

3. elements of classical information theory

3.1 measure of information, Shannon entropy and classical data compression

► What is information?

example: throwing dice

- if a dice player tells you that he has thrown a "5" this is only information for you if you did not watch him throwing the dice
- if he says he has thrown a number between 1 and 6 this has also no information for you

information is a measure of (a priori) ignorance

Let A be a random event with probability $P(A)$.

$\bar{A}$ is called the complementary event of A if $P(\bar{A}) = 1 - P(A)$

A set of events $\{A_1, \dots, A_n\}$ is called complete, if

$$\sum_{i=1}^n P(A_i) = 1$$

Def: $\{A_1, \dots, A_n\}$ complete set of events
 A_i are called elementary events if

$$(1) \quad A_i \Rightarrow \bar{A}_j \quad \forall j \neq i$$

The set of all elementary events is called space of events Ω

principle of maximal ignorance

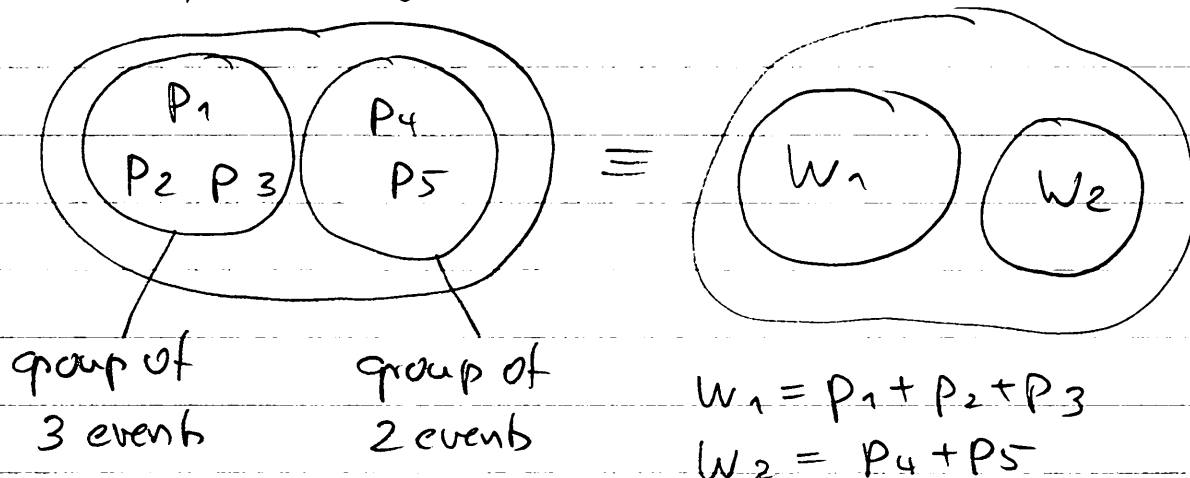
If nothing is known about the possible outcome of a random event, the a priori probability of every elementary event A_i is the same

$$(2) \quad P(A_i) = \frac{1}{n} \quad \forall A_i \quad i=1, \dots, n$$

quantitative measure of ignorance (information) S

n discrete events

- (i) S is a continuous function of the probabilities p_i
 $S = S(p_1, \dots, p_n)$
- (ii) $S = 0$ for a certain event (i.e. $p_j = 1$ for one j)
- (iii) $S = \max$ for maximal ignorance (i.e. all $p_i = \frac{1}{n}$)
and monotonously increasing with n
- (iv) S should fulfill the composition rule
which guarantees that it behaves as an
extensive variable



$$(3) \quad S(p_1, \dots, p_5) = S(w_1, w_2) + w_1 S\left(\frac{p_1}{w_1}, \frac{p_2}{w_1}, \frac{p_3}{w_1}\right) + w_2 S\left(\frac{p_4}{w_2}, \frac{p_5}{w_2}\right)$$

i.e. total information =

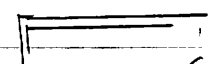
- (i) information if " w_1 or w_2 "
- + (ii) weighted information
 - "if w_1 " then " p_1 or p_2 or p_3 "
 - "if w_2 " then " p_4 or p_5 "

$\Rightarrow$ Shannon Entropy

$$(4) \quad S(p_1, \dots, p_n) = -a \sum_{j=1}^n p_j \ln p_j \quad a > 0$$

rem: the prefactor a can be chosen arbitrarily
often $a = 1/\ln 2$ $a \ln p_j \rightarrow \log_2 p_j = \log p_j$

proof of (i)-(iv):



