

3.7. Quantenfehlerkorrektur

Im Gegensatz zur klassischen Theorie der Informationsübertragung gibt es kein "noisy coding theorem" für Quanteninformation mit allgemeiner Gültigkeit. Auch die Frage nach der Kapazität eines Quantenkanals ist nicht vollständig geklärt. Daher werden wir hier keine quantenmechanische Variante von Shannons noisy coding theorem diskutieren. Es gibt jedoch eine ausgearbeitete Theorie der Quantenfehlerkorrektur die im Folgenden kurz skizziert werden soll.

Besonderheiten einer Quantenfehlerkorrektur

- (i) bei qubits gibt es mehr als nur bit-flip Fehler
- (ii) der Zustand eines qubits darf während einer Übertragung nicht gemessen werden

qubit - Fehler

qubit-Fehler entstehen durch unkontrollierte Wechselwirkung mit der Umgebung

$$\text{qubit } \{|0\rangle, |1\rangle\} \longleftrightarrow \text{Umgebung } \{|E\rangle\}$$

$$\begin{aligned}
 (8') &= (\alpha|0\rangle + \beta|1\rangle) |E_0\rangle \\
 &+ (\alpha|1\rangle + \beta|0\rangle) |E_x\rangle \\
 &- (\alpha|1\rangle - \beta|0\rangle) |E_y\rangle \\
 &+ (\alpha|0\rangle - \beta|1\rangle) |E_z\rangle
 \end{aligned}$$

$$\left. \begin{aligned}
 |E_0\rangle + |E_z\rangle &= |E_{00}\rangle \\
 |E_x\rangle - |E_y\rangle &= |E_{01}\rangle \\
 |E_x\rangle + |E_y\rangle &= |E_{10}\rangle \\
 |E_0\rangle - |E_z\rangle &= |E_{11}\rangle
 \end{aligned} \right\} \begin{aligned}
 |E_0\rangle &= \frac{1}{2} (|E_{00}\rangle + |E_{11}\rangle) \\
 |E_z\rangle &= \frac{1}{2} (|E_{00}\rangle - |E_{11}\rangle) \\
 |E_x\rangle &= \frac{1}{2} (|E_{01}\rangle + |E_{10}\rangle) \\
 |E_y\rangle &= -\frac{1}{2} (|E_{01}\rangle - |E_{10}\rangle)
 \end{aligned}$$

Spin flip

$$(81) \quad \begin{aligned} |0\rangle |E\rangle &\longrightarrow |1\rangle |E_{01}\rangle \\ |1\rangle |E\rangle &\longrightarrow |0\rangle |E_{10}\rangle \end{aligned}$$

Dephasierung

$$(82) \quad \begin{aligned} |0\rangle |E\rangle &\longrightarrow |0\rangle |E_{00}\rangle \\ |1\rangle |E\rangle &\longrightarrow |0\rangle |E_{11}\rangle \end{aligned}$$

Im Laufe der Wechselwirkung werden die Zustände der Umgebung zunehmend orthogonal

$$(83) \quad \begin{aligned} \langle E_{01}(t) | E_{10}(t) \rangle &\xrightarrow{t \rightarrow \infty} 0 \\ \langle E_{00}(t) | E_{11}(t) \rangle &\xrightarrow{t \rightarrow \infty} 0 \end{aligned}$$

Das führt zu folgender Evolution des reduzierten Dichteoperators $|\psi_0\rangle = \alpha|0\rangle + \beta|1\rangle$

$$\rho_0 = \begin{pmatrix} |\alpha|^2 & \alpha\beta^* \\ \alpha^*\beta & |\beta|^2 \end{pmatrix} \xrightarrow[\text{flip}]{\text{spin}} \begin{pmatrix} |\beta|^2 & \alpha^*\beta \langle E_{01}(t) | E_{10}(t) \rangle \\ \alpha\beta^* \langle E_{10}(t) | E_{01}(t) \rangle & |\alpha|^2 \end{pmatrix}$$

$$(84) \quad \xrightarrow{\text{Dephasierung}} \begin{pmatrix} |\alpha|^2 & \alpha\beta^* \langle E_{00}(t) | E_{11}(t) \rangle \\ \alpha^*\beta \langle E_{11}(t) | E_{00}(t) \rangle & |\beta|^2 \end{pmatrix}$$

In beiden Fällen werden die Nichtdiagonalelemente kleiner. Man spricht von Dekohärenz.

