

CDM

Finite Fields and Codes

KLAUS SUTNER

CARNEGIE MELLON UNIVERSITY

FALL 2024



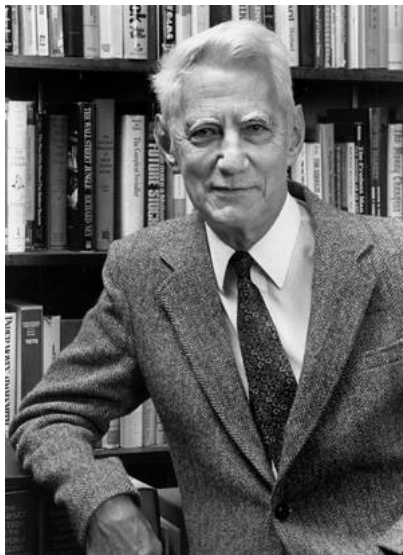
1 Information Theory

2 Codes

3 Linear Codes

4 * Cyclic Codes

5 * Hamming Codes



Two crucial papers:

A Mathematical Theory of Communication

Bell System Technical J., 27(1948)3/4, pp. 379–423, 623–656.

Communication Theory of Secrecy Systems

Bell System Technical J., 28(1949)4, pp. 656–715.

The first paper introduces the notion of a **bit** (binary digit, suggested by John Tukey). Since Shannon was working for the phone company, he was interested in how many phone calls could be transmitted in parallel on the same copper wire. His work relies heavily on probability theory.

Hard Question:

How can we assign a numerical measure to information content?

This may seem like a hopeless task, but it becomes manageable if one stays away from semantics: figuring out whether a text by Immanuel Kant has more information than a text by Albert Einstein is indeed hopeless.

Key Idea: Information has to do with surprise: the more surprising an event is, the more information it carries.

Technically, this approach allows us to study probabilities rather than embark on the hopeless enterprise of defining a semantics for natural languages.

Consider some event E that occurs with probability $\Pr[E]$. Here are some simple axioms for the **information content** $I(E)$.

$$\Pr[E] = 1 \quad \text{implies} \quad I(E) = 0$$

$$\Pr[E] = 0 \quad \text{implies} \quad I(E) = \infty$$

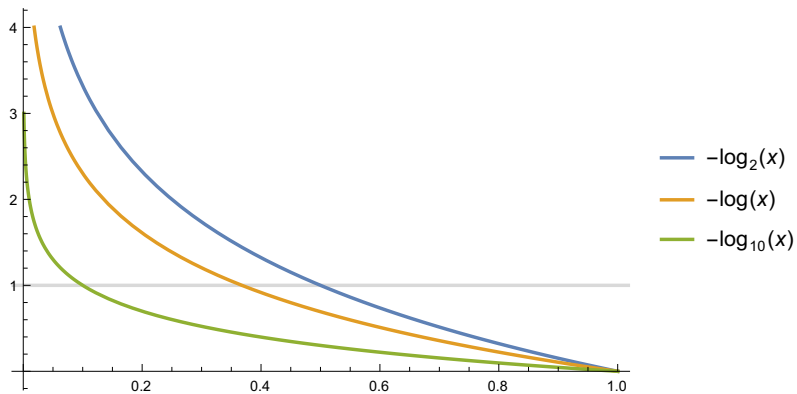
$$\Pr[E_0] \leq \Pr[E_1] \quad \text{implies} \quad I(E_0) \geq I(E_1)$$

$$E_0, E_1 \text{ indep.} \quad \text{implies} \quad I(E_0 \wedge E_1) = I(E_0) + I(E_1)$$

From these basic assumptions it follows easily that

$$I(E) = -\log_b \Pr[E]$$

Base $b = 2$ is the standard choice, corresponding to Yes/No answers.



In most applications, Alice sends a stream of symbols to Bob that constitute a message. The symbols are chosen from a fixed collection, an **alphabet** $\{a_1, a_2, \dots, a_q\}$. Let's call this scenario a **channel**. Many symbols are transmitted, so it makes sense to speak of a_i being sent with probability p_i , where $\sum_i p_i = 1$.

Key Question: How much information is contained in a message?

Well, it's the average information content of the symbols that make up the message.

This is called the **entropy** of the channel:

$$H = - \sum_i p_i \log p_i$$

Here is a simple special case: uniform probabilities $p_i = 1/q$.

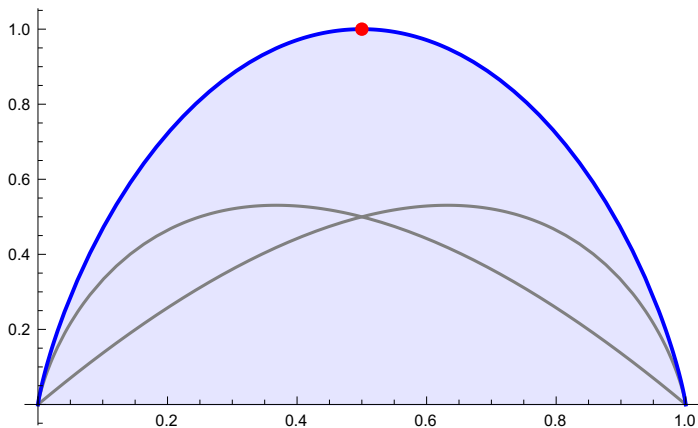
$$H = - \sum_i 1/q \log 1/q = \log q$$

In particular for $q = 2$ we have $H = 1$. Makes some intuitive sense.

How about $q = 2$ with biased probabilities?

$$H = -p \log p - (1 - p) \log(1 - p)$$

This is the **binary entropy function** (and very useful in some combinatorial arguments).



Redundancy is very closely related to entropy and measures the amount of superfluous information being transmitted. Superfluous in the sense that, in principle, we could transmit the same amount of information using fewer bits. In other words, we try to measure possible **compression**.

Why would anyone use redundant communication? Because it provides some degree of protection against noise.