(i) ✓

(ii) ✓ since $p_j = 1 \Rightarrow p_i = 0 \quad i \neq j$
and $0 \ln 0 = 0$

(iii) for this we first prove if $(p_1, \dots, p_n)$ and $(\bar{p}_1, \dots, \bar{p}_n)$ are two (different) probabilities of elementary events it holds

$$(5) \quad \sum_{i=1}^n \bar{p}_i \ln \frac{\bar{p}_i}{p_i} \geq 0$$

equality holds for $p_i = \bar{p}_i$

proof of (5): $\ln x \leq x - 1 \quad \forall x \geq 0$

$$\Rightarrow \ln 1/x = -\ln x \geq 1 - x$$

$$\Rightarrow \ln \frac{\bar{p}_i}{p_i} \geq 1 - \frac{p_i}{\bar{p}_i}$$

$$\Rightarrow \sum_i \bar{p}_i \ln \frac{\bar{p}_i}{p_i} \geq \sum_i \bar{p}_i - \sum_i p_i = 0 \quad \square$$

from (5) with $p_i = \frac{1}{n} \Rightarrow$

$$\sum_{i=1}^n \bar{p}_i (\ln \bar{p}_i - \ln p_i) \geq 0$$

$$\begin{aligned} S(\bar{p}_1, \dots, \bar{p}_n) &= -a \sum_{i=1}^n \bar{p}_i \ln \bar{p}_i \leq -a \sum_{i=1}^n \bar{p}_i \ln \frac{1}{n} \\ &= a \ln n \quad \square \end{aligned}$$

$$(iv) \quad w_1 = p_1 + \dots + p_m$$

$$w_2 = p_{m+1} + \dots + p_n$$

$$S(w_1, w_2) = -a w_1 \ln w_1 - a w_2 \ln w_2$$

$$\begin{aligned} w_1 S\left(\frac{p_1}{w_1}, \dots, \frac{p_m}{w_1}\right) &= -a w_1 \sum_{i=1}^m \frac{p_i}{w_1} \ln \left(\frac{p_i}{w_1}\right) \\ &= -a \sum_{i=1}^m p_i \ln p_i + a w_1 \ln w_1 \end{aligned}$$

$$\Rightarrow S(w_1, w_2) + w_1 S\left(\frac{p_1}{w_1}, \dots, \frac{p_m}{w_1}\right) + w_2 S\left(\frac{p_{m+1}}{w_2}, \dots, \frac{p_n}{w_2}\right)$$

$$= -a w_1 \ln w_1 - a w_2 \ln w_2$$

$$-a \sum_{i=1}^m p_i \ln p_i + a w_1 \ln w_1$$

$$-a \sum_{i=m+1}^n p_i \ln p_i + a w_2 \ln w_2 = \underline{\underline{S(p_1, \dots, p_n)}} \quad \square$$

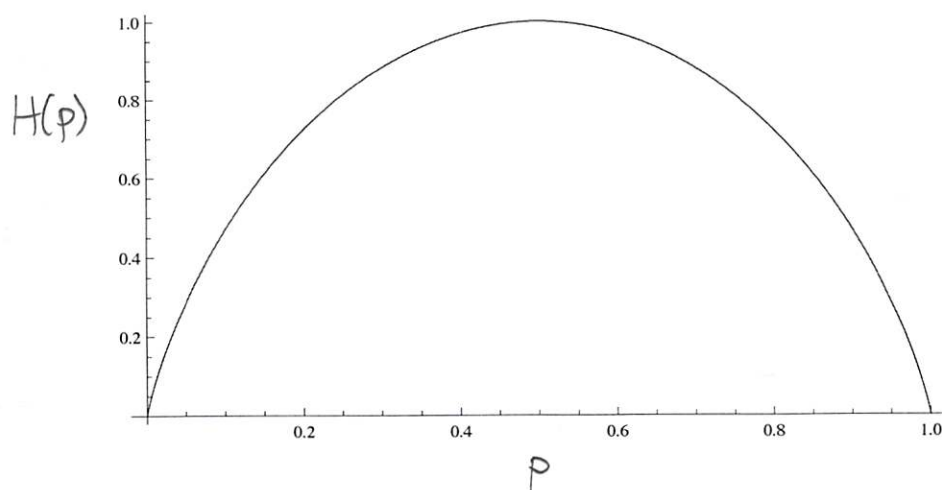
the smallest non-trivial event space has two alternatives

