

# Telecom Network Systems

①

16/06/2015

---

Q1 Nyquist rate I

$$\begin{array}{ccc} SR & = & 2W \\ \uparrow & & \uparrow \\ \text{symbol} & & \text{bandwidth} \\ \text{rate} & & \end{array}$$

Nyquist rate II

$$\begin{array}{ccc} SR & = & 2f \\ \uparrow & & \uparrow \\ \text{sample} & & \text{highest frequency} \\ \text{rate} & & \text{component of signal} \end{array}$$

Q2 See lecture notes (p.10 & p.11 of ch.0)

Q3 Note there are 101 coins that can be H or

L, so total entropy =  $\log(N) = \log(202) =$

7.66 bits. This cannot be measured in

two weighing. Remember: each weighing

with 3 possibilities has max  $\sum p \ln p =$

$\sum \frac{1}{3} \log(\frac{1}{3}) = 1.58$  bits. Two weighing is

max 3.17 bits. Not enough! But we need

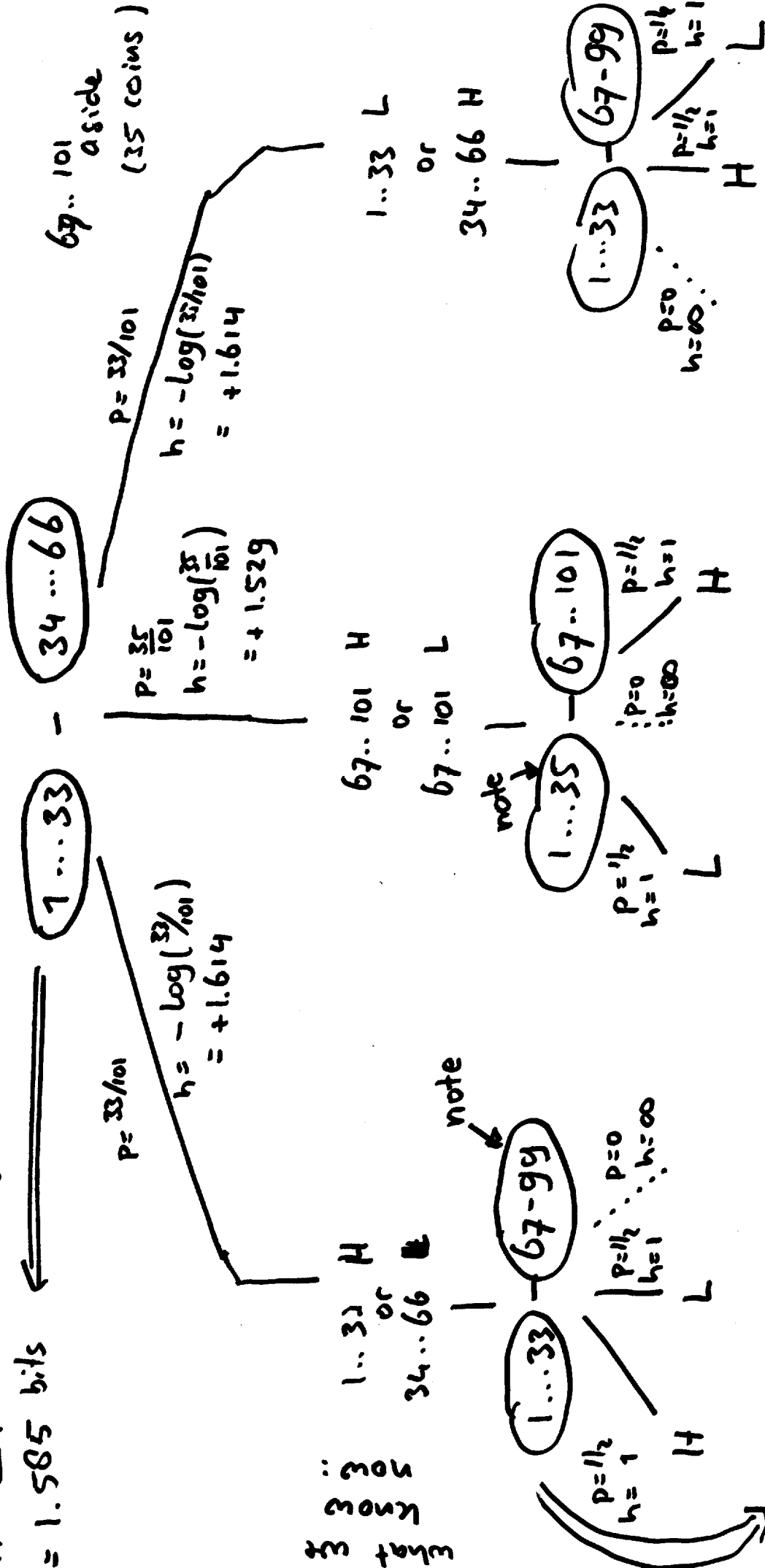
only 1 bit, namely  $H = P_L \log P_L +$

