

Cosmology with Bayesian statistics and information theory

Lecture 3: Information theory

... a.k.a. *how much is there to be learned in my data anyway?*

Florent Leclercq

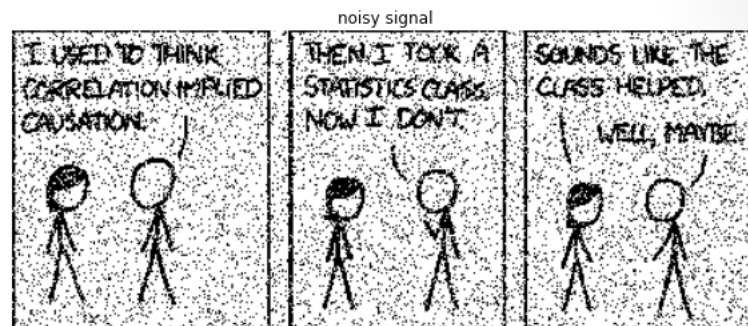
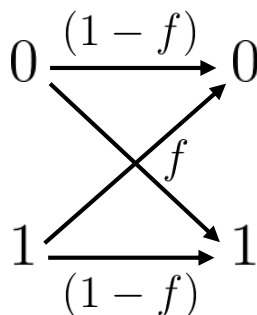
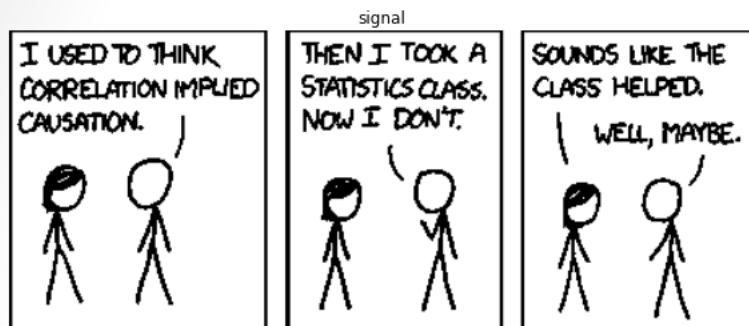
Institute of Cosmology and Gravitation, University of Portsmouth

<http://icg.port.ac.uk/~leclercq/>

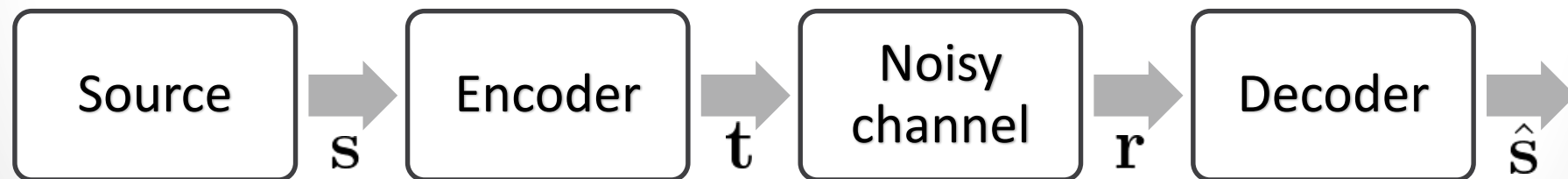


March 10th, 2017

The noisy binary symmetric channel



<https://xkcd.com/552/>



Rate of information transfer: $R = \frac{\#s}{\#t} = \frac{K}{N}$

The R3 code

s	0	0	1	0	1	1	0
t	$\overbrace{000}$	$\overbrace{000}$	$\overbrace{111}$	$\overbrace{000}$	$\overbrace{111}$	$\overbrace{111}$	$\overbrace{000}$
n	000	001	000	000	101	000	000
r	$\overbrace{000}$	$\overbrace{001}$	$\overbrace{111}$	$\overbrace{000}$	$\overbrace{010}$	$\overbrace{111}$	$\overbrace{000}$
$\hat{s}$	0	0	1	0	0	1	0

corrected errors

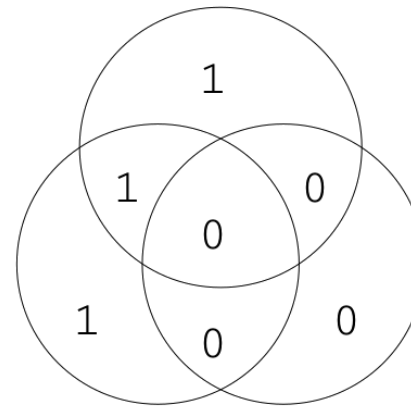
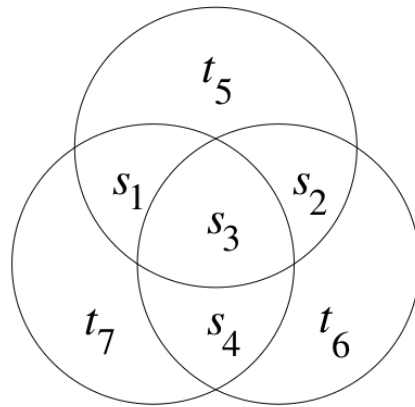
★