$$\{A, \bar{A}\} \hat{=} \{0, 1\}$$

this space is called a classical bit or cbit

for a c-bit holds $(a = 1/\ln 2)$

(6) $S(p, 1-p) = H(p) = -p \log p - (1-p) \log (1-p) = H(1-p)$



classical information is transferred as strings of c-bits so-called bit-words

(7) $X_1, X_2, \dots, X_n$ $X_i \in \{0, 1\}$

The relevance of the Shannon entropy can be seen most easily from the following problem: Given a bit word of length n with $n \rightarrow \infty$. How many cbits are needed to transfer the bit word von A to B?

obviously shorter!

Code:
 10 000 010 1001 10 1010 010 011 10 1000
 011 1001 1101 000 010 001 10 1001 000 1010
 1010 000 010 1101 10 000 010 1001 10 1001
 1011 010 001 10 1001 000 10 1000 000 010
 010 1101 000 10 1000 000 1101 010 010 001
 010 1001 10 1001 010 001 10 1001 000 010
 011 1001 1101 1110 010 001 10 1001 000 010
 10 1001 000 010 011 1001 1101 000 1101 011

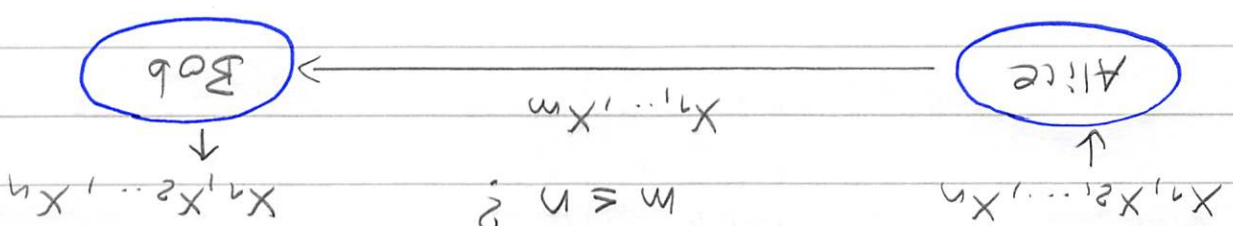
0000 0001 0100 0101 0000 0110 0100 1000 0000 0011
 1000 0101 1100 1101 0100 0010 0000 0101 0001 0100
 0100 0101 0000 0101 0100 0010 0000 0101 0001 0100
 0100 1100 0001 0000 0011 0001 1100 0100 0100 0010
 1001 0100 0010 0000 0101 0001 0000 0011 0001 0100
 0110 0001 0100 1100 0000 0001 0100 0101 0000 0101
 1000 0101 1100 0001 0100 0010 0000 0101 0001 0110
 0000 0001 0100 0101 0000 0110 0100 1000 0000 0011

example (spaces for better reading)

also: there is a unique code containing $1+nH(p)$ b.i.b!
 However, Shannon's theorem does not provide an algorithm to construct this code

Shannon's noiseless coding theorem
 to transfer an n -bit word if it is sufficient in the limit $n \rightarrow \infty$ to send $nH(p)$ c.b.i.b, where p is the probability of a "0" (or a "1", note $H(p) = H(1-p)$) in the n -bit word

the answer is given by



The code used in the example was a bit checking on spaces were used. A code that can be used without spaces is the Huffman code (here for $P(0) = 3/4$ and $P(1) = 1/4$)

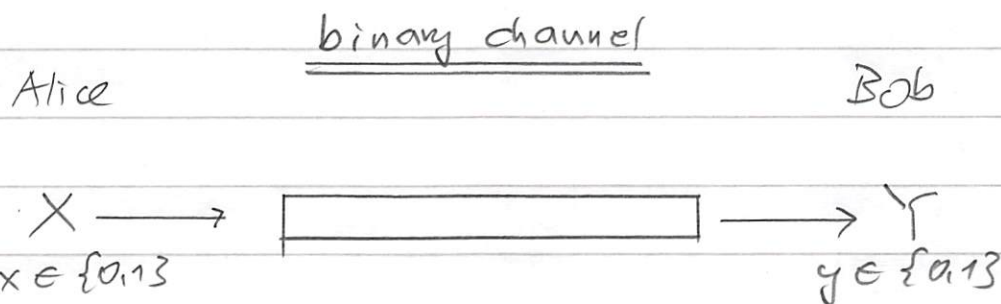
original	code
0000	10
0001	000
0010	001
0100	010
1000	011
0101	1001
0110	1010
1010	1100
1100	1101
0111	1000
1011	1111
1101	1110
1110	11110
1111	11101
11110	111001