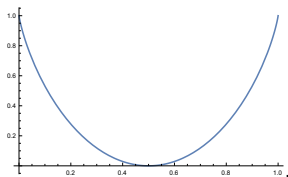
Flipping a few bits in a page of English text is usually no problem, but in a compressed file it will be a disaster.

Unsurprisingly, human speech relies heavily on redundancy to protect against noisy channels.

The redundancy for a q -symbol channel is given by

$$R = \frac{H_{\max} - H}{H_{\max}}$$

Here $H_{\max} = \log q$ represents the maximum entropy of any q -symbol channel. So $0 \leq R \leq 1$, and in the uniform case we have redundancy 0: all possible messages are equally likely.



In the binary case this is just upside-down entropy:

So what is the redundancy of English? Shannon was the first to try to come up with a reasonable answer to this question.

The problem here is that English is not a formal language, so one can try to approximate it by fixing an alphabet (say, 26 letters or 27 if one includes a space symbol) and then looking at probabilities of blocks of N consecutive letters.

onecanalmostalwaysfillinthespacesfromasequenceofwordswithnospaces

Shannon came up with an estimate of $R \approx 0.54$.

The big problem when dealing with natural languages is that one has to make lots of simplifying assumptions.

- English text consists of sequences of blocks of letters (aka words),
- the letters within a block are not equidistributed,
- the arrangement of blocks is controlled by grammar and far from random.

So all one can hope for are some rough approximations.

One approach is to determine probabilities of the next letter assuming knowledge of the current N letters (think of assigning probabilities to the edges of a de Bruijn graph):

$$a_1 a_2 \dots a_N \rightsquigarrow b$$

This is easy for small N and one can hope that the numbers converge rapidly to a limit of sorts. Shannon did concrete calculations and came up with the following estimates:

N	0	1	2	3	word
H_N	4.7	4.14	3.56	3.3	2.62

If all 3-letter blocks were equally likely we would get $H = 14.1$.

The last value for words comes from Shannon's analysis of the entries in a dictionary.

1. XFOML RXKHRJFFJUJ ZLPWCFWKCYJ FFJEYVKCQSGHYD
2. OCRO HLI RGWR NMIELWIS EU LL NBNESEBYA TH EEI
3. ON IE ANTSOUTINYS ARE T INCTORE ST BE S DEAMY ACHIN D ILONASIVE TUCOOWE AT
TEASONARE FUSO |
4. IN NO IST LAT WHEY CRATICT FROURE BIRS GROCID PONDENOME OF DEMONSTURES OF
THE RETAGIN IS REGIACTIONA OF CRE.
5. REPRESENTING AND SPEEDILY IS AN GOOD APT OR COME CAN DIFFERENT NATURAL
HERE HE THE A IN CAME THE TO OF TO EXPERT GRAY COME TO FURNISHES THE LINE
MESSAGE HAD BE THESE.
6. THE HEAD AND IN FRONTAL ATTACK ON AN ENGLISH WRITER THAT THE CHARACTER OF
THIS POINT IS THEREFORE ANOTHER METHOD FOR THE LETTERS THAT THE TIME OF WHO
EVER TOLD THE PROBLEM FOR AN UNEXPECTED.

1 Information Theory

2 **Codes**

3 Linear Codes

4 * Cyclic Codes

5 * Hamming Codes

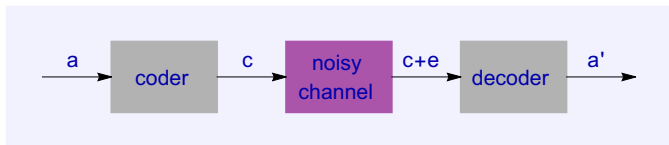
A message a is generated by some source, converted into a signal by the sender (transmitter) which is sent over the channel. The receiver turns the signal into a received message a' which reaches the destination.



Two major scenarios:

- The channel is **noisy**: $c \approx c'$.
Interesting question: how to recover the original message, **coding theory**.
- The channel is **noiseless**: $c = c'$.
Interesting question: how to protect against eaves-dropping, **cryptography**.

We will only deal with the noisy scenario here. The sender here needs to turn a message a into a code c , which is structured in such a way that the decoder has a chance to recover a even though $c \neq c'$.



The decoder receives $c' = c + e$ where e is the error introduced by the channel (for the moment, don't worry what exactly addition means here).

It outputs a decoded message a' which is hopefully identical to a . If only ...

Obviously, if the noisy channel introduces too many errors, there is no way to recover c .

Just think about a channel that returns a stream of 0s, or a sequence of random bits.

We need to make probabilistic assumptions about the likelihood of errors. As long as these probabilities are low, we want to be able to decode. If the error probabilities are large, we are really faced with an issue of channel design, not of coding theory.

Suppose we want to transmit a word over some alphabet Q of size q . This is called an q -ary channel.

In general, all we have is a probability matrix

$$P = (p_{ab}) \in [0, 1]^{q \times q}$$

where p_{ab} indicates the likelihood that symbol a sent into the channel is received at the other end as symbol b . So we would like p_{aa} to be large.

In a civilized channel, we may assume that the matrix is stochastic: all rows and columns sum to 1.

In other words, we will not deal with errors based on insertions and deletions.

Suppose again we have a q -ary channel.

Major Simplification: We assume that there is a single parameter $p \geq 0$ that describes all error probabilities: p is the probability that the channel will transmit any symbol as a different symbol.

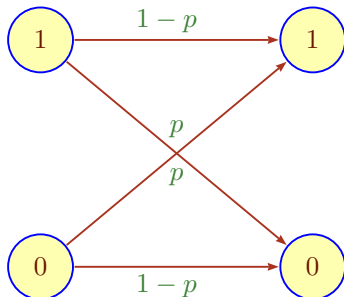
So we are sending a symbol a into the channel, and receive b at the other end. The error probability p is the same for all a and b .

$$p_{aa} = \Pr[a = b] = 1 - (q-1)p$$

$$p_{ab} = \Pr[a \neq b] = (q-1)p$$

Sanity Assumptions:

- We will always assume that p is small relative to q (see below).
- We will not deal with insertions and deletions of symbols.



