(X_i)}{H(X)} = 1 - \frac{\log q}{\log q^k} = 1 - \frac{1}{k}. \quad (39)$$

When the i th column of $\mathbf{G}$ is all zero, however, $d(X, X_i) = d(X, 0) = 1$. Summing up,

$$\sum_{i=1}^n d(X, X_i) = n - \frac{n'}{k} \quad (40)$$

where $n' \leq n$ is the number of non-zero columns in $\mathbf{G}$.

By Theorem 1, if $n - \frac{n'}{k} > n - 1$, that is $n' < k$, then perfect reconstruction is impossible. This is quite natural since, in this case, $(X_1, \dots, X_n)$ entails less q -ary symbols than the vector X , so that it is impossible to reconstruct X from the n' actual symbols in $(X_1, \dots, X_n)$.

In general, if $\mathbf{G}$ has rank $r \leq \min(k, n')$, then, since X is uniformly distributed over $\mathbb{F}_q^k$, the vector $(X_1, \dots, X_n)$ is also uniformly distributed over a subspace of $\mathbb{F}_q^n$ of dimension r . Now, as we have just seen, if the i th column of $\mathbf{G}$ is not the all-zero vector, then X_i is uniformly distributed over $\mathbb{F}_q$. Since uniformly distributed components of a discrete random vector are independent iff the vector is itself uniformly distributed, the only possibility for the components $X_1, \dots, X_n$ to be independent as in Theorem 2 is that $(X_1, \dots, X_n)$ is uniformly distributed over $\mathbb{F}_q^{n'}$, that is $r = n' = k$. In this case $\sum_{i=1}^n d(X, X_i) = n - 1$, and by Theorem 2, perfect reconstruction is possible. Of course, from linear algebra, we know that X can be reconstructed from $(X_1, \dots, X_n)$ as soon as $\mathbf{G}$ has rank $r = k \leq n'$.

Due to the power of linear algebra, this example may appear quite trivial. It would be interesting to generalize it, however, to the case where the vector $(X_1, \dots, X_n)$ is obtained by a *nonlinear* transformation, i.e., each X_i are Boolean functions over $\mathbb{F}_q$ of the components of vector X , e.g., described in algebraic normal form.

6.3. Integer Prime Factorization

Consider an integer-valued random variable X , uniformly distributed over $\{1, 2, \dots, m\}$, and let $n = \pi(m)$ be the number of primes not exceeding m . For every such prime p , let X_p be the p -adic valuation of X , that is the largest exponent of p such that p^{X_p} divides X . We know by the fundamental theorem of arithmetic that the prime factorization of X always exists and is unique: $X = \prod_p p^{X_p}$; hence, X can be reconstructed from the X_p s.

There are $\lfloor \frac{m}{p^k} \rfloor$ values of X divisible by p^k and, therefore, $\lfloor \frac{m}{p^k} \rfloor - \lfloor \frac{m}{p^{k+1}} \rfloor$ values of X such that $X_p = k$. Thus, $H(X|X_p = k) = \log(\lfloor \frac{m}{p^k} \rfloor - \lfloor \frac{m}{p^{k+1}} \rfloor) \leq \log \frac{m}{p^k} = \log m - k \log p$, $H(X|X_p) \leq \log m - \mathbb{E}(X_p) \log p$, and

$$\sum_{p \text{ prime} \leq m} d(X, X_p) = \sum_{p \text{ prime} \leq m} \frac{H(X|X_p)}{H(X)} \leq \sum_{p \text{ prime} \leq m} \frac{\log m - \mathbb{E}(X_p) \log p}{\log m} = n - \frac{\log m!}{m \log m}. \quad (41)$$

In the latter equality, we used the exact value $\sum_p \mathbb{E}(X_p) \log p = \frac{\log m!}{m}$. This can be easily checked from the reconstruction formula itself, since

$$\sum_{p \text{ prime} \leq m} \mathbb{E}(X_p) \log p = \mathbb{E} \log \prod_{p \text{ prime} \leq m} p^{X_p} = \mathbb{E} \log X = \frac{\log m!}{m}. \quad (42)$$