code prefix
zeroth

original	code
0000	10
0001	000
0010	001
0100	010
1000	011
0101	11000
0110	11001
1001	11010
1010	11011
1100	11100
1101	11101
1110	11110
1111	11111
0111	111000
1011	111111
1101	111110
1110	111101
1111	1111001

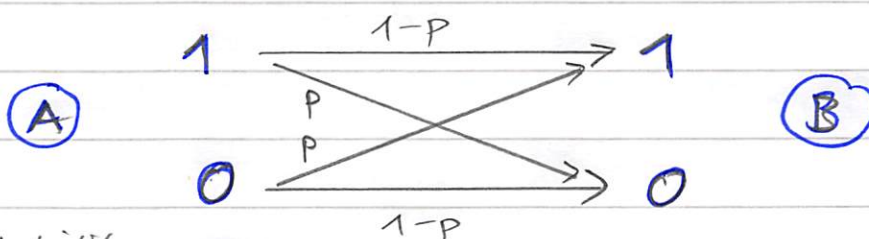
3.2 Communication channels

Shannon's noiseless coding theorem considers information transfer through ideal channels. In reality communication channels have losses and are subject to noise.



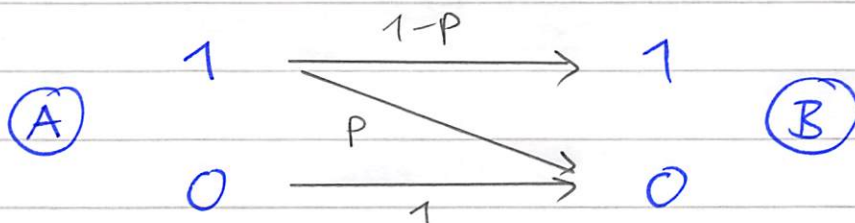
an important class of channels are memoryless channels (Markovian channels)

- Symmetric binary channel



equal error probability $p = P_{\text{error}}$
 $0 \rightarrow 1$ and $1 \rightarrow 0$

- binary decay channel



transformatrix: P_{ij} ($P = P_{\text{error}}$)

$$(8) \begin{pmatrix} 1-p & p \\ p & 1-p \end{pmatrix}$$

symmetric

$$\begin{pmatrix} 1 & p \\ 0 & 1-p \end{pmatrix}$$

decay

An important quantity to characterize channels is the mutual information

$$(9) \quad I(X; Y) = S(X) - S(X|Y) = S(Y) - S(Y|X)$$

where

$$S(Y|X) = - \sum_j p(x_j) \sum_e p(y_e | x_j) \log(p(y_e | x_j))$$

$$(10) \quad = \sum_{e,j} p(x_j, y_e) \log(p(y_e | x_j))$$

is the conditional Shannon entropy

$p(x_j, y_e)$ joint probability of $x = x_j$ and $y = y_e$

$p(y_e | x_j)$ conditional probability of $y = y_e$ given that $x = x_j$ is true

if holds

$$(11) \quad p(x, y) = p(y, x) = p(x|y) p(y) = p(y|x) p(x)$$

if x and y are independent, then

$$(12) \quad p(x|y) = p(x) \Rightarrow p(x, y) = p(x) p(y)$$

one can write the mutual information in the form

$$(13) \quad I(X; Y) = \sum_{x,y} p(x, y) \log \frac{p(x, y)}{p(x) p(y)}$$

I.e. if x and y are independent $\Rightarrow$

$$(14) \quad I(X:Y) = 0 \quad \text{and} \quad S(X|Y) = S(X)$$

thus for independent events there is no information in the knowledge of one about the other event

for a binary symmetric channel

$$(15) \quad \boxed{I(X:Y) = H(P(Y)) - H(P_{\text{error}})}$$

Def: channel capacity

$$(16) \quad C \equiv \max_{p(x)} I(X:Y)$$

Y - output
X - input

$\Rightarrow$ binary symmetric channel

$$(17) \quad \boxed{C = 1 - H(P_{\text{error}})}$$

Now let us ask the question how many bits are needed to transfer an n -bit word through a noisy channel.

