

3. Grundlagen der Quanteninformationsübertragung

3.1. Das klass. Maß der Information, klass. Datenkompression

Zunächst müssen wir uns verständigen, was wir unter Information verstehen wollen.

- Beispiel: Würfel • Wenn ein Würfelspieler Ihnen sagt, er habe eine "5" gewürfelt, so hat diese Aussage für Sie nur Information, wenn Sie nicht beim Würfeln zugehen haben
- Wenn der Würfelspieler sagt er habe eine Zahl zwischen 1 und 6 gewürfelt hat das für Sie überhaupt keine Information

Informationsgehalt einer Aussage
= Maß der a priori Ignoranz

Sei A ein zufälliges Ereignis mit Wahrscheinlichkeit $p(A)$. $\bar{A}$ heißt das Gegenteil von A mit $p(\bar{A}) = 1 - p(A)$.

Eine Menge $\{A_1, \dots, A_n\}$ von Ereignissen heißt vollständig falls $\sum_{i=1}^n p(A_i) = 1$.

(1) Def: Sei $\{A_1, \dots, A_n\}$ eine vollständige Menge von Ereignissen. Die A_i heißen dann Elementarereignisse falls

$$(1) \quad p(A_i) > 0 \quad \text{und} \quad A_i \cap A_j = \emptyset \quad \forall j \neq i.$$

Die Menge aller Elementarereignisse $\Omega = \{A_1, \dots, A_n\}$ heißt Ereignisraum. Jedes Ereignis A ist Teilmenge von Ω .

Prinzip der maximalen Unbestimmtheit

Wenn über ein Ereignis A nichts bekannt ist wird jedem Elementarereignis die a priori Wahrscheinlichkeit

$$(2) \quad p_i = \frac{1}{n}$$

zugeordnet

Nun wollen wir ein quantitatives Maß S der Unbestimmtheit konstruieren. Es soll folgenden Kriterien genügen

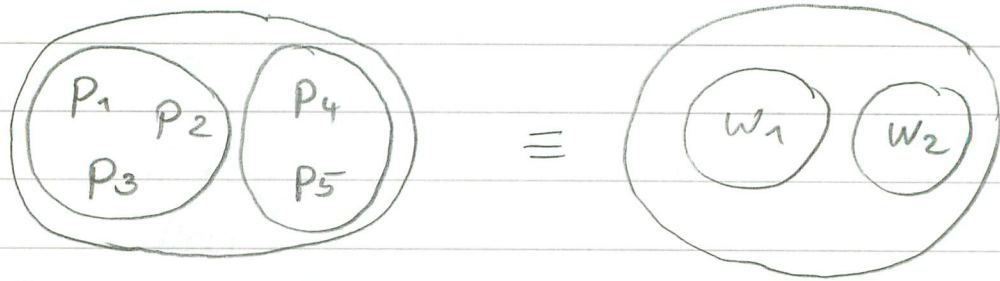
Annahme: n diskrete Ereignisse

(i) S ist Funktion der Wahrscheinlichkeiten p_i
 $S = S(p_1, \dots, p_n)$ und stetig in p_i

(ii) $S = 0$ für sicheres Ereignis (ein $p_i = 1$)

(iii) $S = \max$ für maximale Unbestimmtheit ($p_i = \frac{1}{n}$) und monoton wachsend in n
(Je größer der Ereignisraum über den ich nichts weiß umso mehr Information steckt in der Zufallsvariablen)

(iv) H soll folgende Kompositionsregel erfüllen



Gruppe aus
3 Ereignissen

Gruppe aus
2 Ereignissen

$$w_1 = p_1 + p_2 + p_3$$

$$w_2 = p_4 + p_5$$

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

Gesamtinformation = Information "w₁ oder w₂"

+ gewichtete Information

"falls w₁: p₁ oder p₂ oder p₃"

"falls w₂: p₄ oder p₅"

Die Forderungen (i) – (iv) legen H eindeutig fest.

Man findet die sog. Shannon Entropie

$$(4) \quad S(p_1, \dots, p_n) = -\alpha \sum_{i=1}^n p_i \ln p_i \quad \alpha > 0$$

Beweis:

(i) ✓

(ii) ✓

(iii) ✓

da aus $p_i = 1 \quad p_j = 0 \quad \forall j \neq i$
 $\Rightarrow S = 0$

(iii) Um (iii) zu zeigen, beweisen wir zunächst folgende Aussage: Seien $(p_1, \dots, p_n)$ und $(\bar{p}_1, \dots, \bar{p}_n)$ zwei Wahrscheinlichkeitsverteilungen für n Elementarereignisse. Dann gilt

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

Gleichheit gilt für $\bar{p}_i \equiv p_i$

Bew: $\ln x \leq x - 1 \quad \forall x \geq 0$

$$\Rightarrow \ln \frac{1}{x} = -\ln x \geq 1 - x$$

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

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

Aus (5) folgt mit $p_i \equiv \frac{1}{n}$

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

$$\begin{aligned} S(\bar{p}_1, \bar{p}_2, \dots, \bar{p}_n) &= -a \sum_{i=1}^n \bar{p}_i \ln \bar{p}_i \leq -a \sum_{i=1}^n \bar{p}_i \ln \frac{1}{n} \\ &= a \ln n \quad \square \end{aligned}$$

(iv) Es sei $W_1 = p_1 + \dots + p_m$
 $W_2 = p_{m+1} + \dots + p_n$

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

$$w_1 S\left(\frac{p_1}{w_1}, \dots, \frac{p_m}{w_m}\right) = -a \cancel{w_1} \sum_{i=1}^m \frac{p_i}{\cancel{w_1}} \ln\left(\frac{p_i}{\cancel{w_1}}\right)$$