Since $\log m! \leq m \log m$, the above upper bound is not tight enough to prove Inequality (27) of Theorem 1. It is only satisfied asymptotically as $m \rightarrow +\infty$ since, then, $\frac{\log m!}{m \log m} \rightarrow 1$. Likewise, the independence assumption of Theorem 2 is only true asymptotically: in fact, since, for distinct primes $p_1, \dots, p_\ell$,

$$\mathbb{P}(X_{p_1} \geq k_1, \dots, X_{p_\ell} \geq k_\ell) = \frac{1}{m} \left\lfloor \frac{m}{p_1^{k_1} \dots p_\ell^{k_\ell}} \right\rfloor \rightarrow \frac{1}{p_1^{k_1} \dots p_\ell^{k_\ell}} \quad (43)$$

it follows that the X_p s converge in distribution toward *independent* geometric variables with the respective parameters $1 - \frac{1}{p}$.

6.4. Chinese Remainder Theorem

Consider an integer-valued random variable X , uniformly distributed over $\{0, 1, \dots, k-1\}$, where $k = \prod_{i=1}^n k_i$ is the product of n pairwise coprime factors > 1 , and define the following remainders modulo these factors:

$$\begin{cases} X_1 \equiv X \pmod{k_1} \\ X_2 \equiv X \pmod{k_2} \\ \vdots \\ X_n \equiv X \pmod{k_n}. \end{cases} \quad (44)$$

By the well-known *Chinese remainder theorem*, this system of equations has a unique solution in $\{0, 1, \dots, k-1\}$, i.e., perfect reconstruction of X is possible.

Clearly, since X is uniformly distributed, X_i is likewise uniformly distributed over $\{0, 1, \dots, k_i-1\}$ so that $H(X_i) = \log k_i$, $d(X, X_i) = 1 - \frac{H(X_i)}{H(X)} = 1 - \frac{\log k_i}{\log k}$ and

$$\sum_{i=1}^n d(X, X_i) = \sum_{i=1}^n \left(1 - \frac{\log k_i}{\log k}\right) = n - \frac{\log \prod_{i=1}^n k_i}{\log k} = n - 1. \quad (45)$$

Thus, Inequality (27) of Theorem 1 is achieved with equality, which proves that $X_1, X_2, \dots, X_n$ are independent. Had we proven directly this independence, Theorem 2 would have shown that perfect reconstruction is possible. Thus, an information theoretic proof of the Chinese remainder theorem using this method amounts to proving such an independence. But, this can be performed quite similarly as the Chinese remainder theorem is classically proven.

With our present method, however, it can be easily seen that perfect reconstruction would *not* be possible if we do not use all components $X_1, X_2, \dots, X_n$. Indeed, suppose without loss of generality that one tries to reconstruct X only from $X_1, X_2, \dots, X_{n-1}$. Then, by the above calculation,

$$\sum_{i=1}^{n-1} d(X, X_i) = \sum_{i=1}^{n-1} \left(1 - \frac{\log k_i}{\log k}\right) = n - 1 - \frac{\log \prod_{i=1}^{n-1} k_i}{\log k} = n - 2 + \frac{\log k_n}{\log k} > n - 2 \quad (46)$$

which shows by Theorem 1 that perfect reconstruction of X from less than n remainders is impossible.

6.5. Optimal Sort

In this subsection, we provide a new information theoretic proof of the following.

Theorem 4. Any pairwise-comparison-based sorting algorithm has worst-case computational complexity $\geq \log_2 k! = \Omega(k \log_2 k)$, where k is the cardinality of the list to be sorted.

Recall that $\log$ refers to the logarithm taken to any base, while here, more specifically, $\log_2$ is the logarithm to base two.