$P_H \log P_H = \frac{1}{2} \log \frac{1}{2} + \frac{1}{2} \log \frac{1}{2} = 1$  bit!

Total entropy of system:

$$H = \sum p \cdot h = - \sum p \log(p)$$

$$\downarrow H = \sum p \cdot y = - \sum p \log(p)$$

$$= 1.585 \text{ bits}$$


$$H = \sum p^2 = 1$$

$$H = \sum p_k = 1$$

1  
1  
5  
e  
W  
1  
I

with each weighing we get, actually, (3)  
 too much information (!). Note that in  
 the first weighing we get no useful information  
 whatsoever in terms of entropy. The probability  
 of L and H after first weighing is still  
 50% and 50%. Now the philosophical question:  
 why do we measure, if it gives no useful infor-  
 mation?

Q4 see forward probability (ch. 1, p. 15 of  
 lecture notes).

R<sup>5</sup>. bit flip probability =  $f = 0.01$ . Imagine  
 sending  $0 \rightarrow 00000$  in physical channel

Probability of 0 bit flips =  $(1-f)^5$  and that  
 will be decoded ok.

Probability of 1 bit flip =  $(1-f)^4 \cdot f \cdot \binom{5}{1} \rightarrow \text{OK}$

Probability of 2 bit flips =  $(1-f)^3 \cdot f^2 \cdot \binom{5}{2} \rightarrow \text{OK}$

Probability of 3 bit flips =  $(1-f)^2 \cdot f^3 \cdot \binom{5}{3} \rightarrow \text{ERROR}$

Probability of 4 bit flips =  $(1-f) \cdot f^4 \cdot \binom{5}{4} \rightarrow \text{ERROR}$

Probability of 5 bit flips =  $f^5 \rightarrow \text{ERROR}$

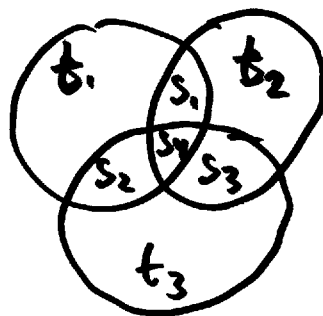
$$\text{BER} = \frac{\sum \text{ERROR}}{\sum \text{ERROR} + \text{OK}} = \frac{(1-f)^2 \cdot f^3}{1}$$

$$= \frac{(1-f)^2 \cdot f^3 \cdot \binom{5}{3} + (1-f) \cdot f^4 \cdot \binom{5}{4} + f^5}{5} \approx 10 f^3 = 10^{-5}$$

Q5 see lecture notes P.18 - P.19 of ch. 1 (4)