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

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

$$= -a \cancel{w_1} \ln w_1 - a \cancel{w_2} \ln w_2$$

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

$$-a \sum_{i=m+1}^n p_i \ln p_i + a \cancel{w_2} \ln w_2$$

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

Beweis Ende

Bem: Der Vorfaktor a kann nach beliebig gewählt werden

Der kleinste nichttriviale Ereignisraum hat zwei Alternativen

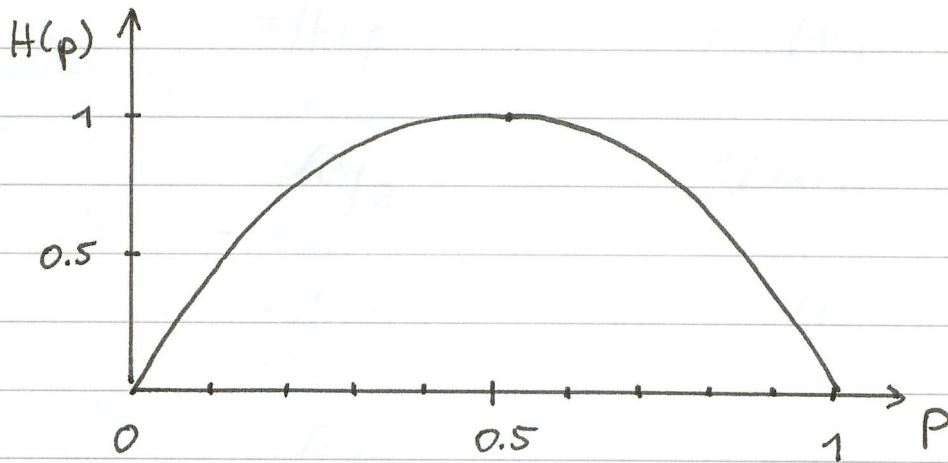
$$\{A, \bar{A}\} \hat{=} \{0, 1\}$$

Diesen Ereignisraum wollen klassisches bit oder cbit nennen.

Für ein cbit ist die Shannon Entropie gegeben durch

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

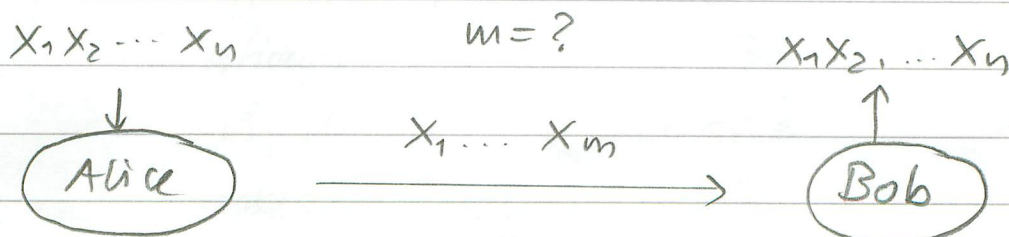
wobei $\log = \log_2$ und $a = 1/\ln 2$ gewählt wurde



Klassische Information kann durch Reihen von cbits, sog. bit-Wörter übertragen werden

$$(7) \quad x_1 x_2 \dots x_n \quad x_i \in \{0,1\}$$

Die Bedeutung der Shannon Entropie als Informationsmaß erläutern wir nun an einem zentralen Problem der Informationsübertragung: Gegeben ein bitwort der Länge n mit $n \rightarrow \infty$. Wieviel cbit werden benötigt um das bitwort von A nach B zu übertragen



Die Antwort wird durch

Shannon

Shannons noiseless coding Theorem

Zur Übertragung eines n bit Wortes genügen für $n \rightarrow \infty$ $n \cdot H(p)$ bits, wobei p die Wahrscheinlichkeit des Auftretens von "0" oder "1" in dem n bit Wort ist

Darüberhinaus gilt: Es gibt immer einen eindeutig decodierbaren Code aus $1 + n \cdot H(p)$ bits. Shannons Theorem gibt aber keine Konstruktionsvorschrift.

Falls "0" und "1" gleichwahrscheinlich $H(p) = \log 2$.

Bspl:

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

$$p(0) = \frac{225}{320} = 0,703 > p(1) = \frac{95}{320} = 0,297$$

Original

Code

0000

10

0001

000

0010

001

0100

010

1000

011

0011

1000

0101

1001

1001

1011

0110

1010

1010

1100

1100

1101

0111

11000

1011

11111

1101

11110

1110

11101

1111

111001

häufige Wörter
erhalten kurzen
code

codierter
Text:

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 11110 010 001 10 1001 000 010

10 1001 000 010 011 1001 1101 000 1101 011

Die Verkürzung ist deutlich zu erkennen.

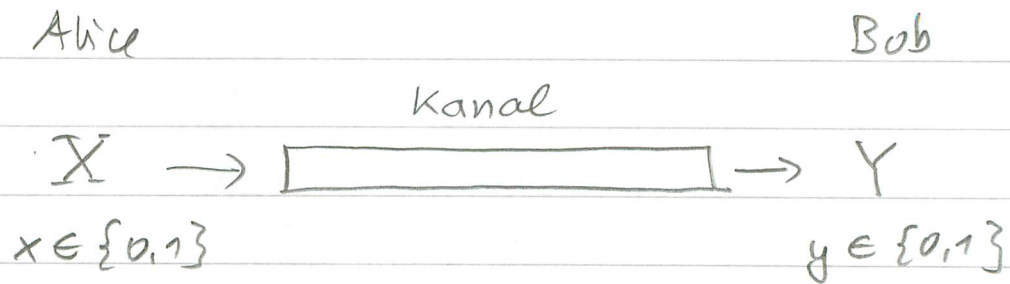
Der in diesem Beispiel verwendete code ist illustrativ aber ein bißchen gemogelt da Leerzeichen im Text beibehalten wurden. Ein code der für Wörter ohne Leerzeichen anwendbar ist, ist der sog. Huffman code

0000	10	für $p(0) = 3/4$ $p(1) = 1/4$
0001	000	
0010	001	
0100	010	
1000	011	Bem.: Huffman code praktisch nicht verwendet.
0011	11000	
0101	11001	
0110	11010	
1001	11011	
1010	11100	
1100	11101	
0111	1111000	
1011	111111	
1101	111110	
1110	111101	
1111	1111001	

3.2. Kommunikationskanäle

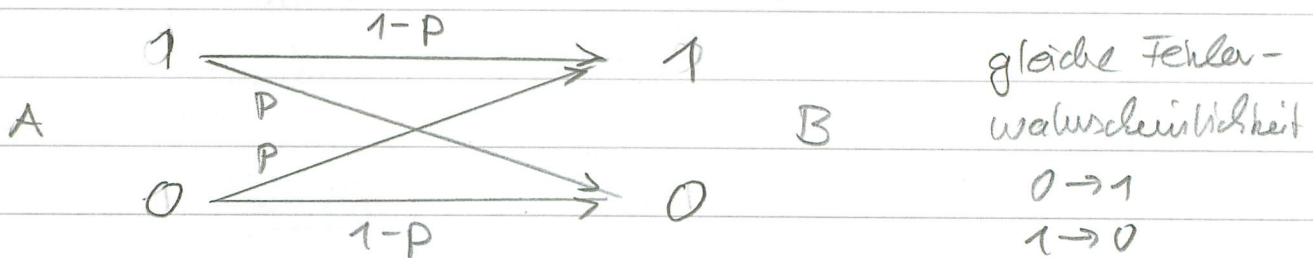
Shannons noiseless coding Theorem macht eine Aussage über die Informationsübertragung durch ideale d.h. verlustfreie Kanäle. Kommunikationskanäle sind aber i.A. verlustbehaftet und führen zu Rauschen.