$q = 2$: This is the most elementary scenario, but already very interesting

$$P = \begin{pmatrix} 1-p & p \\ p & 1-p \end{pmatrix}$$

We will think of messages and codes as words over an alphabet Q :

$$\mathbf{a} = a_1 a_2 \dots a_k$$

$$\mathbf{c} = c_1 c_2 \dots c_n$$

We are mostly interested in the binary case $Q = \mathbf{2}$, but larger alphabets can be handled similarly.

Useful Trick:

Think of Q as being some finite field $\mathbb{F}_q$.

Messages $\mathbf{a}$ and codes $\mathbf{c}$ are then simply vectors over this finite field, of length k and n , respectively.

For example, if we want to use bytes as letters, we pick $\mathbb{F}_{2^8}$.

In characteristic 2 we are just dealing with various bit operations, so efficiency should be good.

Why not simply think of the symbols as purely combinatorial objects, say, letters a , b , c , ...

Because algebra is your friend.

The whole purpose of algebra is to “trivialize” the solution of certain types of problems. If some particular problem can naturally be expressed in terms of algebra, it is probably easier to handle than otherwise.

One advantage of the vector model is that we can model the error as just another vector:

$$\mathbf{c} \mapsto \mathbf{c} + \mathbf{e}$$

The key challenges are

- Error Detection

Detect when $\mathbf{e} \neq \mathbf{0}$ (and ask for a re-transmit).

- Error Correction

Detect an error and correct it (by adding $-\mathbf{e}$ to the transmitted message).

Again, error detection and correction only work if the channel is reasonably well behaved. In practice, the major bottleneck is usually the efficiency of the decoder.

Recall that we think of messages as k -vectors over some finite field $\mathbb{F}$ of size $q = p^\ell$.

Hence we have $q^k = p^{\ell k}$ possible messages and need to assign a unique codeword to each.

Definition

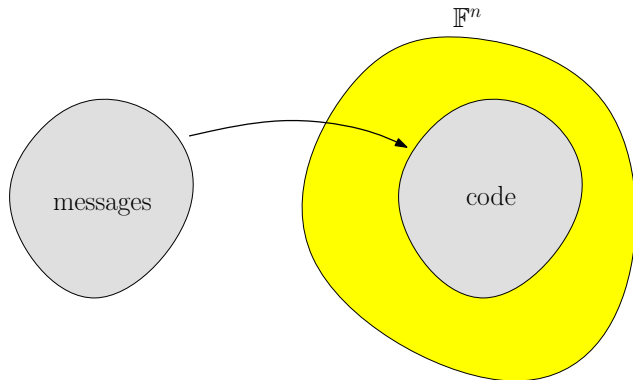
A (q -ary) **code** is a subset $C \subseteq \mathbb{F}^n$ of **size** q^k .

The elements of C are **codewords** and the ambient space $\mathbb{F}^n$ is the **codespace**.

So the key parameters of a code are the **arity** q , the **dimension** k and the **(block) length** $n^\dagger$.

Obviously the length is at least the dimension, but for error detection and correction it needs to be larger, we need to introduce some redundancy.

[†]Bad terminology, as a vector space, $\mathbb{F}^n$ has dimension n .



This picture is slightly misleading: the code words ought to be spaced out nicely in the ambient space.

Once we have fixed arity, dimension and block length, a **coding function** is an injective map

$$E : \mathbb{F}^k \longrightarrow \mathbb{F}^n$$

and a decoding function is a surjection

$$D : \mathbb{F}^n \longrightarrow \mathbb{F}^k$$

The whole point is that not just $D(E(x)) = x$, but

$$c = E(x) \quad \text{implies} \quad D(c + e) = x$$

for as many e as manageable.

To construct a code, at the very least, we need to do the following:

- choose a proper value for the length $n \geq k$,
- pick a subset C of $\mathbb{F}^n$ of cardinality q^k , and
- determine the maps E and D .

Note, though, that unlike with cryptography (and in particular public key cryptography) we only need to do this once, we can use the same code over and over again.

How large should the dimension be? We need to space out codewords so we can recover from some errors, so n should be large—but efficiency requires to keep n small.

Also, we need to specify $C \subseteq \mathbb{F}^n$ explicitly. Since q^n is large this requires some systematic tool, lookup tables are not an option.

Once the code C is fixed, we need to find efficient algorithms to compute the maps E and D .

The good news: There is quite a bit of leeway in how to organize this. Before we talk about specific codes, let's look at some fundamental properties of this setup.

There is a natural way to introduce geometry in the codespace $\mathbb{F}^n$: we can measure distances between points.

Definition

The **Hamming distance** between two vectors $\mathbf{x}, \mathbf{y} \in \mathbb{F}^n$ is defined by

$$\text{dist}(\mathbf{x}, \mathbf{y}) = |\{ i \mid x_i \neq y_i \}|$$

The **weight** (or **support**) of a vector $\mathbf{x} \in \mathbb{F}^n$ is defined by

$$w(\mathbf{x}) = |\{ i \mid x_i \neq 0 \}|$$

Thus $\text{dist}(\mathbf{x}, \mathbf{y}) = w(\mathbf{x} - \mathbf{y})$.

Exercise

Check that d really is a metric.

We are making an attempt to bring geometry to bear on our problem. Careful, though, we are not in the standard, continuous world of Euclidean space $\mathbb{R}^3$, but in some high-dimensional discrete space. Geometric intuition in high-dimensional spaces can be quite limited.

In many ways, this is just fine. E.g., it still makes perfect sense to talk about a ball of radius r centered at x :

$$B_r(x) = \{ z \in \mathbb{F}^n \mid \text{dist}(x, z) \leq r \}$$

But the radius here is integral, and the difference between open and closed balls evaporates: open radius r is closed radius $r-1$.

On the other hand, everything is discrete and we can, say, count the number of points in a ball.