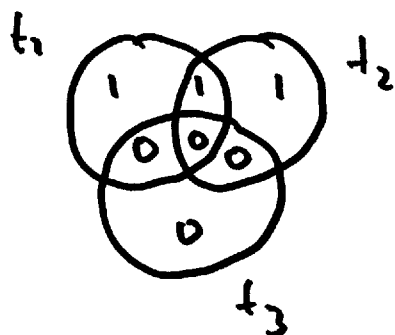
4 source code bit  $s_1, s_2, s_3, s_4$

3 parity bits  $t_1, t_2, t_3$



every circle needs even parity.  
choose  $t$  such that it is so

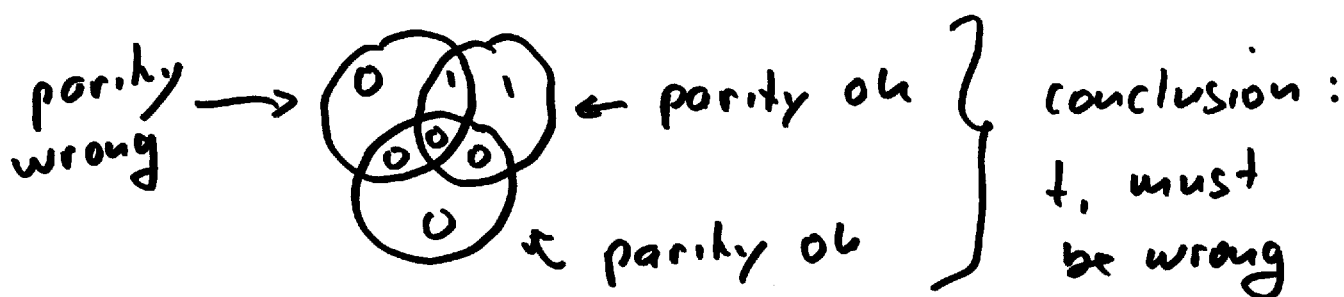
Example :  $s_1, s_2, s_3, s_4 = 1000$



$$t_1, t_2, t_3 = 110$$

Sent : 1000 110

Decoding. imagine  $t_1$  is wrong

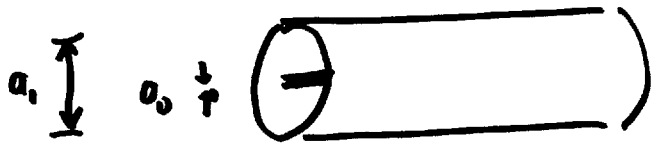


FEC is needed when we want error free uni directional communication (as in TV and radio). Simplex communication

Q6 see p. 13 of ch. 2

(5)

$$\left. \begin{aligned} \iiint \nabla \cdot \vec{F} \, dV &= \oiint \vec{F} \cdot \vec{n} \, dS \\ \nabla \cdot \vec{D} &= \rho \end{aligned} \right\} \Rightarrow \oiint \vec{D} \cdot \vec{n} = \iiint \rho \, dV = Q$$



Place  $Q$  on core and  $-Q$  on shield

A cylinder with radius  $r$  that is around core contains  $Q$

$$\oiint \vec{D} \cdot \vec{n} = Q$$

$$2\pi r L \epsilon E = Q \Rightarrow E = \frac{Q}{2\pi L \epsilon r}$$

$$V(r) = \int E(r) \, dr = \frac{Q}{2\pi L \epsilon} \ln(r) - V_0$$

$$V(a_0) = 0 \Rightarrow V(r) = \frac{Q}{2\pi L \epsilon} \ln\left(\frac{r}{a_0}\right)$$

$$V(a_1) = \frac{Q}{2\pi L \epsilon} \ln\left(\frac{a_1}{a_0}\right)$$

$$C = \frac{Q}{V} = \frac{2\pi \epsilon L}{\ln(a_1/a_0)}$$

Q7

We have been seeking a mathematical model of a source of English text. Such a model should be capable of producing text which corresponds closely to actual English text, closely enough so that the problem of encoding and transmitting such text is essentially equivalent to the problem of encoding and transmitting actual English text. The mathematical properties of the model must be mathematically defined so that useful theorems can be proved concerning the encoding and transmission of the text it produces, theorems which are applicable to a high degree of approximation to the encoding of actual English text. It would, however, be asking too much to insist that the production of actual English text conform with mathematical exactitude to the operation of the model.

The mathematical model which Shannon adopted to represent the production of text (and of spoken and visual messages as well) is the *ergodic source*. To understand what an ergodic source is, we must first understand what a *stationary source* is, and to explain this is our next order of business.