Betrachten binären Übertragungskanal (Alice kann an Bob "0" oder "1" senden)

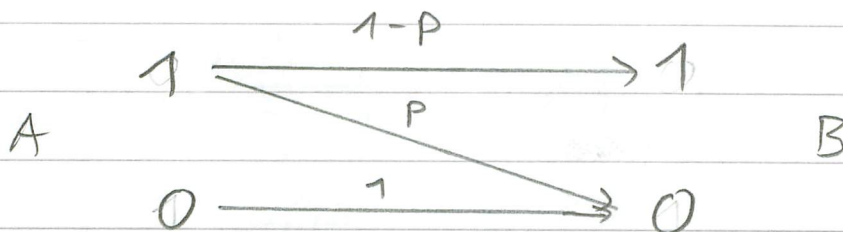


Eine wichtige Klasse von Kanälen sind gedächtnisfreie Kanäle. Bei ihnen hängt die Wahrscheinlichkeit p eines Übertragungsfehlers nicht von der Vorgeschichte ab.

Symmetrischer binärer Kanal



binärer Zufallskanal



Den Kanälen kann eine Übertragungsmatrix P_{ij} zugeordnet werden.

symmetrisch binärer Kanal

Zerfalls kanal

$$(8) \quad \begin{pmatrix} 1-p & p \\ p & 1-p \end{pmatrix} = \begin{pmatrix} 1 & p \\ 0 & 1-p \end{pmatrix}$$

Eine wichtige Größe zur Charakterisierung eines Kanals ist die gegenseitige Information

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

wobei

$$(10) \quad S(Y|X) = - \sum_i p(x_i) \sum_j p(y_j|x_i) \log p(y_j|x_i) \\ = \sum_{ij} p(x_i, y_j) \log p(y_j|x_i)$$

die bedingte Entropie ist. Hierbei sind

$p(x_i, y_j)$ die Verbundwahrscheinlichkeit für das Auftreten von x_i und y_j

$p(y_j|x_i)$ die konditionelle Wahrscheinlichkeit des Auftretens von y_j falls x_i eingehtreten ist

Es gilt

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

$$S(Y|X) = ? = - \sum_x p(x) \sum_y p(y|x) \log p(y|x)$$

symmetrischer binärer Kanal

$$p(y=0|x=0) = p(y=1|x=1) = 1-p_e$$

$$p(y=0|x=1) = p(y=1|x=0) = p_e$$

$$\begin{aligned} S(Y|X) &= - p(x=0) \left[p(y=0|x=0) \log p(y=0|x=0) \right. \\ &\quad \left. + p(y=1|x=0) \log p(y=1|x=0) \right] \\ &\quad - p(x=1) \left[p(y=0|x=1) \log p(y=0|x=1) \right. \\ &\quad \left. + p(y=1|x=1) \log p(y=1|x=1) \right] \end{aligned}$$

$$= - p(x=0) \left[(1-p_e) \log (1-p_e) + p_e \log p_e \right]$$

$$- p(x=1) \left[p_e \log p_e + (1-p_e) \log (1-p_e) \right]$$

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

Für unabhängige Ereignisse ist

$$(12) \quad P(Y|X) = P(Y) \Rightarrow P(X, Y) = P(X)P(Y)$$

Die gegenseitige Information kann geschrieben werden als

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

Falls X und Y unabhängig folgt aus (10)

$$(14) \quad S(X|Y) = S(X) \Rightarrow I(X:Y) = 0$$

und somit

Für einen binären symmetrischen Kanal gilt

$$(15) \quad I(X:Y) = H(p(y)) - H(p_{\text{error}})$$

Wir definieren jetzt die Kanalkapazität

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

Für den binären symmetrischen Kanal erhalten wir somit

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

Nun wollen wir uns die Frage stellen wieviel cbits benötigt werden um ein n bit Wort über einen verrauschten (fehlerhaften) Kanal zu übertragen. Dazu führen wir zunächst die Übertragungsrate R ein:

$$(18) \quad R = \frac{\# \text{ Nachrichtensbits}}{\# \text{ Benutzung des Kanals}} \quad 0 \leq R \leq 1$$

Nehmen wir an der Kanal hat eine Fehlerwahrscheinlichkeit $p_b < 1$. Senden wir die Nachricht mit $R=1$, ist die Fehlerwahrscheinlichkeit im erhaltenen Signal $e = p_b$. Um die Fehlerwahrscheinlichkeit zu verringern benutzen wir die Kodierung

$$(19) \quad 0 \rightarrow 000 \quad 1 \rightarrow 111 \quad (\text{Repetitioncode})$$

D.h. die Übertragungsrate ist $R = \frac{1}{3}$. Wenn nur eines der 3 cbits falsch übertragen wird, kann die Nachricht immer noch exakt rekonstruiert werden. D.h. nur wenn mindestens 2 cbits falsch sind tritt ein Übertragungsfehler auf. Die Wahrscheinlichkeit dafür ist

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

m -facher Repetitioncode

- $e \xrightarrow{m \rightarrow \infty} 0$ 
- $R \xrightarrow{m \rightarrow \infty} 0$ 

Läßt sich Information über einen verrauschten Kanal mit beliebig kleiner Fehlerwahrscheinlichkeit nur mit $R \rightarrow 0$ realisieren? Nein!

Shannons noisy coding theorem (1948)

Für einen Kanal der Kapazität C existiert ein Code der eine Datenübertragung mit beliebig kleiner Fehlerwahrscheinlichkeit und $0 < R < C$ erlaubt.

Fehlerfreie Datenübertragung über verrauschte Kanäle mit nichtverschwindender Übertragungsrate ist möglich!