Im Allgemeinen führt die Wechselwirkung mit der Umgebung zu

$$(85) \quad \begin{array}{l} |0\rangle|E\rangle \\ |1\rangle|E\rangle \end{array} \xrightarrow{U(H)} \begin{array}{l} |0\rangle|E_{00}\rangle + |1\rangle|E_{01}\rangle \\ |0\rangle|E_{10}\rangle + |1\rangle|E_{11}\rangle \end{array}$$

Die Wirkung auf einen allgemeinen qubit Zustand $|\psi_0\rangle = \alpha|0\rangle + \beta|1\rangle$ kann in der Form

$$(86) \quad \begin{aligned} (\alpha|0\rangle + \beta|1\rangle)|E\rangle &\rightarrow (\alpha|0\rangle + \beta|1\rangle)|E_0\rangle \\ &+ (\alpha|1\rangle + \beta|0\rangle)|E_x\rangle \\ &+ (\alpha|0\rangle - \beta|1\rangle)|E_z\rangle \\ &+ (-\alpha|1\rangle + \beta|0\rangle)|E_y\rangle \\ &= (\hat{I}|\psi_0\rangle)|E_0\rangle + (\hat{X}|\psi_0\rangle)|E_x\rangle + (\hat{Y}|\psi_0\rangle)|E_y\rangle \\ &+ (\hat{Z}|\psi_0\rangle)|E_z\rangle \end{aligned}$$

geschrieben werden, wobei

$$(87) \quad X = \begin{pmatrix} 0 & 1 \\ 1 & 0 \end{pmatrix} \quad Y = \begin{pmatrix} 0 & -i \\ i & 0 \end{pmatrix} \quad Z = \begin{pmatrix} 1 & 0 \\ 0 & -1 \end{pmatrix}$$

spin flip phasenfehler

bedeuten und $Y = ZX$ eine Kombination aus spin flip und Phasenfehler darstellt.

Ein fehlerkorrigierender Quantencode muß neben einem bit-flip Fehler $\hat{X}$ auch noch Phasenfehler $\hat{Z}$ und deren Kombination $\hat{Y} = \hat{Z}\hat{X}$ korrigieren.

Um Bedingung (ii) zu genügen wird ein qubit in verschränkten n -qubit Zuständen kodiert, d.h. der Hilbertraum wird durch Anzillazustände erweitert. Die Detektion des Fehlersignals erfolgt dann durch eine Messung an den Anzillazuständen.

3-qubit code

korrigiert 1 spin-flip Fehler $\hat{X}$ (Wahrscheinlichkeit p)

$$(88) \quad \hat{M} = (1-p) \hat{I} + p \hat{X}$$

Kodierung: binärer Lin. Code $k=1, n=3, d=3$

$$(89) \quad |0\rangle \rightarrow |000\rangle \quad |1\rangle \rightarrow |111\rangle$$

Diese Kodierung kann durch zweimaliges Ausführen einer sog. controlled-NOT Operation an dem qubit und zwei Anzillazuständen $|00\rangle$ erzeugt werden ohne daß der qubit Zustand gemessen werden muß

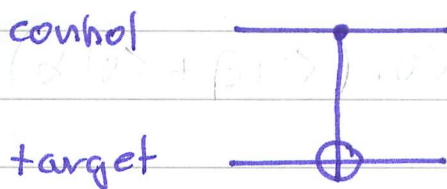
$$(90) \quad \text{CNOT} \quad \begin{array}{c} \text{Projektoren für "control bit"} \\ |0\rangle\langle 0| \otimes \hat{I} + |1\rangle\langle 1| \otimes \hat{X} \\ \text{Operation auf "target bit"} \end{array}$$