transfer rate R

$$(18) \quad R = \frac{\# \text{ bits in message}}{\# \text{ bits in channel}} \quad 0 \leq R \leq 1$$

assume error probability p_b of channel

if $R = 1$ total error probability $e = p_b < 1$

to reduce error probability use code

(19) $0 \rightarrow 000$ $1 \rightarrow 111$ (repetition code)

now $R = \frac{1}{3}$. If only one of the code bits has an error the message can still be reproduced exactly, i.e. the probability of error (two or more code bits have an error)

$$(20) \quad e = p_b^3 + 3p_b^2(1-p_b) < p_b$$

m-fold repetition

$e \xrightarrow{\quad} 0$ 😊
 $R \xrightarrow{\quad} 0$ ☹️
 $m \rightarrow \infty$

Can one do better? Yes!

Shannon's noisy coding theorem (1948)

For any channel with capacity C there is a code that allows a data transfer with arbitrary small error probability and $0 < R < C$

Now we will discuss how data transfer can be achieved through a noisy channel with non-vanishing transfer rate

3.3 classical error correction

Within the last 60 years an extensive theory of error correcting codes has been developed. We will now shortly discuss an important special class: binary, linear codes

binary code k cbits encoded in $n > k$ cbits

of encodable words 2^k

of words 2^n

linear binary code

set C of codewords which form a Galois field (endlicher Körper) with respect to bitwise addition modulo 2

• example: $C = \{(0,1,1), (1,0,1), (1,1,0), (0,0,0)\}$

$$(0,1,1) \oplus (1,0,1) = (1,1,0) \in C$$

$$(0,1,1) \oplus (1,1,0) = (1,0,1) \in C$$

$$(1,0,1) \oplus (1,1,0) = (0,1,1) \in C$$

C has $n=3$ and $k=2$ basis e.g. $\begin{pmatrix} 0,1,1 \\ 1,0,1 \end{pmatrix}$

basis of C

$$\{v_1, v_2, \dots, v_k\}$$

v_i = codeword
consisting of
 $n \geq k$ cbits

every codeword can be written as superposition of $\underline{u}_i$

$$(21) \quad \underline{u} = \bigoplus_{i=1}^k \alpha_i \underline{u}_i \quad \alpha_i \in \{0,1\}$$

(bitwise addition modulo 2)

$\underline{u}$ encodes the k -bit message $(\alpha_1, \alpha_2, \dots, \alpha_k)$

The k basis vectors form the $k \times n$ generator matrix

$$(22) \quad \underline{G} = \begin{bmatrix} \underline{u}_1 \\ \vdots \\ \underline{u}_k \end{bmatrix}$$

in our example above

$$(23) \quad \underline{G} = \begin{pmatrix} 0 & 1 & 1 \\ 1 & 0 & 1 \end{pmatrix}$$

then eq. (21) can be written as

$$(24) \quad \underline{u} = \underline{\alpha} \cdot \underline{G} \quad \underline{\alpha} = (\alpha_1, \alpha_2, \dots, \alpha_k)$$

Alternatively the code can be characterized by the $n-k$ constraints

$$(25) \quad \boxed{\underline{H} \underline{u}^T = 0} \quad (\text{we leave "T" out for simplicity})$$

$\underline{H}$ is an $(n-k) \times n$ matrix called parity-check matrix

The rows of $\underline{H}$ are $(n-k)$ linear independent vectors which are orthogonal to all codewords.

In our example

$$(26) \quad \underline{H} = (1, 1, 1)$$

basis was $\{(0, 1, 1), (1, 0, 1)\}$

$$(0, 1, 1) \cdot (1, 1, 1) = 0 \oplus 1 \oplus 1 = 0 \quad (\oplus \text{ mod } 2)$$

$$(1, 0, 1) \cdot (1, 1, 1) = 1 \oplus 0 \oplus 1 = 0$$

$$\text{code: } \vec{\alpha} = (0, 0) \rightarrow (0, 0, 0)$$

$$(\text{c.g.}) \quad \vec{\alpha} = (1, 0) \rightarrow (0, 1, 1)$$

$$\vec{\alpha} = (0, 1) \rightarrow (1, 0, 1)$$

$$\vec{\alpha} = (1, 1) \rightarrow (1, 1, 0)$$

The parity-check matrix can be used to detect a possible error. For a classical channel the only possible error is a spin flip. A spin-flip error can be represented as addition with an error vector (mod 2)

$$(27) \quad \underline{u} \rightarrow \underline{u} \oplus \underline{e}$$

where $\underline{e}$ has entries "1" at the places where an error happens, e.g.,

$$(28) \quad \underline{e} = (1, 0, 1, 0, 0, 0, \dots)$$

corresponds to a spin-flip error of first and third bit.