3.3. Klassische Fehlerkorrektur

In den letzten 50 Jahren ist ein ganzes Theoriegebäude fehlerkorrigierender (klassischer) Codes entwickelt worden. Ein wichtiger Vertreter sind binäre lineare Codes, die jetzt kurz diskutiert werden sollen.

binärer Code: k cbits sind in $n > k$ cbits kodiert

Anzahl der Codewörter 2^k = Menge der
Anzahl der Wörter insg. 2^n

lineare binärer Code:

Menge G von Codewörtern die bezüglich bitweiser Addition modulo 2 ein Galoisfeld (binärer lin. Raum) (Galois Feld)

Bspl: $\{(0,1,1), (1,0,1), (1,1,0), (0,0,0)\}$

$$(0,1,1) \oplus (1,0,1) = (1,1,0) \quad n=3$$

$$(0,1,1) \oplus (1,1,0) = (1,0,1) \quad k=2$$

$$(1,0,1) \oplus (1,1,0) = (0,1,1) \quad \text{Basis z.B. } \begin{pmatrix} 0,1,1 \\ 1,0,1 \end{pmatrix}$$

Basis von G : $\{v_1, v_2, \dots, v_k\}$ v_i codewörter aus $n \geq k$ bits

Jedes codewort läßt sich dann als Linearkombination der v_i schreiben

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

und $\oplus$ bedeutet bitweise Addition modulo 2.

$\underline{v}$ kodiert die k -bit Nachricht $(\alpha_1, \alpha_2, \dots, \alpha_k)$.

Die k Basisvektoren bilden die $k \times n$ Erzeugende

$$(22) \quad \underline{\underline{G}} = \begin{bmatrix} v_1 \\ \vdots \\ v_k \end{bmatrix}$$

Im obigen Beispiel gibt es zwei Basisvektoren z.B. $(0,1,1)$ und $(1,0,1)$ und somit wäre

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

Dann kann Gl. (21) in der Form

$$(24) \quad \underline{v} = \underline{\underline{x}} \underline{\underline{G}} \quad \underline{x} = (\alpha_1, \alpha_2, \dots, \alpha_k)$$

geschrieben werden.

Alternativ zur Charakterisierung des Codes mittels der Matrix G kann man auch die $n-k$ linearen Bedingungen

$$(25) \quad \underline{\underline{H}} \underline{v}^T = 0$$

benutzen wobei H eine $(n-k) \times n$ Matrix ist.

H heißt parity-check Matrix. Die Zeilen von H sind $(n-k)$ linear unabhängige Vektoren die auf allen Codewörtern senkrecht stehen. Im obigen Beispiel:

Beispiel

$$(26) \quad \underline{\underline{H}} = \begin{pmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \end{pmatrix}$$

Basis: $\{(0,1,1), (1,0,1)\}$

Codewörter $\{(0,0,0), (0,1,1), (1,0,1), (1,1,0)\}$

Range von $\underline{x} = (0,0) : (0,0,0)$

$\underline{x} = (1,0) : (0,1,1)$

$\underline{x} = (0,1) : (1,0,1)$

$\underline{x} = (1,1) : (1,1,0)$

Im Falle klassischer Kanäle ist der einzig mögliche Fehler ein bit-flip. Ein bit-flip Fehler entspricht der Addition eines Fehlervektors

$$(27) \quad \underline{v} \longrightarrow \underline{v} \oplus \underline{e}_i$$

wobei $\underline{e}$ Einträge "1" an den Stellen hat an denen (ein bit flip stattgefunden hat, z.B.)

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

den bit-flip des 1ten bits bedeutet. Fehler $\underline{e}$ entspricht einem flip des ersten und dritten bits.

In einem linearen binären Code können Fehler einfach durch Anwenden der parity check Matrix $\underline{H}$ detektiert werden:

$$(29) \quad \underline{H} (\underline{v} \oplus \underline{e}_i) = \underline{H} \underline{v} \oplus \underline{H} \underline{e}_i = \underline{H} \underline{e}_i$$

$\underline{H} \underline{e}_i$ ist das Fehlersyndrom.

Die Menge aller Fehler $\{\underline{e}_i\}$ die unterschiedliche Fehlersynonyme haben sind mittels des verwendeten codes korrigierbar. Nach eindeutiger Detektion des Fehlers wird dieser dazu einfach noch einmal addiert $\underline{v} + \underline{e}_i \longrightarrow \underline{v} + \underline{e}_i + \underline{e}_i$

$$(30) \quad (\underline{v} \oplus \underline{e}_i) \oplus \underline{e}_i = \underline{v}$$

Falls aber $\underline{H} \underline{e}_1 = \underline{H} \underline{e}_2$ könnten wir den

misinterpretieren und

$$(31) \quad \underline{v} \oplus \underline{c}_1 \rightarrow (\underline{v} \oplus \underline{c}_1) \oplus \underline{c}_2 \neq \underline{v}$$

Der Hamming Abstand eines Codes ist die minimalanzahl von bits die geändert werden müssen um ein Element des Codes in ein anderes zu überführen.

Ein linearer Code mit Abstand $d = 2t + 1$ kann t bit-flip Fehler korrigieren

Der oben eingeführte Code hat $d=2$ und kann daher keinen Fehler korrigieren, denn

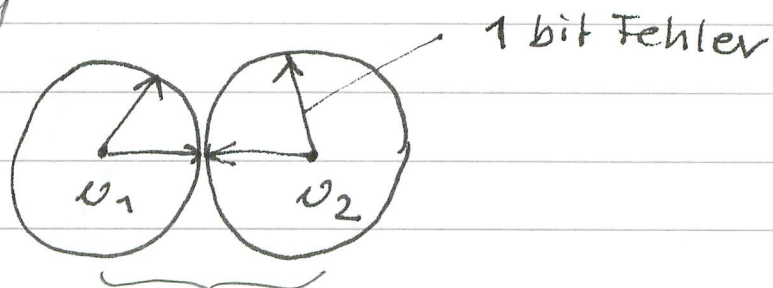
$$v_1 = (0, 1, 1) \oplus c_1 = (1, 0, 0)$$

||

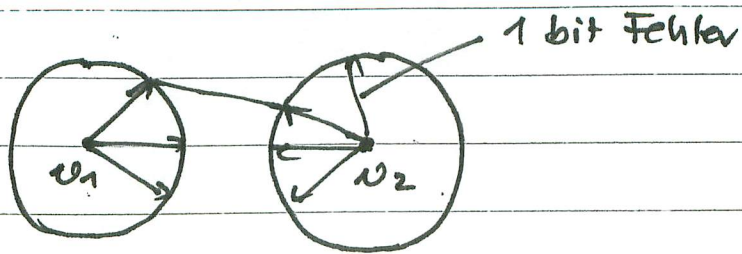
(32)

$$v_2 = (1, 0, 1) \oplus c_2 = (0, 1, 0)$$

Hier hat man



Für $d=3$ hat man



$$d = 3$$

Beispiel für $[n=7, k=4, d=3]$ Code

<u>Message</u>	<u>Huffman</u>	<u>Hamming</u>
0000	10	0000000
0001	000	1010101
0010	001	0110011
0011	11000	1100110
0100	010	0001111
0101	11001	1011010
0110	11010	0111100
0111	1111000	1101001
1000	011	1111111
1001	11011	0101010
1010	11100	1001100
1011	111111	0011001
1100	11101	1110000
1101	111110	0100101
1110	111101	1000011
1111	1111001	0010110