A basic error in the channel is thus represented by an error vector of weight 1:

$$e = (0, \dots, 0, x, 0, \dots, 0)$$

where $x \neq 0$. In the characteristic 2 case $x = 1$, the error is a unit vector.

We can think of the channel as introducing a number of basic errors sequentially. An error vector of weight e is thus referred to as “ e errors.” This makes perfect sense in symmetric channels where each flip is independent and equally likely.

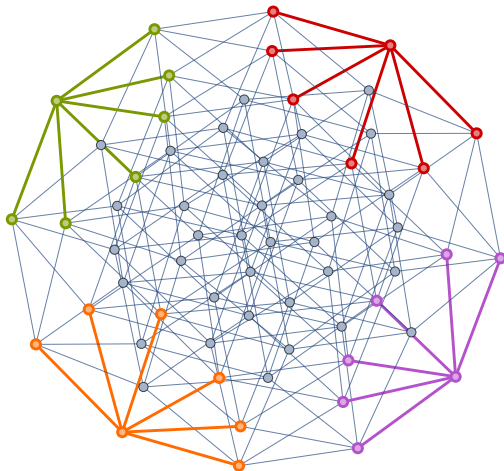
Goal: We want codes that can correct reasonably many errors.

Strategy: We will use redundancy to protect against errors.

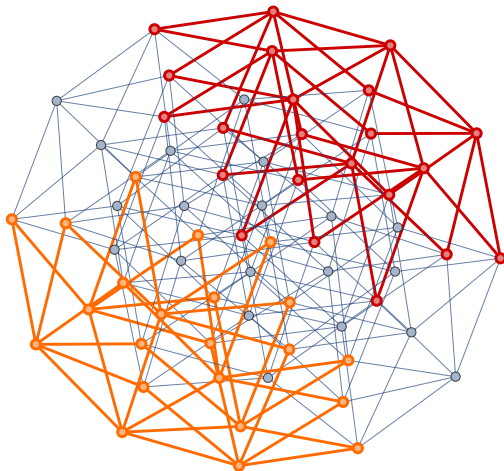
In coding theory texts, when one says a code **detects e errors**, what is actually meant is this: For any codeword c and any error $e \neq 0$ of weight at most e :

- one can detect that $z = c + e \notin C$, and
- one can determine the weight of e from z .

So this better than just realizing that there was an error, but slightly weaker than being able to correct it: we know the weight e of the error, but not necessarily the actual error e .

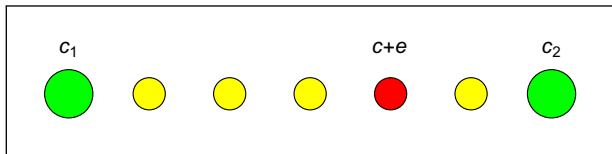


4 balls of radius 1 in the 6-dimensional hypercube 2^6 .



2 balls of radius 2 in the 6-dim hypercube 2^6 , each containing 22 points.

Suppose the received noisy transmission $c + e$ winds up “between” two codewords at distance 6. For simplicity, let’s ignore all other codewords for the time being and organize points by distance.



We have $d(c_1, c_2) = 6$ and the transmission could come from a weight-2 or a weight-4 error. We can only detect errors of weight 3 and we can correct errors of weight 2.

Suppose we have an q -ary symmetric channel with (single) error probability $(q - 1)p$.

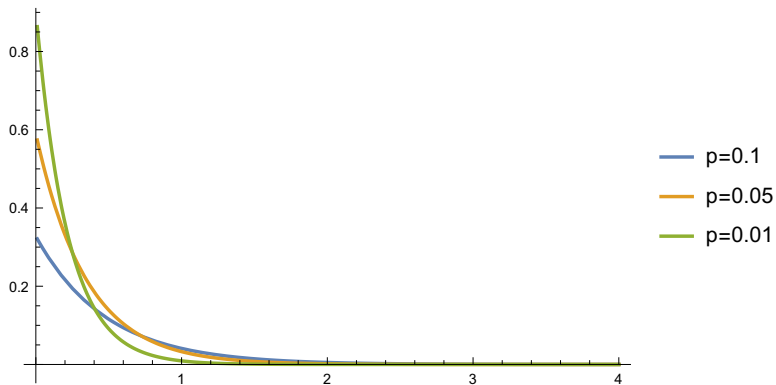
The probability that $\mathbf{y}$ is received when $\mathbf{x}$ as been transmitted is

$$\Pr[\mathbf{y} \mid \mathbf{x}] = p^d (1 - (q - 1)p)^{n-d}$$

where $d = \text{dist}(\mathbf{x}, \mathbf{y})$.

As long as $p < 1/q$ this function is sharply decreasing in d , so it is unlikely that we will hit a $\mathbf{y}$ far away from $\mathbf{x}$.

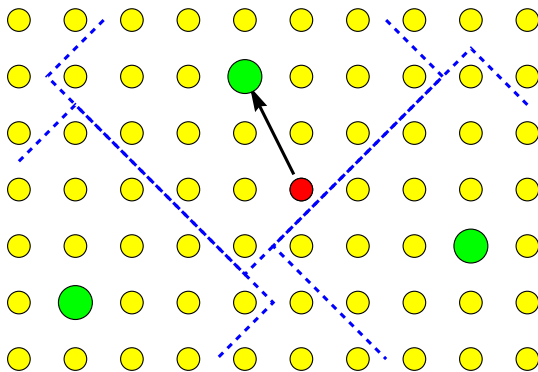
As already mentioned, if p is too large one has to redesign the channel.



So $d = 2$ is already quite unlikely, even for an error probability of 10%.

If x is transmitted and $y \notin C$ is received, the natural decoding strategy based on $\Pr[y | x]$ is to find the codeword $c \in C$ closest to y .

This makes sense in particular when this c is uniquely determined.



Suppose we have chosen a code C . We define the **minimal distance** of C to be

$$\text{md}(C) = \min(\text{dist}(x, y) \mid x \neq y \in C)$$

Thus a ball of radius $\text{md}(C) - 1$ around a codeword in C contains no other points in C .