detection of an error

$$(29) \quad \boxed{\underline{H}(\underline{u} \oplus \underline{e}) = \underbrace{\underline{H}\underline{u}}_{=0} \oplus \underline{H}\underline{e} = \underline{H}\underline{e}} \quad \text{error detection}$$

$\underline{H}\underline{e}$ is the error syndrome.

The set $\mathcal{E}$ of all errors $\{e_i\}$ which have different error syndromes can be corrected. Once the error has been unambiguously detected it is corrected by adding the same error for a second time

$$(30) \quad \boxed{(\underline{u} \oplus \underline{e}_i) \oplus \underline{e}_i = \underline{u}} \quad \text{error correction}$$

If however $\underline{H}\underline{e}_1 = \underline{H}\underline{e}_2$ we could misinterpret the error syndrome and correct wrongly!

$$(31) \quad \underline{u} \oplus \underline{e}_1 \rightarrow (\underline{u} \oplus \underline{e}_1) \oplus \underline{e}_2 \neq \underline{u}$$

Hamming distance

The Hamming distance of a code is the minimum number of bits that need to be changed to transfer one element of the code into another one.

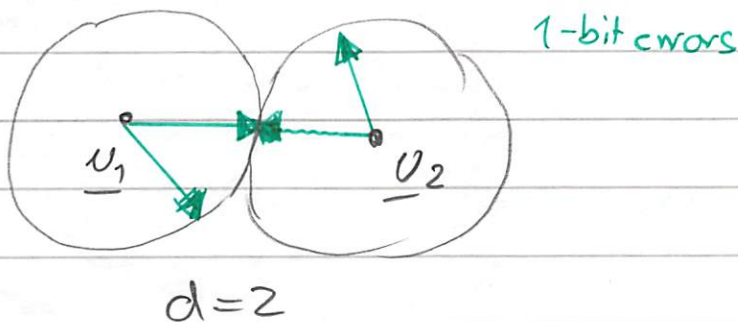
A linear binary code of Hamming distance $d=2t+1$ can correct t bit flip errors

the code introduced above has $d=2$ and can thus not correct an error

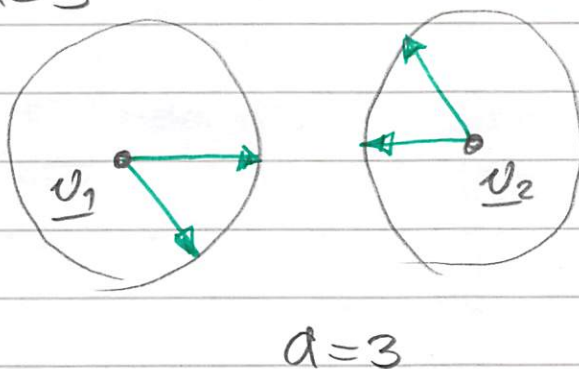
$$e_1 = (1, 0, 0) \quad e_2 = (0, 1, 0)$$

$$\begin{aligned} \underline{u}_1 \oplus \underline{e}_1 &= (0, 1, 1) \oplus (1, 0, 0) = (1, 1, 1) \\ (32) \quad &= (1, 0, 1) \oplus (0, 1, 0) = \underline{u}_2 \oplus \underline{e}_2 \end{aligned}$$

here



for $d=3$



example of an $[n=7, k=4, d=3]$ code

<u>original</u>	<u>Hamming $[n=7, k=4, d=3]$ code</u>
0000	0000000
0001	1010101
0010	0110011
0011	1100110
0100	0001111
0101	1011010
0110	0111100
0111	1101001
1000	1111111
1001	0101010
1010	1001100
1011	1110000
1101	0100101
1110	1000011
1111	0010110

the $[n=7, k=4, d=3]$ code allows to correct
all 1-bit errors in an $n=7$ bit code