$$(91) \quad \text{CNOT } |x\rangle|y\rangle = |x\rangle|y+x \bmod 2\rangle$$

control target

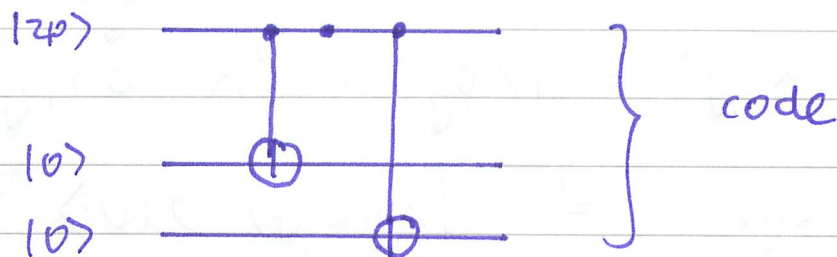
	$ 00\rangle$	$\longrightarrow$	$ 00\rangle$
	$ 01\rangle$	$\longrightarrow$	$ 01\rangle$
(92)	$ 10\rangle$	$\longrightarrow$	$ 11\rangle$
	$ 11\rangle$	$\longrightarrow$	$ 10\rangle$

graphische Darstellung



• Kodierung:

$$(93) (\alpha|0\rangle + \beta|1\rangle)|00\rangle \xrightarrow{2 \times \text{CNOT}} \alpha|000\rangle + \beta|111\rangle$$



• Wechselwirkung mit Umgebung

	$\alpha 000\rangle + \beta 111\rangle$	$(1-p)^3$
	$\alpha 100\rangle + \beta 011\rangle$	$p(1-p)^2$
(94)	$\alpha 010\rangle + \beta 101\rangle$	$p(1-p)^2$
	$\alpha 001\rangle + \beta 110\rangle$	$p(1-p)^2$
	andere	$O(p^2)$

Endzustände

Wahrscheinlichkeit

klass. Fehler $\underline{v} \rightarrow \underline{v} \oplus \underline{e} = \underline{v}'$

klass. Fehlerdefinition

$$\underline{H}(\underbrace{\underline{v} \oplus \underline{e}}_{\underline{v}'}) = \underline{H} \underline{e}$$

$\hat{X}$ Fehler im i-ten Bit $\hat{X}_i$

$$\hat{X}_2 |x0y\rangle = |x1y\rangle = |x, 0 \oplus 1, y\rangle$$

↑
Addition mod 2

$$\hat{X}_2 |x1y\rangle = |x0y\rangle = |x, 1 \oplus 1, y\rangle$$

$$\hat{X}_2 |xyz\rangle \stackrel{?}{=} |xyz \oplus 010\rangle$$

Spin-Flop Fehler abbildbar auf Addition eines
klass. Fehlervektors zum klass. Vektor der Bit Werte

$$\underline{H} \overset{e_1}{\begin{pmatrix} 1 \\ 0 \\ 0 \end{pmatrix}} = \begin{pmatrix} 1 \\ 1 \end{pmatrix} \quad \underline{H} \overset{e_2}{\begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix}} = \begin{pmatrix} 1 \\ 0 \end{pmatrix} \quad \underline{H} \overset{e_3}{\begin{pmatrix} 0 \\ 0 \\ 1 \end{pmatrix}} = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$$

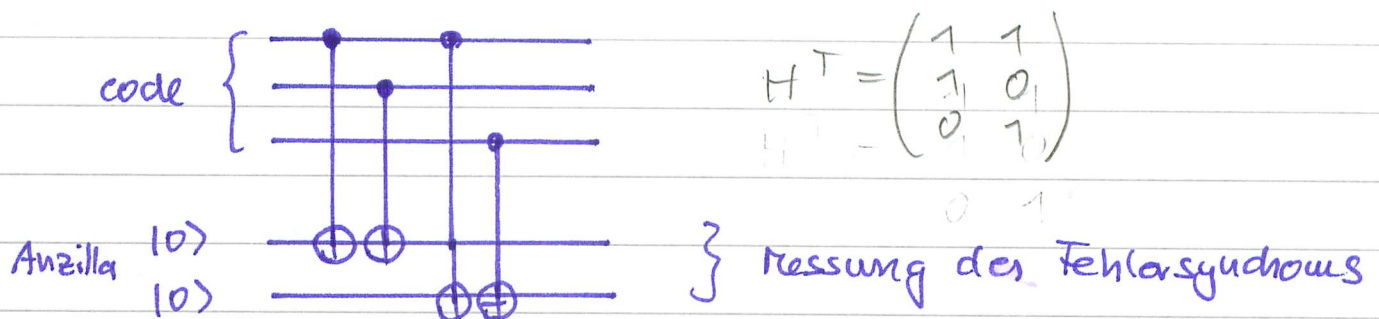
$$\underline{H} \underline{v}_i = \begin{pmatrix} 0 \\ 1 \end{pmatrix}$$