The general idea of a stationary source is well conveyed by the name. Imagine, for instance, a process, i.e., an imaginary machine, that produces forever after it is started the sequences of characters

A E A E A E A E A E, etc.

Clearly, what comes later is like what has gone before, and *stationary* seems an apt designation of such a source of characters. We might contrast this with a source of characters which, after starting, produced

A E A A E E A A A E E E, etc.

Here the strings of A's and E's get longer and longer without end; certainly this is not a stationary source.

Similarly, a sequence of characters chosen at random with some assigned probabilities (the first-order letter approximation of example 1 above) constitutes a stationary source and so do the digram and trigram sources of examples 2 and 3. The general idea of a stationary source is clear enough. An adequate mathematical definition is a little more difficult.

The idea of stationarity of a source demands no change with time. Yet, consider a digram source, in which the probability of the second character depends on what the previous character is. If we start such a source out on the letter A, several different letters can follow, while if we start such a source out on the letter Q, the second letter must be U. In general, the manner of starting the source will influence the statistics of the sequence of characters produced, at least for some distance from the start.

To get around this, the mathematician says, let us not consider just one sequence of characters produced by the source. After all, our source is an imaginary machine, and we can quite well imagine that it has been started an infinite number of times, so as to produce an infinite number of sequences of characters. Such an infinite number of sequences is called an *ensemble* of sequences.

These sequences could be started in any specified manner. Thus,

in the case of a digram source, we can if we wish start a fraction, 0.13, of the sequences with E (this is just the probability of E in English text), a fraction, 0.02, with W (the probability of W), and so on. *If we do this*, we will find that the fraction of E's is the same, averaging over all the *first* letters of the ensemble of sequences, as it is averaging over all the *second* letters of the ensemble, as it is averaging over all the *third* letters of the ensemble, and so on. No matter what position from the beginning we choose, the fraction of E's or of any other letter occurring in that position, taken over all the sequences in the ensemble, is the same. This independence with respect to position will be true also for the probability with which TH or WE occurs among the first, second, third, and subsequent *pairs* of letters in the sequences of the ensemble.

This is what we mean by stationarity. If we can find a way of assigning probabilities to the various starting conditions used in forming the ensemble of sequences of characters which we allow the source to produce, probabilities such that any statistic obtained by averaging over the ensemble doesn't depend on the distance from the start at which we take an average, then the source is said to be stationary. This may seem difficult or obscure to the reader, but the difficulty arises in giving a useful and exact mathematical form to an idea which would otherwise be mathematically useless.

In the argument above we have, in discussing the infinite ensemble of sequences produced by a source, considered averaging over-all *first* characters or over-all *second* or *third* characters (or pairs, or triples of characters, as other examples). Such an average is called an *ensemble* average. It is different from a sort of average we talked about earlier in this chapter, in which we lumped together all the characters in *one* sequence and took the average over them. Such an average is called a *time* average.

The time average and the ensemble average can be different. For instance, consider a source which starts a third of the time with A and produces alternately A and B, a third of the time with B and produces alternately B and A, and a third of the time with E and produces a string of E's. The possible sequences are

1. A B A B A B A B, etc.
2. B A B A B A B A, etc.
3. E E E E E E E E, etc.



We can see that this is a stationary source, yet we have the probabilities shown in Table V.