Proposition

- *If $\text{md}(C) \geq 2e$ then C detects e errors[†].*
- *If $\text{md}(C) \geq 2e + 1$ then C corrects e errors.*

Note the active voice; we really should have said: it is possible to detect/correct—we don't yet know how to do this efficiently.

[†]Remember the meaning of “detects.”

Lemma (Hamming)

Suppose the q -ary code C has length n and corrects e errors.

Let $V = \sum_{i=0}^e \binom{n}{i} (q-1)^i$. Then

$$V |C| \leq q^n$$

Proof. Since the code corrects e errors, all balls of radius e centered at $c \in C$ are necessarily disjoint. The number of points at distance i to c is $\binom{n}{i} (q-1)^i$, so V is just the volume of a ball of radius e . $\square$

A code is **perfect** iff it realizes this bound. Needless to say, Hamming found a way to construct perfect codes.

Time for some concrete examples.

In all cases, we introduce redundancy in a more or less clever way to help with error detection and correction.

One could object that these codes are really just hacks, for a more systematic approach look at the next section.

This is a bit embarrassing, but bear with me. In $\mathbb{F}_2^n$ define E by

$$x \mapsto \underbrace{(x, x, \dots, x)}_n$$

So the only codewords are $00 \dots 0$ and $11 \dots 1$.

The minimum (and only) distance is n and balls of radius $\lfloor (n-1)/2 \rfloor$ around the codewords are still disjoint.

More importantly, it easily corrects $\lfloor (n-1)/2 \rfloor$ errors: we use a majority count to determine the original codeword.

For odd n , these codes are perfect.

The obvious fatal problem with this is the low rate of transmission: the codeword is n -times longer than the original message.

To encode, add a parity “bit” to the message:

$$x \mapsto \left(x, \sum x_i \right)$$

so that $n = k + 1$. Here we are exploiting the fact that the x_i are field elements, so the sum naturally makes sense.

A parity check code detects 1 error.

Alas, it corrects none.

This may sound next to useless, but it is not: variants of this scheme are used in ISBN, good enough for a low key sanity check.

Consider $q = 2$, dimension $k = 4$ and length $n = 7$.

To encode a 4-bit vector $\mathbf{x} = (x_1, x_2, x_3, x_4)$, we transmit 7 bits $\mathbf{c} = (c_1, c_2, c_3, c_4, c_5, c_6, c_7)$ where

$$c_3 = x_1 \quad c_5 = x_2 \quad c_6 = x_3 \quad c_7 = x_4$$

are the original message bits. The weird numbering will become clear in a moment.

The missing bits c_1, c_2, c_4 are determined by

$$c_1 = c_3 + c_5 + c_7 = x_1 + x_2 + x_4$$

$$c_2 = c_3 + c_6 + c_7 = x_1 + x_3 + x_4$$

$$c_4 = c_5 + c_6 + c_7 = x_2 + x_3 + x_4$$

We can think of the encoding process as a simple vector-matrix multiplication $x \mapsto xG$ where $G \in \mathbf{2}^{4 \times 7}$:

$$G = \begin{pmatrix} 1 & 1 & 1 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 1 & 1 & 0 & 0 \\ 0 & 1 & 0 & 1 & 0 & 1 & 0 \\ 1 & 1 & 0 & 1 & 0 & 0 & 1 \end{pmatrix}$$

G has full rank 4, so it generates a 4-dimensional subspace.

This is easy to see if we move the columns around to get an identity matrix plus the parity part.

$$\left(\begin{array}{cccc|ccc} 1 & 0 & 0 & 0 & 1 & 1 & 0 \\ 0 & 1 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 & 1 & 1 & 1 \end{array} \right)$$

OK, but why mess things up as in G ?

For decoding, calculate 3 bits α , β and γ as follows:

$$\gamma = c_1 + c_3 + c_5 + c_7$$

$$\beta = c_2 + c_3 + c_6 + c_7$$

$$\alpha = c_4 + c_5 + c_6 + c_7$$

If there is no error, $\alpha = \beta = \gamma$ and we can read off x directly.

If there is a single error, then $1 \leq \alpha\beta\gamma \leq 7$ is the binary expansion of the place where the error occurred.

Example:

$x = (1, 1, 0, 0)$ yields $c = (0, 1, 1, 1, 1, 0, 0)$.

Let $c' = (0, 1, 1, 1, 0, 0, 0)$.

Then $\alpha\beta\gamma = 101$, the error is in position 5.

The last code does not handle 2 errors in a civilized way: the following codewords have identical 1-error and 2-error versions:

$$\begin{array}{ccc} 0000000 & 0101010 & \\ & \downarrow & \downarrow \\ 0000010 & = & 0000010 \end{array}$$

We could fix this problem by adding another parity bit

$$c_0 = c_1 + c_2 + c_3 + c_4 + c_5 + c_6 + c_7$$

This extended code detects two errors (recall the definition) and still corrects one as before.

1 Information Theory

2 Codes

3 **Linear Codes**

4 * Cyclic Codes

5 * Hamming Codes

So far, everything is a bit ad-hoc. Here comes the first real idea:

Use linear algebra to define codes and handle coding/decoding.

So instead of some clever hack, we would like C to be determined by some algebraic definition.

Hopefully this approach will automatically take care of spacing out the the codewords, and provide some computational support for coding/decoding, including error correction.

Definition

A **linear code** is a linear subspace $C \subseteq \mathbb{F}^n$.

A linear code of length n and dimension k is an $[n, k]$ **code**.

A **generator matrix** G for C is a k by n matrix over $\mathbb{F}_q$ whose rows form a basis for the code space C .