undetected errors

★

The (7,4) Hamming code: Encoder

- Introducing the concept of **parity-checks**:

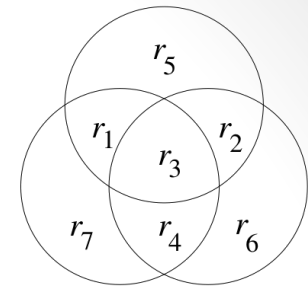


$$\mathbf{t} = \mathbf{G}\mathbf{s}$$

$$\mathbf{G} =$$

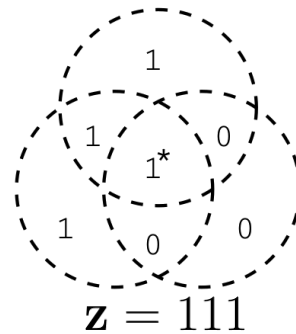
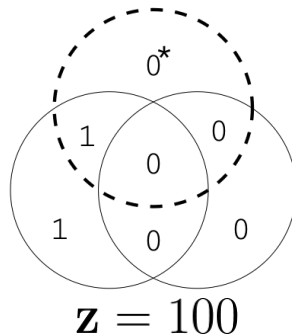
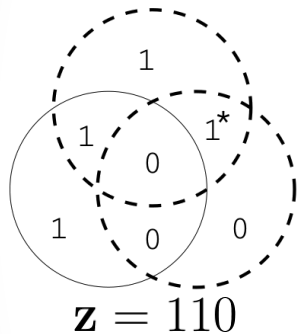
$$\begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \\ 1 & 1 & 1 & 0 \\ 0 & 1 & 1 & 1 \\ 1 & 0 & 1 & 1 \end{bmatrix}$$

The (7,4) Hamming code: Decoder



syndrome $\rightarrow$ $\mathbf{z} = \mathbf{H}\mathbf{r}$

- Introducing the concept of **syndrome**:
- Error correction:

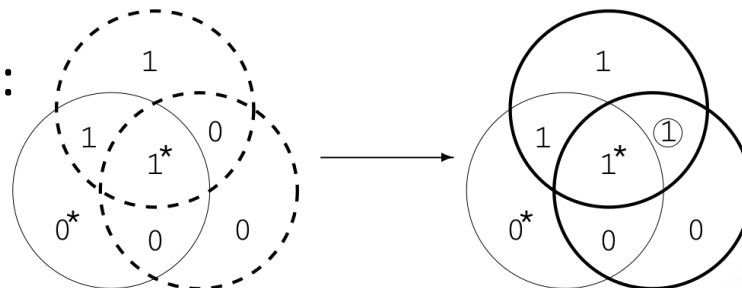


$$\mathbf{H} = \left[\begin{array}{c|c} \left\{ \begin{array}{cccc} 1 & 1 & 1 & 0 \\ 0 & 1 & 1 & 1 \\ 1 & 0 & 1 & 1 \end{array} \right\} & \left\{ \begin{array}{ccc} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{array} \right\} \end{array} \right]$$

parity-checks identity

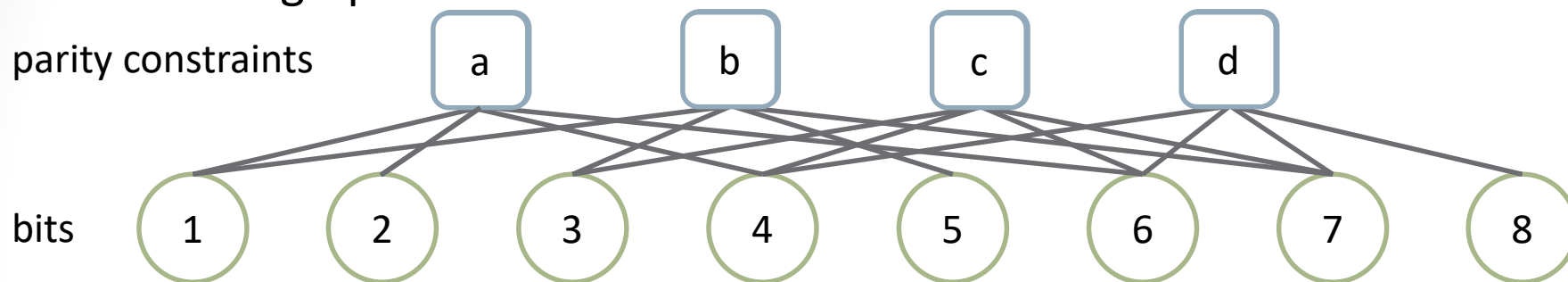
Syndrome $\mathbf{z}$	000	001	010	011	100	101	110	111
Unflip this bit	None	r_7	r_6	r_4	r_5	r_1	r_2	r_3