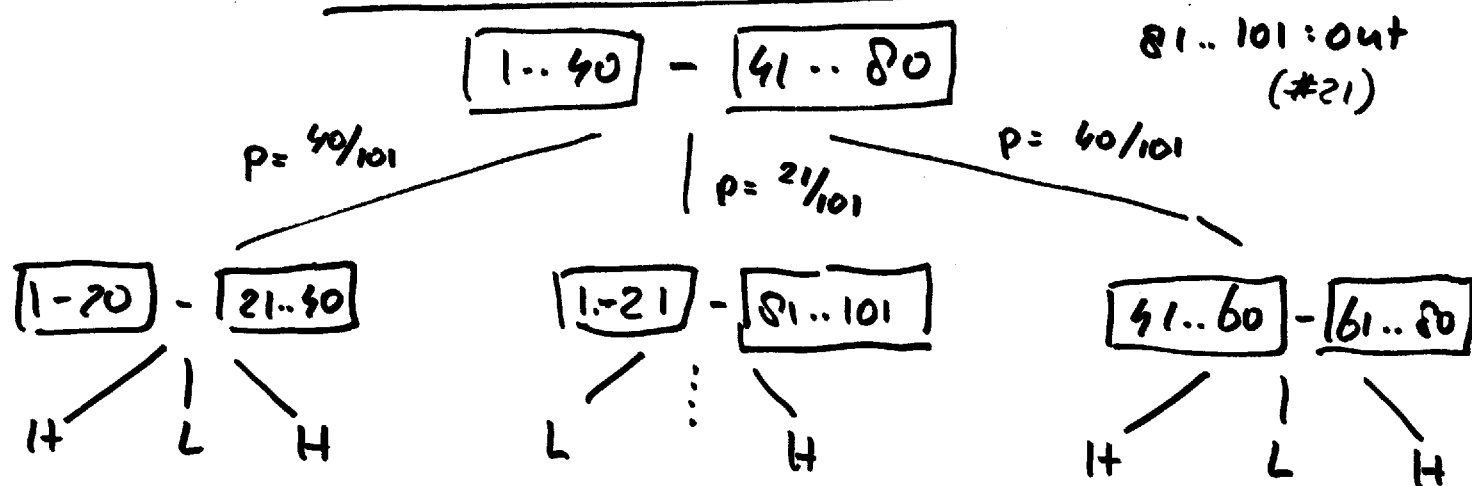
TABLE V

| <i>Probability<br/>of</i> | <i>Time Average<br/>Sequence (1)</i> | <i>Time Average<br/>Sequence (2)</i> | <i>Time Average<br/>Sequence (3)</i> | <i>Ensemble<br/>Average</i> |
|---------------------------|--------------------------------------|--------------------------------------|--------------------------------------|-----------------------------|
| A                         | $\frac{1}{2}$                        | $\frac{1}{2}$                        | 0                                    | $\frac{1}{3}$               |
| B                         | $\frac{1}{2}$                        | $\frac{1}{2}$                        | 0                                    | $\frac{1}{3}$               |
| E                         | 0                                    | 0                                    | 1                                    | $\frac{1}{3}$               |

When a source is stationary, and when every possible ensemble average (of letters, digrams, trigrams, etc.) is equal to the corresponding time average, the source is said to be ergodic. The theorems of information theory which are discussed in subsequent chapters apply to ergodic sources, and their proofs rest on the assumption that the message source is ergodic.<sup>1</sup>