• Fehlerdetektion

Der Code (89) ist linear aus $n=3$ qubit der $k=1$ qubit kodiert. Der Hammingabstand beträgt $d=3$, D.h. es existiert eine $(n-k=2)$ $\times$ $(n=3)$ parity check Matrix

$$(95) \quad \underline{H} = \begin{pmatrix} 1 & 1 & 0 \\ 1 & 0 & 1 \end{pmatrix}$$

$\underline{H}$ hat die wichtige Eigenschaft daß sie angewandt auf beide codewörter $(0,0,0)$ und $(1,1,1)$ null ergibt. Dies kann angewendet werden um das Fehlersyndrom zu detektieren ohne den kodierten Zustand zu messen



Multiplikation mit 2 Anzillazuständen $|0,0\rangle$ und CNOT 1.+2. code bit mit 1. Anzillabit (gemäß erstem Vektor aus $\underline{H}$) sowie 1.+3. codebit mit 2. Anzillabit (gemäß zweitem Vektor aus $\underline{H}$) liefert

$$(96) \quad \begin{aligned} &(\alpha|000\rangle + \beta|111\rangle) && |00\rangle \\ &(\alpha|100\rangle + \beta|011\rangle) && |11\rangle \\ &(\alpha|010\rangle + \beta|101\rangle) && |10\rangle \\ &(\alpha|100\rangle + \beta|110\rangle) && |01\rangle \end{aligned}$$

Messung beider Anzillabits in der Basis $\{|0\rangle, |1\rangle\}$ liefert eindeutig Fehler-syndrom.

• Fehlerkorrektur (wie bei klassischem Code) kann durch unitäre Operationen

Nachdem der aufgetretene Fehler detektiert wurde kann er durch eine unitäre Operation korrigiert werden

$$(97) \quad \begin{array}{ll} |00\rangle \rightarrow \hat{I}_1 \hat{I}_2 \hat{I}_3 & |10\rangle \rightarrow \hat{I}_1 \hat{X}_2 \hat{I}_3 \\ |11\rangle \rightarrow \hat{X}_1 \hat{I}_2 \hat{I}_3 & |01\rangle \rightarrow \hat{I}_1 \hat{I}_2 \hat{X}_3 \end{array}$$

Allgemein gilt (wie bei klassischen Codes) die sog. Hamming Ungleichung:

Ein linearer Code für k qubits der 3 unabh. abhängige Fehler $(\hat{X}, \hat{Y}, \hat{Z})$ (klassisch nur $\hat{X}$) in y qubits korrigieren soll muß mindestens n qubits haben mit

$$(98) \quad 2^k \sum_{i=0}^y 3^i \binom{n}{i} \leq 2^n$$

Der kleinste Code der alle auftretenden Fehler $(\hat{X}, \hat{Y}, \hat{Z})$ in einem qubit korrigieren kann, muß also $(k=2, y=1)$

$$(99) \quad 2 \cdot 3 \binom{n}{1} = 6n \leq 2^n$$

$n=5$ qubits haben. Ein solcher Code ist konstruiert worden ist aber nicht sehr übersichtlich.

Man kann durch folgende Überlegung zeigen, daß kein $n=4$ qubit code existieren kann, da einen spin-flip Fehler an einer beliebigen Stelle korrigieren kann. Ein solcher würde nämlich das no-cloning Theorem verletzen: Man kann zeigen daß ein code der t spin-flip Fehler an beliebigen Stellen korrigiert, benutzt werden kann um $2t$ Fehler an wohldefinierten Stellen zu korrigieren. (Beispiel: Wenn ein 3-qubit code (89) ein Fehler nur im ersten und zweiten bit auftreten kann, kann die Menge der fehlerhaften codewörter in zwei aufgeteilt werden, die sich im 3. bit unterscheiden.