- Unsuccessful error correction:



Low-density parity check (LDPC) codes

- Tanner graph:



- N bits, $M=(1-R)N$ parity-check constraints
 - 2^{RN} possible “words”: the *dictionary*
 - A *sparse* parity-check matrix
- Decoding LDPC codes: the general theory borrows from *statistical physics*: Ising spins in interaction and the BP mean field approximation (Bethe-Peierls – Belief Propagation).

Measures of entropy and information

- **Information content:** $I[X] \equiv - \sum_{x \in \mathcal{X}} \log_2 p(x)$

- **Entropy:**

$$H[X] \equiv \langle I[X] \rangle_{p(X)} = - \sum_{x \in \mathcal{X}} p(x) \log_2 p(x)$$

- **Conditional entropy:**

$$H[X|Y] \equiv \langle H[X|Y = y] \rangle_{p(Y)} = \sum_{x \in \mathcal{X}, y \in \mathcal{Y}} p(x, y) \log_2 \left(\frac{p(y)}{p(x, y)} \right)$$

- **Properties:**

- chain rule: $H[X|Y] = H[X, Y] - H[Y]$
- “Bayes’s theorem”: $H[X|Y] + H[Y] = H[Y|X] + H[X]$

Measures of entropy and information

- **Mutual information:** $I[X:Y] \equiv \sum_{x \in \mathcal{X}, y \in \mathcal{Y}} p(x, y) \log_2 \left(\frac{p(x, y)}{p(x)p(y)} \right)$
- **Properties:**
 - $I[X:Y] \geq 0$
 - $I[X:X] = H[X]$
 - $I[X:Y] = H[X] - H[X|Y]$
 $= H[Y] - H[Y|X]$
 $= H[X] + H[Y] - H[X, Y]$
 $= H[X, Y] - H[X|Y] - H[Y|X]$
 - Consequence: $H[X|Y] \leq H[X]$
- **Relative entropy/Kullback-Leibler divergence/Information gain:** $D_{\text{KL}}[p||q] \equiv \sum_{x \in \mathcal{X}} p(x) \log_2 \left(\frac{p(x)}{q(x)} \right)$
- **Properties:**
 - Gibbs's inequality: $D_{\text{KL}}[p||q] \geq 0$
 - Relation to mutual information:

$$I[X:Y] = D_{\text{KL}}[p(x, y)||p(x)p(y)] = \langle D_{\text{KL}}[p(x|y)||p(x)] \rangle_{p(Y)}$$

Measures of entropy and information

- **Symmetrization procedures:**

- Jeffreys symmetrization: $C_s[A:B] = \frac{1}{2}C_a[A||B] + \frac{1}{2}C_a[B||A]$

- Jensen symmetrization:

$$C_s[A:B] = \frac{1}{2}C_a[A||M] + \frac{1}{2}C_a[B||M] \quad \text{with} \quad M = \frac{A+B}{2}$$

- **Jeffreys divergence:** $D_J[p:q] = \frac{1}{2}D_{\text{KL}}[p||q] + \frac{1}{2}D_{\text{KL}}[q||p]$

- **Jensen-Shannon divergence:**

$$D_{\text{JS}}[p:q] = \frac{1}{2}D_{\text{KL}}[p||r] + \frac{1}{2}D_{\text{KL}}[q||r] \quad \text{with} \quad r \equiv \frac{p+q}{2}$$

- **Properties:**

- $D_{\text{JS}}[p:q] = H[r] - \frac{1}{2}H[p] - \frac{1}{2}H[q]$

- $0 \leq D_{\text{JS}}[p:q] \leq 1$

- $D_{\text{JS}}[p:q] = I[m:z] \quad \text{with} \quad m \equiv zp + (1-z)q \quad z = \begin{cases} 0 & (p = 1/2) \\ 1 & (p = 1/2) \end{cases}$

