

Vorlesungsskript
Kryptologie
Sommersemester 2018

Prof. Dr. Johannes Köbler
Humboldt-Universität zu Berlin
Lehrstuhl Komplexität und Kryptografie

21. Juli 2016

Inhaltsverzeichnis

1	Kryptografische Hashverfahren	1
1.1	Einführung	1
1.2	Schlüssellose Hashfunktionen (MDCs)	3
1.2.1	Vergleich von Sicherheitsanforderungen	4
1.2.2	Das Zufallsorakelmodell (ZOM)	5
1.2.3	Iterierte Hashfunktionen	8
1.2.4	Die Merkle-Damgaard-Konstruktion	9
1.2.5	Die MD4-Hashfunktion	10
1.2.6	Die MD5-Hashfunktion	11
1.2.7	Die SHA-1-Hashfunktion	12
1.2.8	Die SHA-2-Familie	13
1.2.9	Kryptoanalyse von Hashfunktionen	14
1.2.10	Die Sponge-Konstruktion	15
1.2.11	SHA-3	17
1.3	Nachrichten-Authentikationscodes (MACs)	18
1.3.1	Angriffe gegen symmetrische Hashfunktionen	19
1.3.2	Informationstheoretische Sicherheit von MACs	19

1 Kryptografische Hashverfahren

1.1 Einführung

Durch kryptographische Verfahren lassen sich unter anderem die folgenden **Schutzziele** realisieren.

- *Vertraulichkeit*
 - Geheimhaltung
 - Anonymität (z.B. Mobiltelefon)
 - Unbeobachtbarkeit (von Transaktionen)
- *Integrität*
 - von Nachrichten und Daten
- *Zurechenbarkeit*
 - Authentikation
 - Unabstreitbarkeit
 - Identifizierung
- *Verfügbarkeit*
 - von Daten
 - von Rechenressourcen
 - von Informationsdienstleistungen

Kryptografische Hashverfahren sind ein wirksames Werkzeug zur Sicherstellung der Integrität von Nachrichten oder generell von digitalisierten Daten. Sie nehmen somit beim Schutz der Datenintegrität eine ähnlich herausragende Stellung ein wie sie Kryptosystemen bei der Wahrung der Vertraulichkeit zukommt. Daneben finden kryptografische Hashfunktionen aber auch vielfach als Bausteine von komplexeren Systemen Verwendung. Wie wir noch sehen werden, sind kryptografische Hashfunktionen etwa bei der Erstellung von digitalen Signaturen sehr nützlich. Auf weitere Anwendungsmöglichkeiten werden wir später eingehen.

Vielen Anwendungen von kryptografischen Hashfunktionen h liegt die Idee zugrunde, dass sie zu einem vorgegebenen Text x eine zwar kompakte aber dennoch repräsentative Darstellung $h(x)$ liefern, die unter praktischen Gesichtspunkten als eine eindeutige Identifikationsnummer von x fungieren kann. Die Berechnungsvorschrift für h muss somit „charakteristische Merkmale“ von x in den Hashwert $h(x)$ einfließen lassen. Da der Fingerabdruck eines Menschen ganz ähnliche Eigenschaften besitzt (was ihn für Kriminalisten bekanntlich so wertvoll macht), wird der Hashwert $h(x)$ auch oft als ein **digitaler Fingerabdruck** von x bezeichnet. Gebräuchlich sind auch die Bezeichnungen **kryptografische Prüfsumme** oder *message digest* (englische Bezeichnung für „Nachrichtenextrakt“).

Typische Schutzziele, die sich mittels Hashfunktionen realisieren lassen, sind die Nachrichten- und Teilnehmerauthentikation.

- „Nachrichtenauthentikation“ (message authentication)

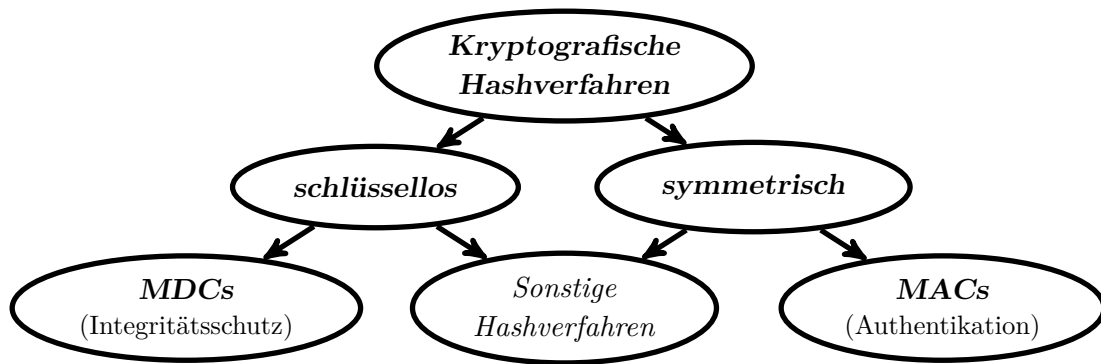


Abbildung 1.1: Eine grobe Einteilung von kryptografischen Hashverfahren.

- Wie lässt sich sicherstellen, dass eine Nachricht (oder eine Datei) während einer (räumlichen oder auch zeitlichen) Übertragung nicht verändert wurde?
- Wie lässt sich der Urheber (oder Absender) einer Nachricht zweifelsfrei feststellen?
- „Teilnehmerauthentikation“ (entity authentication, identification)
 - Wie kann sich eine Person (oder ein Gerät) anderen gegenüber zweifelsfrei ausweisen?