# Question 3

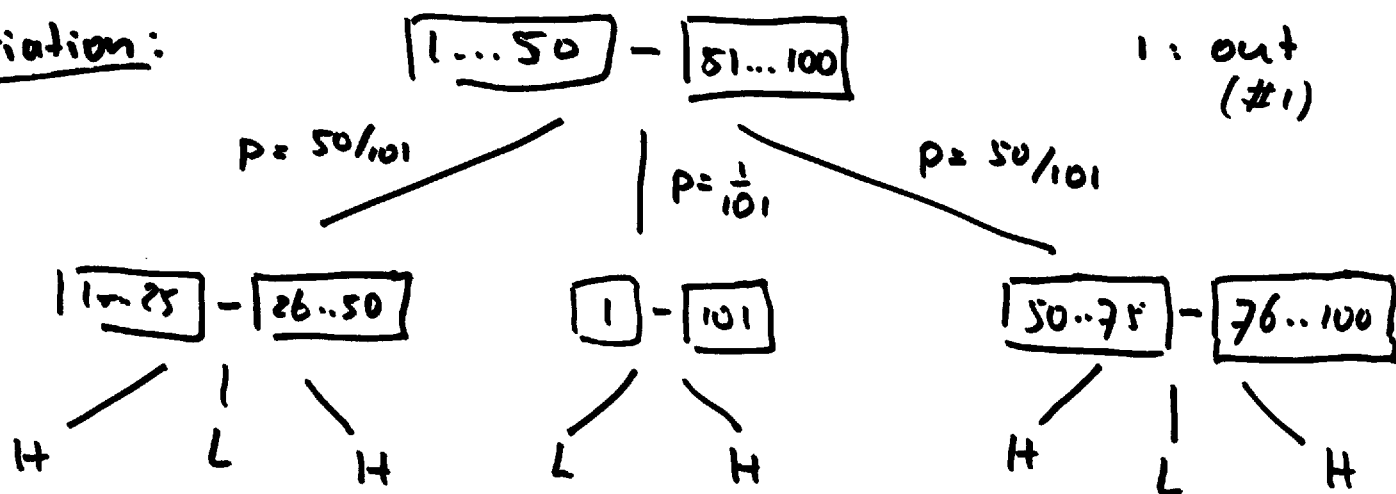
## Alternatives



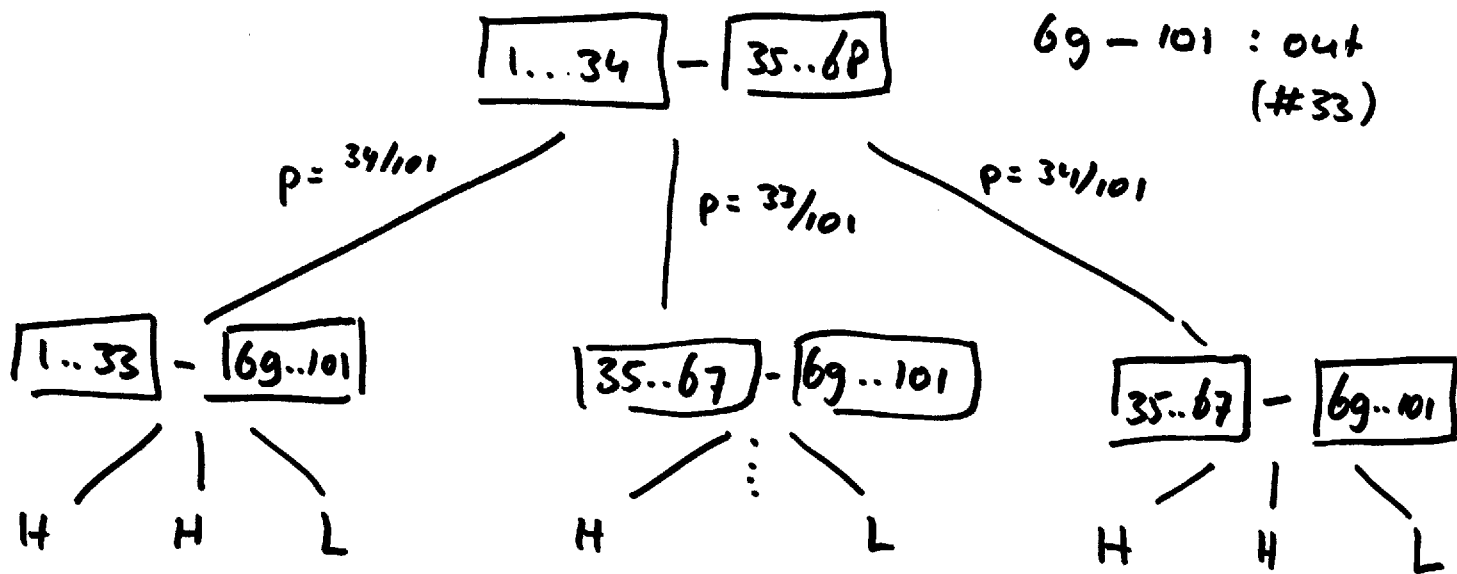
works. But note that entropy of first measurement is not optimized.

$H = -\sum p \log(p)$  is max when all  $p$ 's are more or less equal.

variation:



Same comment as above. Entropy of first measurement is even worse.



↑ does not work  
also goes here if one of 35..68 was lighter.