CNOT Operationen der ersten zwei bits mit einem Auxilla $|00\rangle$ erlaubt dann eindeutig festzustellen welcher Fehler aufgetreten ist.) Dann könnten wir ein einzelnes unbekanntes qubit in 4 qubits kodieren und diese in zwei Paare aufteilen. Jedes dieser Paare besteht aus 2 qubits und kann durch zwei weitere Auxilla qubits z.B. im Zustand $|00\rangle$ ergänzt werden. Die Auxilla bits entsprechen 2 bit Fehlern an bekannten Stellen. Der code sollte erlauben beide eindeutig zu korrigieren. Dann hätten wir aber zwei perfekte Klone erzeugt im Widerspruch zum no-cloning Theorem.

Um zu sehen wie neben dem spin-flip Fehler ($\hat{X}$) auch andere Fehler ($\hat{Y}, \hat{Z}$) korrigiert werden können, wollen wir eine spezielle Klasse von Codes betrachten die sog. CSS-Codes (Calderbank-Shor-Stean) einführen, die es erlauben diese Fehler

subsequentiv zu korrigieren

CSS codes

Sei C_1 ein (klassischer) linearer Code $[n, k_1]$
mit $(n-k_1) \times n$ parity check Matrix H_1 .
Nun sei C_2 ein Untercode von C_1 mit $k_2 < k_1$

$$(100) \quad C_2 \subset C_1$$

mit $(n-k_2) \times n$ parity check Matrix H_2 .

Alle Elemente w aus C_2 erfüllen die $(n-k_1)$ Bedingungen $H_1 w = 0$ aber darüberhinaus noch $k = k_1 - k_2$ extra Bedingungen (extra Zeilen in H_2).

Der Untercode C_2 definiert Äquivalenzklassen:

Def: $u, v \in C_1$ heißen äquivalent

$$(\Leftrightarrow) \exists w \in C_2 \text{ mit } u = v + w$$

Ein CSS code ist ein Quantencode der jedem codeword eine ganze Äquivalenzklasse zuordnet.
Es gibt genau $k = k_1 - k_2$ Äquivalenzklassen

Basisvektoren des CSS codes ($v \in C_1$ aus versch. Äq.-kl.)

$$(101) \quad |\bar{v}\rangle = \frac{1}{\sqrt{2^{k_2}}} \sum_{w \in C_2} |v+w\rangle$$

Falls v und v' in verschiedenen Äquivalenzklassen liegen, gilt

$$(102) \quad \langle \bar{v} | \bar{v}' \rangle = 0$$

Wir werden im Folgenden sehen, daß ein Phasenfehler ($\hat{Z}$) in einem CSS code korrigiert werden kann wie ein spin-flip Fehler ($\hat{X}$) wenn vorher eine bestimmte bitweise unitäre Transformation ausgeführt wird, die sog.

Hadamard Transformation

$$(103) \quad \text{Had} = \frac{1}{\sqrt{2}} (\hat{Z} + \hat{X}) = \frac{1}{\sqrt{2}} \begin{pmatrix} 1 & 1 \\ 1 & -1 \end{pmatrix}$$

$$(104) \quad \text{Had} |0\rangle \longrightarrow \frac{1}{\sqrt{2}} (|0\rangle + |1\rangle)$$

$$\text{Had} |1\rangle \longrightarrow \frac{1}{\sqrt{2}} (|0\rangle - |1\rangle)$$

graphische Darstellung



Man erkennt aus (104) daß ein Phasenfehler ($\hat{Z}$) einem spin-flip der Hadamard Transformation entspricht

$$(105) \quad |1\rangle \rightarrow -|1\rangle \quad \hat{=} \quad \text{Had} |0\rangle \leftrightarrow \text{Had} |1\rangle$$

7-qubit code (einfacher CSS code)

Ausgangspunkt: klassischer Hamming code $[n=7, k=4, d=3]$ mit parity-check Matrix $(n-k=3 \times 7)$

$$(106) \quad \underline{H} = \begin{pmatrix} 0 & 0 & 0 & 1 & 1 & 1 & 1 \\ 0 & 1 & 1 & 0 & 0 & 1 & 1 \\ 1 & 0 & 1 & 0 & 1 & 0 & 1 \end{pmatrix}$$