Information-theoretic experimental design

(expected) data

experimental design

- **General problem:** maximize $U(\xi) = \langle U(d, \xi) \rangle_{p(d|\xi)} = \int p(d|\xi) U(d, \xi) dd$

- **1- Parameter inference utility functions:**

- Maximal information gain:

$$U(d, \xi) \equiv D_{\text{KL}}[p(\theta|d, \xi) || p(\theta|\xi)]$$

$$U(\xi) = I[\theta : d|\xi]$$

- A-optimality:

$$U_A(d, \xi) \equiv \frac{1}{\text{tr}(\text{cov}(\theta|d, \xi)^{-1})}$$

- D-optimality:

$$U_D(d, \xi) \equiv \det(\text{cov}(\theta|d, \xi))$$

Information-theoretic experimental design

- **2- Model selection utility functions:**

- Maximal Bayes factor:

$$U(\xi) \equiv \mathcal{B}_{12}(\xi) = \frac{p(d|\xi, \mathcal{M}_1)}{p(d|\xi, \mathcal{M}_2)} \quad \text{It is hard to predict Bayes factors...}$$

but see [Trotta 2007, arXiv:astro-ph/0703063](#)

- Maximal mutual information between the model indicator and the data:

$$U(\xi) \equiv I[\mathcal{M}:d|\xi] \quad \text{e.g. Cavagnaro *et al.* 2010}$$

- Maximal Jensen-Shannon divergence between posterior predictive distributions:

$$U(\xi) \equiv D_{\text{JS}}[p_1:p_2] = I[\mathcal{M}:r|\xi] \quad \text{with} \quad r \equiv \frac{p_1 + p_2}{2}$$

[Vanlier *et al.* 2014, FL, Lavaux, Jasche & Wandelt 2016, arXiv:1606.06758](#)

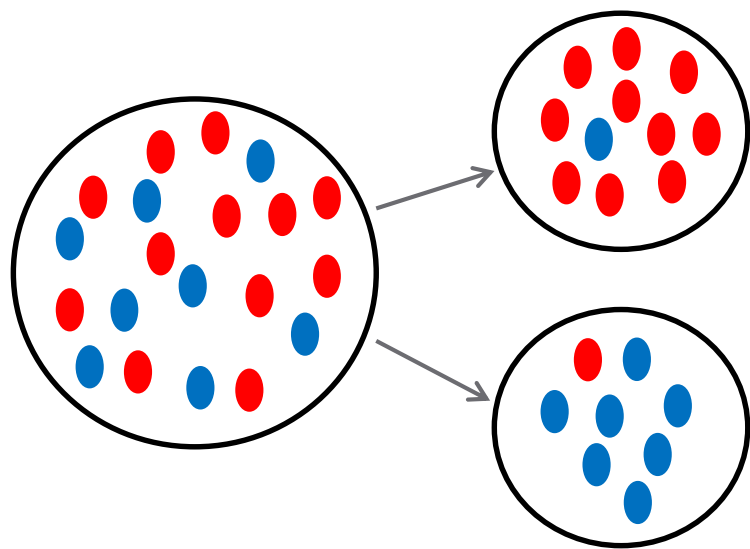
- **3- Utilities for prediction of future observations:**

$$U(d, \xi) \equiv D_{\text{KL}}[p(t|d, \xi) || p(t|\xi)] \quad U(\xi) = I[t:d|\xi] = H[p(t|\xi)] - H[p(t|d, \xi)]$$

i.e. the “supervised machine learning” utility: $U(a) = H[T] - H[T|a]$

Supervised machine learning basics

- How to compute the information gain?



child1 entropy:

$$H = -\frac{10}{11} \log_2 \left(\frac{10}{11} \right) - \frac{1}{11} \log_2 \left(\frac{1}{11} \right) = 0.4395$$

child2 entropy:

$$H = -\frac{8}{9} \log_2 \left(\frac{8}{9} \right) - \frac{1}{9} \log_2 \left(\frac{1}{9} \right) = 0.5033$$

parent entropy:

$$H = -\frac{8}{20} \log_2 \left(\frac{8}{20} \right) - \frac{12}{20} \log_2 \left(\frac{12}{20} \right) = 0.9709$$

weighted average entropy of children:

$$\frac{11}{20} \times 0.4395 + \frac{9}{20} \times 0.5033 = 0.4682$$

$$\text{information gain for this split: } 0.9709 - 0.4682 = 0.5027 \text{ Sh}$$