In other words, the subspace C is the row space of G :

$$C = \{ \mathbf{x} G \mid \mathbf{x} \in \mathbb{F}^k \}$$

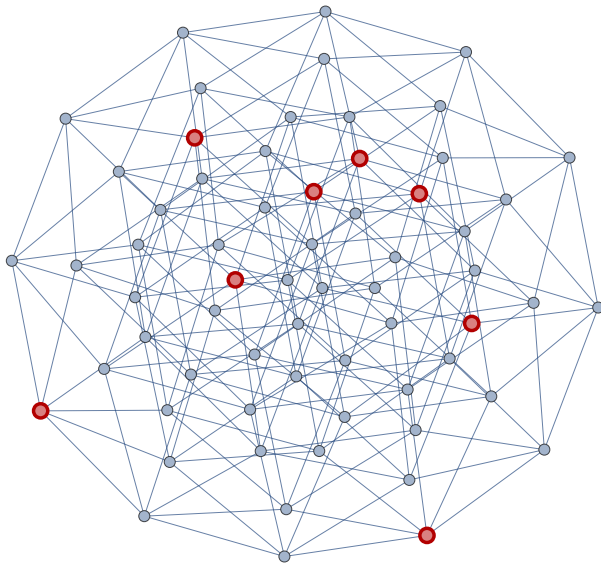
Let $q = 2$, $k = 3$, $n = 6$ and

$$G = \begin{pmatrix} 1 & 0 & 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 1 & 1 & 0 \end{pmatrix}$$

The encoding map is

$$\mathbf{x} \mapsto \mathbf{y} = (x_1, x_2, x_3, x_2 + x_3, x_1 + x_3, x_1 + x_2)$$

Decoding in the absence of errors is entirely trivial since $\mathbf{x}$ is just the first 3 components of $\mathbf{y}$, but what if the decoder receives $\mathbf{y} + \mathbf{e}$ instead?



Note that our generator matrix has the form

$$G = (I_k \ P) \in \mathbb{F}^{k \times n}$$

where I_k is the identity matrix of order k , so G is in reduced echelon form and has full rank $n-k$. Such codes are said to be in **standard form**. Alternatively the columns of I_k can be spread out in G ; this is called **systematic form**. This is essentially the same, but recall the hack from above.

Correspondingly, in a standard form code, one calls

- the first k bits **information symbols**
- the last $n - k$ bits **parity check symbols**

In this sense all linear codes are parity check codes.

Suppose we are in characteristic 2. Let $G = (I_k \ P)$ be the $k \times n$ generator matrix of a standard form $[n, k]$ linear code. Define the $(n-k) \times n$ matrix H by

$$H = (P^T \ I_{n-k})$$

Definition

H is the **parity check matrix** for code C .

The vector $\text{syn}(z) = z \cdot H^T \in \mathbb{F}^{n-k}$ is the **syndrome** of $z \in \mathbb{F}^n$.

Since we are working in characteristic 2 we have $G \cdot H^T = \mathbf{0}$ so that

$$z \in C \iff \text{syn}(z) = \mathbf{0}$$

and we can test membership in C via a simple vector-matrix multiplication.

For the $[6, 3]$ code from above we have

$$G = \begin{pmatrix} 1 & 0 & 0 & 0 & 1 & 1 \\ 0 & 1 & 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 1 & 1 & 0 \end{pmatrix}$$

and

$$H = \begin{pmatrix} 0 & 1 & 1 & 1 & 0 & 0 \\ 1 & 0 & 1 & 0 & 1 & 0 \\ 1 & 1 & 0 & 0 & 0 & 1 \end{pmatrix}$$

In this case, the P -part of G is symmetric and thus invariant under transpose.

For our example, the syndrome looks like $\text{syn}(\mathbf{y}) = (s_1, s_2, s_3)$ where

$$s_1 = e_2 + e_3 + e_4$$

$$s_2 = e_1 + e_3 + e_5$$

$$s_3 = e_1 + e_2 + e_6$$

and here are the syndromes associated with 1-bit errors (really H^T):

e_1	0	1	1
e_2	1	0	1
e_3	1	1	0
e_4	1	0	0
e_5	0	1	0
e_6	0	0	1

So we can fix all 1-bit errors.

Decoding here hinges on the following observation (which is not hard to prove):

Claim

If $(s_1, s_2, s_3) \neq (1, 1, 1)$ there is a unique choice for such an error vector e (of weight at most 1).

Alas, if $(s_1, s_2, s_3) = (1, 1, 1)$ all the following 3 weight 2 vectors work:
 $(1, 0, 0, 1, 0, 0)$, $(0, 1, 0, 0, 1, 0)$, $(0, 0, 1, 0, 0, 1)$.

In this case we can either ask for a retransmit, or we can simply pick one of the 3 possibilities—which will be right 1/3 of the time.

How much have we gained?

If we transmit 3 bits in a symmetric binary channel with error probability $p = 0.001$, the likelihood of correct transmission is 0.997003.

But if we use our $[6, 3]$ code, the likelihood of a correct transmission is

$$\begin{aligned}(1-p)^6 + 6(1-p)^5p + (1-p)^4p^2 = \\ 0.994015 + 0.00597006 + 9.9600610^{-7} = 0.999986\end{aligned}$$

corresponding to $s = 0$, $0 < s < 1$ and $s = 1$.

The new error probability is smaller by a factor of about 214, at a cost of having to send twice as many bits. A nice improvement.

According to our definition the minimum distance is the least $\text{dist}(x, y)$ where $x \neq y \in C$. In a linear code we can do better.

Lemma

For a linear code we have

$$\text{md}(C) = \min(w(x) \mid 0 \neq x \in C)$$

This may not look too impressive, but at least it yields a linear time method as opposed to the obvious quadratic one that works for an any code.

Let H be the check matrix of linear code C with columns $\mathbf{h}_1, \mathbf{h}_2, \dots, \mathbf{h}_n$.

Then for $\mathbf{z} = \mathbf{c} + \mathbf{e}$ we have $\text{syn}(\mathbf{z})^T = H \mathbf{e}^T = \sum \mathbf{h}_i e_i$.

This linear combination is 0 iff $\mathbf{z} \in C$.

Hence we have the following alternative (and often superior) way to compute the minimum distance.

Lemma

The minimum distance of C is the least d such that there exists a set of d linearly dependent column vectors in H .