Proof. Consider a finite, totally ordered list of k elements. It can be seen as a permutation of the uniquely sorted elements, and sorting this list amounts to finding this permutation. Let $X = (X_1, X_2, \dots, X_k)$ be a (uniformly chosen) random permutation on $\{1, 2, \dots, k\}$.

For $i, j \in \{1, \dots, k\}$ with $i \neq j$, let $X_{i,j}$ be the binary random variable taking the value 1 if $X_i < X_j$ and 0 otherwise. Clearly, $X_{i,j} \leq X$ for any i, j .

Since there are as many permutations such that $X_i < X_j$ such that $X_i > X_j$, every $X_{i,j}$ is a Bernoulli (1/2) variable (equiprobable bit). Therefore,

$$d(X, X_{i,j}) = 1 - \frac{\log 2}{\log k!} = 1 - \frac{1}{\log_2 k!}. \quad (47)$$

Assuming n pairwise comparisons are made to sort the complete list, this gives

$$\sum_{i,j \text{ (} n \text{ terms)}} d(X, X_{i,j}) = n - \frac{n}{\log_2 k!}. \quad (48)$$

By Theorem 1, it is necessary that this value does not exceed $n - 1$, i.e., $n \geq \log_2 k!$ for perfect reconstruction to hold. In other words, the worst-case complexity to achieve the complete sort for *any* possible realization of the initial unsorted list requires at least $\lceil \log_2 k! \rceil$ pairwise comparisons. $\square$

Remark 32. *This example illustrates a method to find a lower bound on the worst-case complexity of a problem. The first step is to express the instance of the problem as a random variable X . Second, one determines which pieces of information one is allowed to extract from X and models them as “observed” random variables $X_i \leq X$. Third, for each i , we compute the Rajski distance $d(X, X_i)$. Finally, we use Theorem 1 to find a lower bound on the number of “observed” variables X_i that are required to reconstruct X . We feel that such a method is interesting because it is often harder to find a lower bound on the complexity of a problem than to find an upper bound on it.*

7. Conclusions and Perspectives

It is an understatement to say that the “true” information theory of 1953 was not as popular as the classical theory of 1948. John Pierce, a colleague of Shannon, wrote that, “apparently the structure was not good enough to lead to anything of great value” [21]. We find two possible reasons for this pessimism: the fact that the lattice is not Boolean, which does not facilitate the calculations, and the discontinuous nature of the common information with respect to the entropy metric.

However, as we have shown in this paper, this lattice structure is quite helpful to understand reconstruction problems. As shown in Section 6, the implications of the resolution of perfect reconstruction problems go beyond signal processing, since the concept of information is pervasive in all fields of mathematics and of science. Thus, we believe it is important to deepen this theory, defining *information* per se, and to further generalize the reconstruction problems. It would indeed be of great interest to find a simple sufficient condition to reconstruct a variable X from the (not necessarily independent) components $X_1, X_2, \dots, X_n$.

One may legitimately argue that most examples (except in Section 6.1) assume uniform distributions, where the entropy is just a logarithmic measure of the alphabet size, and since all considered processings are deterministic, the essence of the present reconstruction problem appears more combinatorial than probabilistic. Indeed, a desirable perspective is to go beyond perfect reconstruction of discrete quantities by considering the possibility of the noisy reconstruction of discrete and/or continuous sources of information.

In a perspective closer to computer science, we used our theorems to find a lower bound on the complexity of the comparison-based sorting problem. It would be interesting to find other problems for which a lower bound on complexity can be found using our technique, especially for decision problems that are not known to be in P.

Finally, as another practical perspective for security problems, one may assume that X models all the possible values that can take a secret key in a given cryptographic device and that an attacker can observe k random values that are deterministically obtained from X . Such important problems have been studied, e.g., in [22] to evaluate information leakage in the execution of deterministic programs. One may use the theorems of Section 5 to find a lower bound on k for the attacker to be able to reconstruct the secret.