Die zugehörige erzeugende Matrix (Gl. 22) ist

$$(107) \quad \underline{G} = \begin{pmatrix} \underline{v}_1 \\ \underline{v}_2 \\ \underline{v}_3 \\ \underline{v}_4 \end{pmatrix} = \begin{pmatrix} 0 & 0 & 0 & 1 & 1 & 1 & 1 \\ 0 & 1 & 1 & 0 & 0 & 1 & 1 \\ 1 & 0 & 1 & 0 & 1 & 0 & 1 \\ 1 & 1 & 1 & 0 & 0 & 0 & 0 \end{pmatrix}$$

Wir sehen, daß $\underline{H}$ in $\underline{G}$ enthalten ist. (Da $\underline{H} \underline{v}_i^T = 0$ (Gl. 25) sieht man die in der üblichen linearen Algebra etwas ungewöhnliche Eigenschaft des $\mathbb{F}_2$ -Raumes $\mathbb{F}_2$ (bgl. Addition modulo 2) daß Vektoren zu sich selbst orthogonal sein können, was natürlich nur an der Addition modulo 2 liegt:

$$\underline{v}_1 \circ \underline{v}_1 = \underline{v}_2 \circ \underline{v}_2 = \underline{v}_3 \circ \underline{v}_3 = 0$$

Nun treffen wir folgende Zuordnung

$$(108) \quad \begin{array}{ll} C_1 : & \text{Generator } G \\ & \text{parity check } H \end{array} \quad 2^{k_1} = 16 \text{ Elemente}$$

(109) $C_2 \subset C_1$: Generator H (Untermatrix von G_1)
 parity check G_1 ; $2^{k_2} = 8$ Elemente
 gerader Untercode von C_1
 $[n=7, k_2=3, d=3]$

Der aus C_1 und C_2 konstruierte CSS code kodiert also $k_1 - k_2 = 1$ qubit.

7 qubit code

(109) $|\bar{0}\rangle = \frac{1}{\sqrt{2^3}} \sum_{v \in C_2} |0+v\rangle = \frac{1}{\sqrt{2^3}} \sum'_{v \in C_1} |u\rangle$
 mit gerader Anzahl "1"

(110)

(110) $|\bar{1}\rangle = \frac{1}{\sqrt{2^3}} \sum_{v \in C_2} |1+v\rangle = \frac{1}{\sqrt{2^3}} \sum'_{v \in C_1} |u\rangle$
 mit ungerader Anzahl "1"

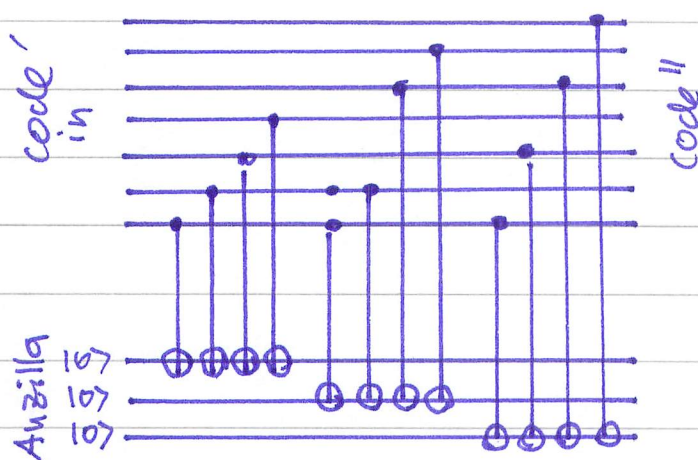
Da beide Elemente des 7 qubit codes $|\bar{0}\rangle$ und $|\bar{1}\rangle$ Superpositionen aus $|u\rangle$ mit $u \in C_1$ ($d=3$) kann ein einzelner spin flip detektiert und korrigiert werden. Das Verfahren dazu hatten wir beim 3-qubit code kennengelernt. Phasenfehler werden damit jedoch noch nicht korrigiert. Phasenfehler entsprechen aber einem spin-flip Fehler der Hadamard-transformierten qubits $\text{Had}|0\rangle$ und $\text{Had}|1\rangle$.