Definition

For $\mathbf{x}, \mathbf{y} \in \mathbb{F}^n$ define the **inner product**

$$\mathbf{x} \cdot \mathbf{y} = \sum x_i y_i$$

Given a subset $X \subseteq \mathbb{F}^n$ define its **orthogonal complement** by

$$X^\perp = \{ \mathbf{y} \in \mathbb{F}^n \mid \forall \mathbf{x} \in X (\mathbf{x} \cdot \mathbf{y} = 0) \}$$

This is similar to the notion of orthogonal complements in real spaces but the geometry here is more complicated. For example, it may happen that $X^\perp = X$: consider $X = \{ (x_1, x_1, x_2, x_2, \dots, x_k, x_k) \mid \mathbf{x} \in \mathbb{F}_2^k \}$.

Our codes are linear subspaces

$$C = \{ \mathbf{x}G \mid \mathbf{x} \in \mathbb{F}^k \}$$

and G is an $k \times n$ matrix of full rank (in this case rank k).

Alternatively we can use the $(n - k) \times n$ parity check matrix to describe C like so:

$$C = \{ \mathbf{z} \in \mathbb{F}^n \mid \mathbf{z}H^T = \mathbf{0} \}$$

So C is the space orthogonal to another which is described by the parity check matrix.

We can push this idea one step further:

Parity check matrices are actually generators for a closely related code.

Definition

Let C be a linear $[n, k]$ code. The **dual code** is $C^\perp$.

- The dual code has dimension $n - k$.
- The dual of the dual is the code: $C^{\perp\perp} = C$.
- If C has generator matrix G then $C^\perp$ has parity check matrix G .
- If C has parity check matrix H then $C^\perp$ has generator matrix H .

How do we decode a linear code systematically?

Suppose we receive $z = c + e$. Note that

$$z + C = e + C \subseteq \mathbb{F}^n$$

since $c \in C$ and C is a linear subspace. Hence, the e lies in the same coset as z . We can decompose $\mathbb{F}^n$ into $n - k$ disjoint cosets (affine subspaces)

$$C, a_2 + C, \dots, a_{n-k} + C$$

Two vectors x and y are in the same coset iff they have the same syndrome:

$$\text{syn}(x) = \text{syn}(y) \iff x - y \in C$$

Definition

A **coset leader** is an element of the coset of minimal weight.

So if we want to do maximum likelihood decoding we should pick a minimum weight vector in the coset, the coset leader, as the error vector.

Unfortunately, “minimum weight” is not a property expressible in terms of linear algebra, it adds a component of combinatorics. So we have to precompute a minimum weight vector for each possible syndrome.

For our standard $[6, 3]$ example we have the following coset leaders, parametrized by syndrome (the s -vectors from above):

syndrome	leader
000	000000
001	000001
010	000010
100	000100
011	100000
101	010000
110	001000
111	001001, 010010, 100100

As we have seen, the syndrome function is constant on each coset and differs on each coset (if you like, the coset partition is just the kernel relation induced by the syndrome function).

But then it suffices to store a **syndrome table**: for each syndrome store the corresponding coset leader.

To decode $z = c + e$:

- compute $s = \text{syn}(z)$,
- look up the corresponding error vector e ,
- return $z - e$, presumably just the message c .

1 Information Theory

2 Codes

3 Linear Codes

4 * Cyclic Codes

5 * Hamming Codes

Definition

A linear code C is **cyclic** if for all $c \in C$ the cyclic shift of c is also in C .

As we will see, cyclic codes have a very simple description (a single generator). They are used mostly for error detection: **cyclic redundancy check (CRC)**. A retransmit is requested if an error is found.

Example

$C = \{000, 011, 101, 110\}$ is cyclic code of length 3.

This generalizes to all words in $\mathbb{F}^n$ with even weight.

Instead of thinking about vectors $\mathbf{a} \in \mathbb{F}^n$ we will interpret codewords as polynomials, so-called **code polynomials**:

$$a(x) = a_0 + a_1x + \dots + a_{n-1}x^{n-1}$$

Write $\mathbf{a}'$ for the cyclic right-shift of $\mathbf{a}$ and note that the corresponding code polynomial has the form

$$a_{n-1} + a_0x + a_1x^2 + \dots + a_{n-1}x^n$$

But then $a'(x) = x \cdot a(x) \bmod (x^n - 1)$.

So if we think of code polynomials as elements of the quotient ring $\mathbb{F}_q[x]/(x^n - 1)$, then the combinatorial shift operation corresponds simply to the algebraic operation “multiplication by x .”

Is there a good description for cyclic codes, viewed as subsets of $\mathbb{F}_q[x]/(x^n - 1)$?

Lemma

A linear code C is cyclic iff C is an ideal in $\mathbb{F}_q[x]/(x^n - 1)$.

Proof. If C is cyclic then the set of corresponding code polynomials is closed under multiplication by x , and hence under multiplication by any polynomial. Closure under addition follows from C being a code.

In the opposite direction, C must be additively closed since the ideal is so closed. And C must be invariant under shifts since the ideal is closed under multiplication by x . □

It is known from algebra that $\mathbb{F}_q[x]/(x^n - 1)$ is a **principal ideal domain**: all ideals in this ring are principal ideals. Thus, $C = (g(x)) = \mathbb{F}_q[x] g(x)$ for some polynomial $g(x)$.

More precisely, the **generator polynomial** $g(x)$ is a monic polynomial in $C - \{0\}$ of minimal degree (this polynomial is uniquely determined by C).

Note that C has dimension $n - \deg(g)$.

Once we have a generator polynomial $g(x)$, coding is very easy:

$$a \mapsto \left(\sum a_i x^i \right) g(x)$$

This can be implemented with the help of dedicated hardware (feedback shift register), so coding can be very fast.

The generator polynomial $g(x)$ must divide $x^n - 1$: otherwise $x^n - 1 = g(x)h(x) + r(x)$ where $\deg(r) < \deg(g)$. But $r(x) \in C$, contradiction.

So let $g(x)h(x) = x^n - 1$, so g and h are zero-divisors in $\mathbb{F}_q[x]/(x^n - 1)$.