Author Contributions: Conceptualization, O.R., I.D., J.B. and A.S.; Formal analysis, I.D. and O.R.; Writing—original draft, I.D. and O.R.; Writing—review & editing, O.R., I.D., J.B., V.R. and A.S.; Supervision, O.R., V.R. and A.S. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Conflicts of Interest: The authors declare no conflicts of interest.

References

- Shannon, C.E. A mathematical theory of communication. *Bell Syst. Tech. J.* **1948**, *27*, 379–423. 623–656. [[CrossRef](#)]
- Shannon, C.E. The lattice theory of information. *Trans. Ire Prof. Group Inf. Theory* **1953**, *1*, 105–107. [[CrossRef](#)]
- Fano, R.M. Interview by Aftab, Cheung, Kim, Thkhar, Yeddapanudi, 6.933 Project History, Massachusetts Institute of Technology. November 2001.
- Fano, R.M. *Class Notes for Course 6.574: Transmission of Information*; MIT: Cambridge, MA, USA, 1952.
- Cherry, E.C. A history of the theory of information. *Proc. Inst. Electr. Eng.* **1951**, *98*, 383–393.
- Shannon, C.E. The bandwagon (editorial). In *IRE Transactions on Information Theory*; Institute for Radio Engineers, Inc.: New York, NY, USA, 1956; Volume 2, p. 3.
- Shannon, C.E. Some Topics on Information Theory. In *Proceedings of the International Congress of Mathematicians*, Cambridge, MA, USA, 30 August–6 September 1950; Volume II, pp. 262–263.
- Rioul, O.; Béguinot, J.; Rabiet, V.; Souloumiac, A. La véritable (et méconnue) théorie de l’information de Shannon. In *Proceedings of the 28e Colloque GRETSI 2022*, Nancy, France, 6–9 September 2022.
- Rajski, C. A metric space of discrete probability distributions. *Inf. Control* **1961**, *4*, 371–377. [[CrossRef](#)]
- Gács, P.; Körner, J. Common information is far less than mutual information. *Probl. Control Inf. Theory* **1973**, *2*, 149–162.
- Gamal, A.E.; Kim, Y.-H. *Network Information Theory*; Cambridge University Press: Cambridge, UK, 2011.
- Wyner, A.D. The common information of two dependent random variables. *IEEE Trans. Inf. Theory* **1975**, *21*, 163–179. [[CrossRef](#)]
- Nakamura, Y. Entropy and semivaluations on semilattices. *Kodai Math. Semin. Rep.* **1970**, *22*, 443–468. [[CrossRef](#)]
- Yeung, R.W. *Information Theory and Network Coding*; Springer: Berlin/Heidelberg, Germany, 2008.
- Horibe, Y. A note on entropy metrics. *Inf. Control* **1973**, *22*, 403–403. [[CrossRef](#)]
- Jaccard, P. Distribution de la flore alpine dans le bassin des Dranses et dans quelques régions voisines. *Bull. Société Vaudoise Des Sci. Nat.* **1901**, *37*, 241–272.
- Csiszár, I.; Körner, J. *Information Theory. Coding Theorems for Discrete Memoryless Systems*, 2nd ed.; Cambridge University Press: Cambridge, UK, 2011.
- Donderi, D.C. Information measurement of distinctiveness and similarity. *Percept. Psychophys.* **1988**, *44*, 576–584. [[CrossRef](#)] [[PubMed](#)]
- Donderi, D.C. An information theory analysis of visual complexity and dissimilarity *Perception* **2006**, *35*, 823–835. [[CrossRef](#)] [[PubMed](#)]
- Rioul, O. *Théorie de l’information et du Codage*; Hermes Science—Lavoisier: London, UK, 2007.
- Pierce, J.R. The early days of information theory. *IEEE Trans. Inf. Theory* **1973**, *19*, 3–8. [[CrossRef](#)]
- Malacaria, P. Algebraic foundations for quantitative information flow. *Math. Struct. Comput. Sci.* **2015**, *25*, 404–428. [[CrossRef](#)]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.