Der Hamming $[7,4,3]$ Code erlaubt Korrektur aller 1-bit Fehler.

3.4. Quanteninformation

In Abschnitt 3.1. hatten wir klassische Zufallsgrößen X betrachtet, die mit Wahrscheinlichkeit p_i die Werte (oder klassischen Zustände) x_i annehmen können. Wir hatten dann die von Shannon Entropie

$$(33) \quad S_{sh} = - \sum_i p_i \ln p_i$$

als Maß für unser Unwissen über den tatsächlichen Zustand x_i eingeführt.

Nun wollen wir die x_i durch quantenmechanische Zustände ersetzen. Zunächst wollen wir verschiedenen Werten x_i orthogonale quantenmechanische Zustände zuordnen

$$(34) \quad x_i \longleftrightarrow |x_i\rangle \langle x_i| \text{ mit } \langle x_i | x_j \rangle = \delta_{ij}$$

die mit Wahrscheinlichkeit p_i auftreten sollen.

Quantenmechanisch entspricht dies einem Dichtoperator

$$(35) \quad \rho = \sum_i p_i |x_i\rangle \langle x_i|$$

und (33) kann geschrieben werden als

$$(36) \quad S_{sh} = - \text{Tr} \{ \rho \ln \rho \} = S_N$$

was gerade der von-Neumann Entropie entspricht.

Als quantitatives Maß für die Unbestimmtheit eines Quantenzustandes $\hat{\rho}$ wollen wir daher die von-Neumann Entropie

$$(37) \quad S_N = - \text{Tr} \{ \hat{\rho} \log \hat{\rho} \}$$

verwenden.

Eigenschaften von S_N

$$(i) \quad \text{reiner Zustand } \hat{\rho} = |\varphi\rangle\langle\varphi| \Rightarrow S_N = 0 \quad (38)$$

$$(ii) \quad \begin{array}{l} \text{maximaler gemischter Zustand } \hat{\rho} = \frac{1}{N} \mathbb{I}_N \\ N = \dim(\mathcal{H}) \end{array} \Rightarrow S_N = \log N \quad (39)$$

$$(iv) \quad \text{allg.} \quad 0 \leq S_N \leq \log(\dim(\mathcal{H})) \quad (40)$$

$$(v) \quad \begin{array}{l} S_N \text{ ist konvex, d.h. } \lambda_1, \lambda_2, \dots, \lambda_N \geq 0 \\ \text{und } \lambda_1 + \lambda_2 + \dots + \lambda_N = 1 \end{array}$$

$$(41) \quad S_N(\lambda_1 \hat{\rho}_1 + \dots + \lambda_n \hat{\rho}_n) \geq \lambda_1 S_N(\hat{\rho}_1) + \dots + \lambda_n S_N(\hat{\rho}_n)$$

Die Eigenschaft (v) bringt folgenden Sachverhalt zum Ausdruck: Wenn wir wissen, daß ein Quantenzustand aus n anderen Zuständen $\hat{\rho}_i$ gemischt ist

zustand $\rho = \lambda_1 \rho_1 + \dots + \lambda_N \rho_N$ aus einem Gemisch der Zustände $\rho_1, \dots, \rho_N$ mit Gewichten (Wahrscheinlichkeiten) λ_i konstruiert wurde besitzen wir ^{1.A.} inher apriori Information über das System als wenn uns nur ρ bekannt ist. Das liegt daran, daß dasselbe ρ auf verschiedene Arten zusammengesetzt werden kann.

(vi) unitäre Evolutionen lassen S_N unverändert

(42)

$$S_N(U \rho U^{-1}) = S_N(\rho)$$

(vii) es gilt die Subadditivität

(43)

$$S_N(\rho_{AB}) \leq S_N(\rho_A) + S_N(\rho_B)$$

wobei $\rho_A = \text{Tr}_B \{ \rho_{AB} \}$ und $\rho_B = \text{Tr}_A \{ \rho_{AB} \}$

(44) ~~Fürdies~~ $S_N(\rho_{AB}) = S_N(\rho_A) + S_N(\rho_B)$

falls $\rho_{AB} = \rho_A \otimes \rho_B$

Diese Eigenschaft ist analog zur Shannon Entropie

$$(45) \quad S_{Sh}(X, Y) = S_{Sh}(X) + S_{Sh}(Y) - S_{Sh}(X|Y)$$

$$\text{d.h.} \quad S_{Sh}(X, Y) \leq S_{Sh}(X) + S_{Sh}(Y)$$

und $S_{Sh}(X|Y) = 0$ für unkorrelierte Variable X, Y

(viii) es gilt weiterhin die starke Subadditivität

$$(46) \quad S_N(S_{ABC}) + S_N(S_B) \leq S_N(S_{AB}) + S_N(S_{AC})$$

Diese für die von Neumann Entropie nichttriviale Bzghg. ist die Grundlage vieler formaler Beweise in der Q.-Inf., (die wir aber in dieser Vorlesung i.A. nicht behandeln werden).

(ix) Araki-Lieb Ungleichung

$$(47) \quad S_N(S_{AB}) \geq |S(S_A) - S(S_B)|$$

Diese Ungleichung ist im scharfen Kontrast zur klassischen Ungleichung!

$$S_{Sh}(X, Y) \geq S_{Sh}(X), S_{Sh}(Y)$$

D.h. insbesondere kann $S_N(S_{AB}) < S(S_A), S(S_B)$.