But then

$$c \in C \iff c(x)h(x) = 0 \pmod{x^n - 1}$$

More generally, we can compute the syndrome of a received message y by multiplying $y(x)$ and $h(x)$, and reducing modulo $x^n \mapsto 1$. Note that this reduction is easier than the ones we saw in connection with finite fields.

From now on assume that the length of the code n and p , the characteristic of the underlying field, are coprime. In this case $x^n - 1$ factors into distinct irreducible polynomials:

$$x^n - 1 = f_1(x) f_2(x) \dots f_{t-1}(x) f_t(x)$$

We can produce 2^t generators by choosing $I \subseteq [t]$:

$$g(x) = \prod_{i \in I} f_i(x)$$

A maximal cyclic code has $I = \{i\}$: we pick exactly one of these irreducible factors.

Dually, a minimal cyclic code has $I = [n] - \{i\}$, we choose all but one.

Let $q = 2$ and $n = 9$. Then

$$x^9 + 1 = (x + 1)(x^2 + x + 1)(x^6 + x^3 + 1)$$

The maximum code based on $g = x^6 + x^3 + 1$ has codewords

$$(c_0, c_1, c_2, c_0, c_1, c_2, c_0, c_1, c_2)$$

and minimum distance 3.

The minimum code based on $h = (x + 1)(x^2 + x + 1) = x^3 + 1$ is a $[9, 6]$ code with minimum distance 2; it has codewords

$$(c_0, c_1, c_2, c_4, c_5, c_6, c_0 + c_4, c_1 + c_5, c_2 + c_6)$$

Note that these codewords are in fact cyclic (which is not entirely obvious from the definition).

We can compute a cyclic code by polynomial multiplication
 $c(x) = a(x)g(x) \bmod (x^n - 1)$.

Slightly better performance can be obtained by coding

$$\mathbf{a} \mapsto (\mathbf{a}, -\mathbf{s})$$

where $\mathbf{s}$ is determined by

$$x^{n-k}a(x) = f(x)g(x) + s(x)$$

The reason this is slightly better is that one can send the first k information symbols while computing the following $n - k$ parity symbols.

1 Information Theory

2 Codes

3 Linear Codes

4 * Cyclic Codes

5 * **Hamming Codes**



The purpose of computation is insight, not numbers.

Suppose H is the check matrix of a binary $[n, k]$ code.

Note that the syndrome of a received word $\mathbf{x} = \mathbf{c} + \mathbf{e}$ is the sum of the columns of H in positions where an error occurred:

$$\text{syn}(\mathbf{x}) = \text{syn}(\mathbf{c}) + \text{syn}(\mathbf{e}) = H \mathbf{e}^T$$

In the special case where there is only **one** error we just get a single column in H . If the columns are all distinct this tells us the position of the error and we can correct it.

Let $Q = 2$, $k = 4$, $n = 7$ and define a code C by

$$H = \begin{pmatrix} 1 & 1 & 0 & 1 & 1 & 0 & 0 \\ 1 & 0 & 1 & 1 & 0 & 1 & 0 \\ 0 & 1 & 1 & 1 & 0 & 0 & 1 \end{pmatrix} \quad G = \begin{pmatrix} 1 & 0 & 0 & 0 & 1 & 1 & 0 \\ 0 & 1 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 1 & 0 & 0 & 1 & 1 \\ 0 & 0 & 0 & 1 & 1 & 1 & 1 \end{pmatrix}$$

H is organized so as to make G simple (standard form, the first 4 bits are information bits). Note, though, that every 3-bit word other than $\mathbf{0}$ appears as a column in H .

It is not hard to see that $\text{md}(C) = 3$, so we can correct one error by computing the syndrome.

A little care is needed to lift the binary example to the q -ary case.

Definition

A linear code over $\mathbb{F}_q$ is a **Hamming code of redundancy r** if it has a check matrix H that contains a unique column ax for all $0 \neq x \in \mathbb{F}^r$ and some $0 \neq a \in \mathbb{F}$.

Thus every non-zero vector in $\mathbb{F}^r$ appears as exactly one column in H , up to multiplication by a non-zero scalar.

One can determine the multiplicative factor for example by insisting that the first non-zero entry in each column be 1.

Theorem

Let C be a Hamming code of redundancy $r \geq 2$. Then C is a

$$\left[\frac{q^r - 1}{q - 1}, \frac{q^r - 1}{q - 1} - r \right]$$

code. The minimum distance is 3.

Proof. There are $q^r - 1$ non-zero vectors of length r over $\mathbb{F}$. Since each column in H is a non-zero multiple there are $n = (q^r - 1)/(q - 1)$ columns, the length of the code. The dimension then is $n - r$.

The claim about distance follows from the lemma about linear dependence of columns in H above. $\square$

We can add a parity bit to the Hamming $[7, 4]$ to get a $[8, 4]$ code. This code

- has minimum distance 4,
- corrects 1 error,
- detects 2 errors (in the strong sense above, 3 in the weak sense).

This type of code is sometimes called a SECDED code (single error correction, double error detection) and used in memory systems.

So far we have focused on codewords being sufficiently far from each other. Another reasonable condition is that no point in codespace should be far from a codeword.

Definition

A code with minimum distance $2e + 1$ is **perfect** if for all $z \in \mathbb{F}^n$ the ball of radius e contains exactly one codeword:

$$B_e(z) \cap C = \{c\}.$$

In other words, a code is perfect if we have equality in Hamming's bound.

Here is a bad example for a perfect code: the repetition code of length $n = 2e + 1$ has $\rho(C) = e$ and $\text{md}(C) = 2e + 1$.

Recall that in the q -ary case we have length $n = (q^k - 1)/(q - 1)$ and dimension $n - r$.

Claim

Every Hamming code is perfect.

We have minimum distance 3, so $e = 1$. We need to check that for any sphere S of radius 1 we have $|C| \cap S = q^n$. But the latter is equivalent to

$$q^{n-r}(1 + (q - 1)n) = q^n$$

which is easily verified.