Eine bitweise Hadamard Transformation des 7-bit Codes liefert (wie man nicht ganz so schnell nachrechnet)

$$\begin{aligned} (111) \quad (\text{Had})^{(7)} \quad |\bar{0}\rangle &= \frac{1}{\sqrt{2}} (|\bar{0}\rangle + |\bar{1}\rangle) \equiv |\bar{0}\rangle_p \\ (\text{Had})^{(7)} \quad |\bar{1}\rangle &= \frac{1}{\sqrt{2}} (|\bar{0}\rangle - |\bar{1}\rangle) \equiv |\bar{1}\rangle_p \end{aligned}$$

Da sowohl $|\bar{0}\rangle_p$ als auch $|\bar{1}\rangle_p$ Superpositionen aus Zuständen $|u\rangle$ mit $u \in C_1$ sind, können wiederum ein einzelner Spin flip Fehler irgendeiner qubits detektiert und korrigiert werden. Da ein Spin flip Fehler in der Hadamard rotierten Basis einem Phasenfehler in der ursprünglichen Basis entspricht wird somit ein Phasenfehler korrigiert. Die Extraktions- und Korrekturprozedur sind identisch zu der für einen Spin flip Fehler. Zusammen ergibt sich folgendes Schema

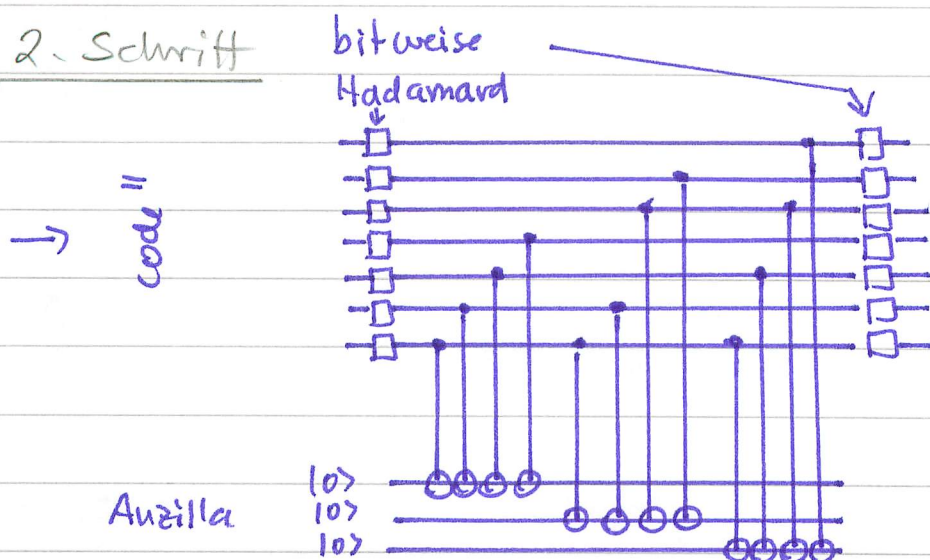
1. Schritt



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

Detektion von spin flip (max. 1)

2. Schritt



Detection von spin-flips in Hadamard-transformierter Basis = Detection von Phasenfehlern in ursprünglicher Basis

Der 7 qubit code kann einen einzelnen $\hat{X}$ Fehler und einen einzelnen $\hat{Z}$ Fehler korrigieren. Treffen beide am selben qubit auf entspricht das der Komplexität eines $\hat{Y}$ Fehlers.

Stabilisator codes

Um einen $[5,1,3]$ Quantencode zu konstruieren benötigt man ein effizientere Prozedur als das CSS Verfahren. Ein solches sind die sog. Stabilisator codes die ausnutzen, daß die Menge der Fehleroperatoren

$$(112) \quad \{\pm I, \pm X, \pm Y, \pm Z\} = \pm \{I, X, Y, Z\}$$

eine Gruppe bilden. Das n -fache Tensorprodukt

der Paulioperatoren bildet ebenfalls eine Gruppe

$$(113) \quad G_n = \pm \{I, X, Y, Z\}^{\otimes n}$$

der Ordnung 2^{2n+1} . Die einzige nichttriviale Darstellung von G_n sind 2^n dimensionale Tensorprodukte aus 2×2 Matrizen. Die Elemente dieser Darstellung haben die Form (z.B.)