Bsp1: reiner Zweiparteienzustand

$$|\phi_+\rangle = \frac{1}{\sqrt{2}} (|00\rangle + |11\rangle) \quad S_A = \frac{1}{2} \mathbb{1} = S_B$$

$$S(S_A) = S(S_B) = 1 \quad \text{aber} \quad S(S_{AB}) = 0$$

In einem verschränkten reinen Zstd. steckt alle Infor-

mation in nichtlokalen Korrelationen.

Aus Gleichung (40) sehen wir, daß der kleinste Hilbertraum für den $S_N \neq 0$ möglich ist, der eines Spin $\frac{1}{2}$ Teilchens ist. In Analogie zum klassischen cbit als Einheit der Information wollen wir

$$(48) \quad \mathcal{H} = (|0\rangle, |1\rangle) \quad \text{qubit}$$

heissen.

Der Quantenzustand eines qubits kann stets in der Form

$$(49) \quad \rho = \frac{1}{2} (1 + \vec{S} \cdot \vec{\sigma})$$

wobei $\vec{S} \in \mathbb{R}^3$, $|\vec{S}| \leq 1$ der Polarisationsvektor des qubits bedeutet und $\vec{\sigma} = \sigma_x \vec{e}_x + \sigma_y \vec{e}_y + \sigma_z \vec{e}_z$ der Pauli vektoroperator ist,

Es gilt

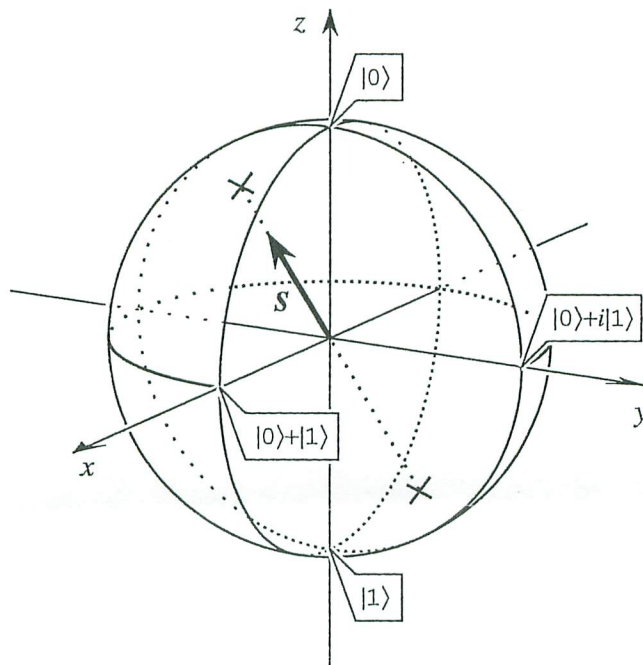
$$(50) \quad S_i = \text{Tr} \{ \vec{\sigma}_i \rho \}$$

(49) kann auf Diagonalform gebracht werden

$$(51) \quad \rho = \rho_+ |\uparrow_s\rangle \langle \uparrow_s| + \rho_- |\downarrow_s\rangle \langle \downarrow_s|$$
$$\rho_{\pm} = \frac{1}{2} (1 \pm |\vec{S}|)$$

wobei $|\uparrow_s\rangle, |\downarrow_s\rangle$ Eigenzustände von $\vec{S} \cdot \vec{e} = \hat{S}_s$ sind. Die Darstellung ist eindeutig falls $|\vec{S}| \neq 0$.

Der Quantenzustand eines qubits kann mittels der Bloch Kugel dargestellt werden:



$|\vec{S}|=1$ entspricht den reinen Zuständen. $|\vec{S}|=0$ ist der maximal gemischte Zustand.

3.5. Quanteninformation vs. Information aus der Erzeugung bzw. der Messung von Quantenzuständen

In der klassischen Inf.-theorie ist die Information die in einer Zufallsvariable X steckt identisch zu der die wir durch eine Messung von X erhalten können bzw. die wir durch Kenntnis der Erzeugungsprozedur von X erlangen können. Quantenmechanisch ist dies

nicht mehr der Fall:

Gegeben ein Quantenzustand g . Nun messen wir die Observable $\hat{A}$

$$(52) \quad \hat{A} = \sum_y |a_y\rangle a_y \langle a_y|$$

was uns die Ergebnisse a_y mit Wahrscheinlichkeit $p_y = \langle a_y | g \rangle \langle a_y |$ liefert. Die Meßergebnisse $\{a_y, p_y\}$ sind jetzt eine klassische Zufallsvariable der wir eine Shannon Entropie zuordnen können. Diese wird entropy of measurement genannt. Es gilt:

$$(53) \quad S_{sh}(a_y) \geq S_N(g)$$

Gleichheit gilt falls $[\hat{A}, g] = 0$, d.h. wenn die $|a_y\rangle$ Eigenzustände von g sind.

Umgekehrt kann man die Erzeugung eines Quantenzustandes g aus Mischung reiner Zustände $|\varphi_x\rangle$ mit Wahrscheinlichkeiten p_x betrachten

$$(54) \quad g = \sum_x p_x |\varphi_x\rangle \langle \varphi_x|$$

Wenn die $|\varphi_x\rangle$ jetzt nicht notwendig mehr orthogonal sind gilt für die Shannon Entropie der Zufallsvariable $\{x, p_x\}$:

(55)

$$S_{sh}(\varphi_x) \geq S_N(S)$$

Gleichheit gilt falls alle $|\varphi_x\rangle$ orthogonal (so hatten wir ja in Abschnitt 3.4. die von Neumann Entropie eingeführt). $S_{sh}(\varphi_x)$ wird entropy of preparation genannt.

Bspl:

Quelle die qubits in $|\uparrow_{\vec{a}}\rangle$ mit Wahrscheinlichkeit p_a und in $|\uparrow_{\vec{b}}\rangle$ mit p_b erzeugt. Sei $p_a = p_b = 0.5$. Falls $\vec{a}$ und $\vec{b}$ verschieden sind, sind sie klassisch unterscheidbar!!!

• Shannon entropy of preparation

(56)

$$H_{\text{prep}} = -p_a \log p_a - p_b \log p_b = 1$$

• Q.-Zustand

(57)

$$\rho = p_a |\uparrow_{\vec{a}}\rangle \langle \uparrow_{\vec{a}}| + p_b |\uparrow_{\vec{b}}\rangle \langle \uparrow_{\vec{b}}|$$

Da jeder qubit Zustand sich in der Form (51) schreiben läßt, gilt

(58)

$$\rho = \rho_+ |\uparrow_{\vec{c}}\rangle \langle \uparrow_{\vec{c}}| + \rho_- |\downarrow_{\vec{c}}\rangle \langle \downarrow_{\vec{c}}|$$

mit
$$\rho_{\pm} = \frac{1}{2} (1 \pm |\langle \uparrow_{\vec{a}} | \uparrow_{\vec{b}} \rangle|)$$

$\Rightarrow$

$$(59) \quad H_N = -S_+ \log S_+ - S_- \log S_- \leq 1$$

• Messung: z.B. Richtung $\vec{a}$

$$(60) \quad H_{\text{Mess}} = -p_{a+} \log p_{a+} - p_{a-} \log p_{a-} \leq 1$$

$$p_{a+} = \langle T_{\vec{a}} | S | T_{\vec{a}} \rangle = p_a + p_b |\langle T_{\vec{a}} | T_{\vec{b}} \rangle|^2$$

$$p_{a-} = p_b |\langle T_{\vec{a}} | T_{\vec{b}} \rangle|^2$$

$$(61) \Rightarrow \boxed{0 \leq H_N \leq H_{\text{Mess}}, \quad H_{\text{prep}} \leq 1}$$

Übertragung klassischer Information über einen Quantenkanal, d.h. mittels qubits ist nur dann ohne Verlust von Information möglich, falls die

(i) klassische Information in orthogonalen Zuständen kodiert wird

$$H_N = H_{\text{prep}}$$

(ii) bei der Messung die dabei verwendete Basis bekannt ist

$$H_N = H_{\text{Mess}}$$

Diese Besonderheit der klass. Informationsübertragung willtes Quantenkandidaten werden wir uns in der Quantenkryptographie zu Nutze machen.

3.6. Quantendatenkompression

In der klass. Inf.-theorie hatten wir Shannons noiseless coding theorem als nachträgliche Rechtfertigung dafür kennengelernt, daß die Shannon Entropie ein vernünftiges Informationsmaß ist. Jetzt wollen wir uns die Frage stellen: Was ist das Quantenanalogen von Shannons Theorem? Dazu müssen wir zwei Fälle unterscheiden

(a) Übertragung von reinen Zuständen

(b) Übertragung von gemischten Zuständen

(a) Quantendatenkompression für reine Zustände

Nehmen wir an Alice benutzt ein Alphabet aus reinen Zuständen $\{|\varphi_x\rangle\langle\varphi_x|\}$ die nicht notwendig orthogonal sein müssen. In einem (langen) Wort taucht $|\varphi_x\rangle$ mit Wahrscheinlichkeit p_x auf, d.h. wir haben

$$(62) \quad \rho = \sum_x p_x |\varphi_x\rangle\langle\varphi_x|$$

Welche Dimension müßte ein Hilbertraum haben,

um ρ exakt zu speichern? Genaue:
Wie groß muß die Dimension des Hilbertraums
sein um n Kopien von ρ zu speichern?

Schumacher $n \rightarrow \infty$ reine Zustände

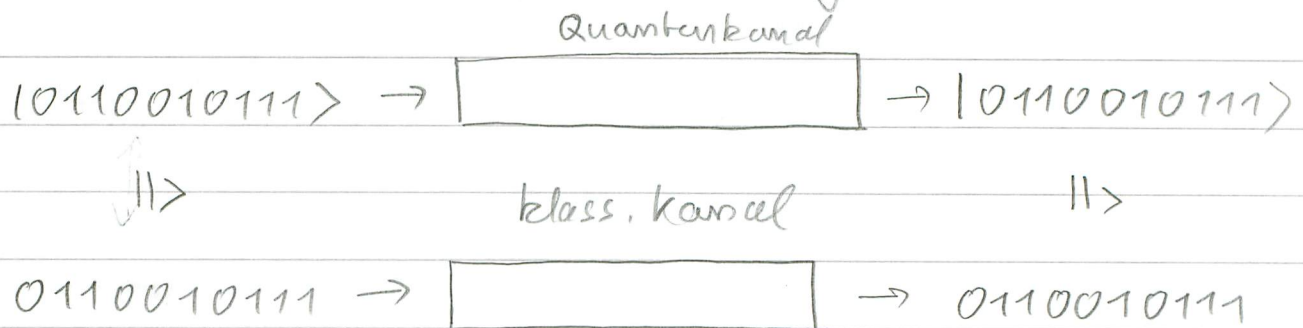
$$(63) \quad \log(\dim \mathcal{H}) = n S_N(\rho)$$

D.h. um n qubits über einen verlustfreien
Kanal zu schicken benötigt man im Mittel
nur $n H_N(\rho)$ qubits.

Der Beweis des Theorems von Schumacher ist
einfach wenn die Zustände des Alphabets orthogonal
sind. In diesem Fall können wir einfach die 1:1
Abbildung auf eine klassische Zufallsvariable
und $S_N(\rho) = S_{Sh}(X)$ verwenden

$$\{|\varphi_x\rangle, P(\varphi_x)\} \iff \{X, P_X\}$$

sowie Shannons noiseless coding Theorem



Wenn die Zustände $|\varphi_x\rangle$ nicht orthogonal sind, ist Schumachers Theorem nicht ganz so einseitig.

Beispiel: Alice Alphabet

$$(64) \quad |\uparrow_z\rangle = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \quad |\uparrow_x\rangle = \begin{pmatrix} 1/\sqrt{2} \\ 1/\sqrt{2} \end{pmatrix}$$

Beide treten mit $p_x = p_z = \frac{1}{2}$ auf. Die Entropy of preparation H_{prep} ist offensichtlich gleich 1.