$$(114) \quad M = \underbrace{XYIZY \dots X}_n$$

Die M 's haben die Eigenschaften

$$(i) \quad M^{-1} = M^+$$

$$(ii) \quad M^2 = \pm I^{\otimes n}$$

man beachte

$$X^2 = Z^2 = I \text{ aber } Y^2 = -I$$

$$M^2 = +I^{\otimes n} \text{ bei gerader Anzahl von } Y$$

$$M^2 = -I^{\otimes n} \text{ bei ungerader Anzahl von } Y$$

$$(iii) \quad M \text{ ist entweder hermitisch oder anti-hermitisch}$$

$$\text{man beachte } X^+ = X, Z^+ = Z, Y^+ = -Y$$

$$M^2 = I \Rightarrow M \text{ hermitisch}$$

$$M^2 = -I \Rightarrow M \text{ anti-hermitisch}$$

(iv) es gilt $M_1 M_2 = \pm M_2 M_1$ d.h. die Elemente der Darstellung kommutieren oder anti-kommutieren.

Sei S eine abelsche (kommutative) Untergruppe von G_n . Dann können alle Elemente M aus S gleichzeitig diagonalisiert werden. Die Menge der gleichzeitigen Eigenzustände aller Elemente M aus S zum Eigenwert $+1$ ist der Stabilisator code.

D.h. $|2^n\rangle \in \text{Stabilisator code}$

(115)

$$\Leftrightarrow M |2^n\rangle = |2^n\rangle \quad \forall M$$

Die abelsche Gruppe S wird durch ihre Generatoren $\{M_i\}$ charakterisiert. Jedes Element aus S kann durch Produkte der Generatoren produziert werden, aber kein Generator durch Produkte der anderen.

Man kann zeigen: Besitzt S $n-k$ Generatoren dann hat der Code-Raum die Dimension 2^k , d.h. k qubits können kodiert werden.

Jeder Fehler E_a ist selber wieder ein Element der Pauli-Gruppe G_n . D.h. für jedes E_a gilt es kommutiert oder antikommutiert mit den Generatoren M_i .

D. h.

$$(11b) \quad M_i E_a = (-1)^{S_{ia}} E_a M_i$$

$S_{ia} = \pm 1$. Die S_{ia} $i = 1, \dots, n-k$ bilden das Fehlensyndrom.

Fehlerdetektion ist dann eindeutig möglich wenn für jedes Paar von Fehlern E_a, E_b folgendes gilt

(i) entweder $E_a^\dagger E_b \in S$

(ii) oder $\exists M \in S$ das mit $E_a^\dagger E_b$ antikommutiert

Beweis: zu (i) Fehler a $|\psi\rangle \rightarrow E_a |\psi\rangle$
Fehler b $|\psi\rangle \rightarrow E_b |\psi\rangle$

$$\langle \psi | E_a^\dagger E_b | \psi \rangle \stackrel{(i)}{=} \langle \psi | \psi \rangle = 1$$

Die auftretenden Fehler verschieben alle Codewörter auf die gleiche Art.

$$\text{zu (ii)} \quad \langle \psi | E_a^\dagger E_b | \psi \rangle \stackrel{(115)}{=} \langle \psi | E_a^\dagger E_b M | \psi \rangle$$

$$\stackrel{(ii)}{=} - \langle \psi | M E_a^\dagger E_b | \psi \rangle$$

$$\stackrel{(115)}{=} - \langle \psi | E_a^\dagger E_b | \psi \rangle = 0$$

Man kann zeigen, daß "gute" stabilisator codes existieren, die folgende Eigenschaft haben

$$[n, k, d=2t+1]$$

$$n \rightarrow \infty \quad R = \frac{k}{n} \rightarrow \text{endlicher Wert} \\ (\text{bit-Rate})$$

$$p = \frac{t}{n} \rightarrow \text{endlicher Wert} \\ (\text{max. Fehlerwahrscheinlichkeit pro qubit})$$

Insbesondere kann man zeigen, daß gute codes existieren für $p < 0.0946$.

Bem: Bei allen Überlegungen in diesem Kapitel wurde angenommen, daß alle Operationen zur Detektion und Korrektur von Fehlern perfekt implementierbar sind. Es ist möglich, die Ideen der Quantenfehlerkorrektur auf imperfekte Operationen zu erweitern. Man spricht dann von fehlertoleranter Fehlerdetektion und Korrektur. Man kann zeigen, daß dafür gute codes existieren falls

$$p < 1 - 10^{-4}$$