Um die von Neumann Entropie zu berechnen diagonalisieren wir den Dichteparameter

$$\rho = \frac{1}{2} |\uparrow_z\rangle\langle\uparrow_z| + \frac{1}{2} |\uparrow_x\rangle\langle\uparrow_x|$$

$$(65) \quad = \begin{pmatrix} \frac{3}{4} & \frac{1}{4} \\ \frac{3}{4} & \frac{1}{4} \end{pmatrix}$$

ρ hat die Eigenvektoren und Eigenwerte

$$(66) \quad |0'\rangle = \begin{pmatrix} \cos \frac{\pi}{8} \\ \sin \frac{\pi}{8} \end{pmatrix} \quad \rho_{0'} = \cos^2 \frac{\pi}{8}$$

$$(67) \quad |1'\rangle = \begin{pmatrix} -\sin \frac{\pi}{8} \\ \cos \frac{\pi}{8} \end{pmatrix} \quad \rho_{1'} = \sin^2 \frac{\pi}{8}$$

Somit ist

$$(68) \quad H_N(\rho) = -\text{Tr}\{\rho \log \rho\} = -\rho_{0'} \log \rho_{0'} - \rho_{1'} \log \rho_{1'}$$

$$= 0,600876 \dots < 1$$

Der Überlapp von $|0'\rangle$ mit $|1_x\rangle$ und $|1_z\rangle$ ist deutlich größer als der von $|1'\rangle$

$$(69) \quad |\langle 0' | 1_x \rangle|^2 = |\langle 0' | 1_z \rangle|^2 = \cos^2 \frac{\pi}{8} = 0,8535$$

$$|\langle 1' | 1_x \rangle|^2 = |\langle 1' | 1_z \rangle|^2 = \sin^2 \frac{\pi}{8} = 0,1465$$

Wenn Alice ein qubit Wort der Länge 3 an Bob senden will

$$(70) \quad |2\rangle = |2_1\rangle |2_2\rangle |2_3\rangle$$

wobei jedes der qubits $|2_i\rangle$ entweder $|1_x\rangle$ oder $|1_z\rangle$ darstellt kann sie dies am effektivsten tun indem sie nur die 4 Zustände

$$(71) \quad \{|0',0',0'\rangle, |1',0',0'\rangle, |0',1',0'\rangle, |0',0',1'\rangle\}$$

verwendet anstelle der $2^3 = 8$ möglichen 3er Kombinationen aus $|1_x\rangle$ und $|1_z\rangle$. Man rechnet leicht nach, daß gilt

$$(72) \quad |\langle 0',0',0' | 2 \rangle|^2 = \left[\cos^2 \left(\frac{\pi}{8} \right) \right]^3 = 0,6219$$

$$|\langle 1',0',0' | 2 \rangle|^2 = |\langle 0',1',0' | 2 \rangle|^2 = |\langle 0',0',1' | 2 \rangle|^2 = \left[\cos^2 \left(\frac{\pi}{8} \right) \right]^2 \sin^2 \frac{\pi}{8} = 0,1067$$

D.h. die Wahrscheinlichkeit ein beliebiges 3 qubit Wort in dem durch (71) definierten Unterraum zu finden ist

$$(73) \quad \text{Erfolg} = 0,6219 + 3 \cdot 0,1067 = 0,9419$$

Im asymptotischen Grenzfall von n -qubit Wörtern mit $n \rightarrow \infty$ kann man einen Quantencode aus Zuständen $|0'\rangle$ und $|1'\rangle$ finden (ähnlich zu (71)) mit $n \cdot H(\cos^2 \frac{\pi}{8}) \approx 0,600876 n$ qubits und

$$(74) \quad \text{Erfolg} \xrightarrow{n \rightarrow \infty} 1$$

(b) Quantendatenkompression für gemischte Zustände

Jetzt wollen wir den Fall betrachten, daß das Alphabet von Zuständen das Alice benutzt aus gemischten Zuständen besteht. Gilt dann ebenfalls Schumachers Theorem (63)? Man kann leicht sehen daß das nicht der Fall ist:

Bsp1: Alice verwendet nur einen "Buchstaben" (qubit Zustand) ρ_0 mit $S_N(\rho_0) > 0$. Alice Nachricht hat also stets die Form

$$(75) \quad \rho_0 \otimes \rho_0 \otimes \rho_0 \otimes \rho_0 \dots$$

In dieser Nachricht steckt gar keine Information und es wird überhaupt kein Quantenzustand zur Übertragung benötigt, obwohl $S_N(\rho_0) > 0$.

Um eine Verallgemeinerung von Schumachers Theorem für gemischte Zustände zu finden, erinnern wir uns daran, daß für reine Zustände gilt

$$(76) \quad S_N(\{|\psi_i\rangle\}) = S_{Sh}(p(\psi_i))$$

falls die Zustände $|\psi_i\rangle$ die mit Wahrscheinlichkeit $p(\psi_i)$ im Alphabet auftreten (d.h. $\sum p(\psi_i) |\psi_i\rangle \langle \psi_i|$) orthogonal sind. Besteht das Alphabet nun aus orthogonalen gemischten Zuständen ρ_i , d.h. ist

$$(77) \quad \text{Tr} \{ \rho_i \rho_j \} = 0 \quad \text{für } i \neq j$$

so sollte die mögliche Kompression ebenfalls durch die Shannon Entropie $S_{Sh}(p(\rho_i))$ gegeben sein,

$$\text{Istri} \quad S_N(\rho) = S_{Sh}(p(\rho_i)) = S_{Sh}(p(\rho_i))$$

$$(78) \quad \rho = \sum_i p_i \rho_i$$

und die mögliche Kompression mit Zuständen ρ_i gilt

$$(79) \quad \begin{aligned} S_N(\rho) &= -\text{Tr} \{ \rho \log \rho \} \\ &= -\sum_i \text{Tr} \{ p_i \rho_i \log (p_i \rho_i) \} \end{aligned}$$

$$= - \sum_i p_i \log p_i - \sum_i p_i \operatorname{Tr} \{ s_i \log s_i \}$$

$$= S_{\text{Sh}}(p(s_i)) + \sum_i p_i S_N(s_i)$$

Die Größe

$$(80) \quad S_{\text{Sh}}(p(s_i)) = S_N(s) - \sum_i p_i S_N(s_i) \equiv \mathcal{X}(\mathcal{E})$$

heißt Holevo Information und stellt eine Verallgemeinerung der von Neumann Entropie dar. $\mathcal{X}(\mathcal{E})$ hängt offensichtlich nicht nur vom Quantenzustand s sondern auch von der Art wie s durch ein Ensemble von gemischten Zuständen zusammengesetzt ist, d.h. $\mathcal{E} = \{p_i, \hat{s}_i\}$. Für ein Alphabet aus reinen Zuständen ist $\mathcal{X}(\mathcal{E}) = S_N(s)$.

Holevo Vermutung

Zur Übertragung eines n qubit Wortes aus gem. Zuständen s_i die mit Wahrscheinlichkeit p_i in dem qubit Wort auftreten sind im Limit $n \rightarrow \infty$

$$n \mathcal{X}(\{p_i, \hat{s}_i\}) \leq n S_N(s = \sum_i p_i s_i)$$

qubits notwendig