Klassifikation von Hashverfahren

Kryptografische Hashverfahren lassen sich grob danach klassifizieren, ob der Hashwert lediglich in Abhängigkeit vom Eingabetext berechnet wird oder zusätzlich von einem symmetrischen Schlüssel abhängt (siehe Abbildung 1.1).

Kryptografische Hashfunktionen, bei deren Berechnung keine Schlüssel benutzt werden, dienen vornehmlich der Erkennung von unbefugt vorgenommenen Manipulationen an Dateien oder Nachrichten. Daher werden sie auch als **MDC** bezeichnet (**Manipulation Detection Code** [englisch] = Code zur Erkennung von Manipulationen). Zuweilen wird das Kürzel **MDC** auch als eine Abkürzung für **Modification Detection Code** verwendet. Seltener ist dagegen die Bezeichnung **MIC** (**message integrity codes**). Abbildung 1.2 zeigt eine typische Anwendung von MDCs.

Um die Integrität eines Datensatzes x sicherzustellen, der über einen ungesi-

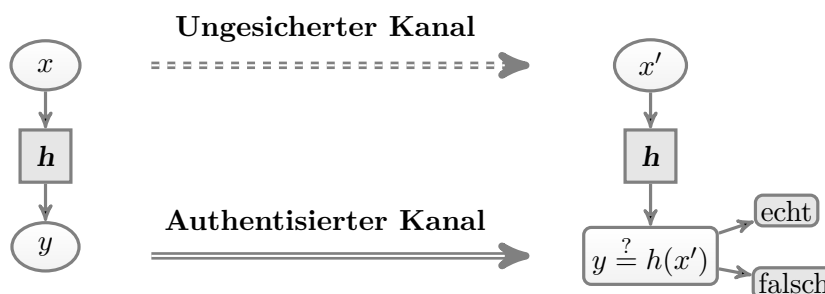


Abbildung 1.2: Einsatz eines MDC h zur Überprüfung der Integrität eines Datensatzes x .

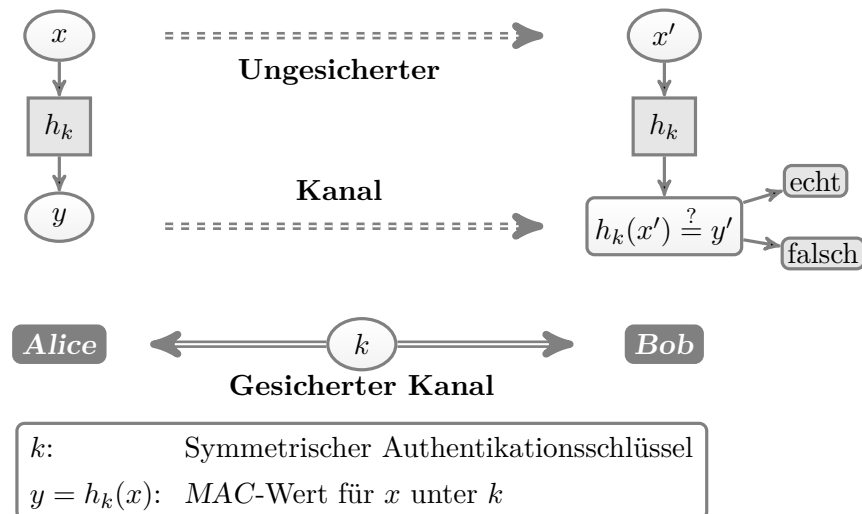


Abbildung 1.3: Verwendung eines MAC zur Nachrichtenauthentikation.

cherten Kanal gesendet (bzw. auf einem vor Manipulationen nicht sicheren Webserver abgelegt) wird, kann man wie folgt verfahren. Man sendet den MDC-Hashwert von x über einen authentisierten Kanal und prüft, ob der Datensatz nach der Übertragung noch denselben Hashwert liefert.

Kryptografische Hashverfahren mit symmetrischen Schlüsseln finden hauptsächlich bei der Authentifizierung von Nachrichten Verwendung. Diese werden daher auch als **MAC** (**m**essage **a**uthentication **c**ode [englisch] = Code zur Nachrichtenauthentifizierung) oder als **Authentifikationscode** bezeichnet. Daneben gibt es auch Hashverfahren mit asymmetrischen Schlüsseln. Diese werden jedoch der Rubrik der Signaturverfahren zugeordnet, da mit ihnen ausschließlich digitale Unterschriften gebildet werden. Abbildung 1.3 zeigt, wie sich Nachrichten mit einem MAC authentisieren lassen. Man beachte, dass nun auch der Hashwert über den unsicheren Kanal gesendet wird.

Möchte Alice eine Nachricht x an Bob übermitteln, so berechnet sie den zugehörigen MAC-Wert $y = h_k(x)$ und fügt diesen der Nachricht x hinzu. Bob überprüft die Echtheit der empfangenen Nachricht (x', y') , indem er seinerseits den zu x' gehörigen Hashwert $h_k(x')$ berechnet und das Ergebnis mit y' vergleicht. Der geheime Authentifikationsschlüssel k muss hierbei genau wie bei einem symmetrischen Kryptosystem über einen gesicherten Kanal vereinbart werden.

Indem Alice ihre Nachricht x um den Hashwert $y = h_k(x)$ ergänzt, gibt sie Bob nicht nur die Möglichkeit, anhand von y die empfangene Nachricht auf Manipulationen zu überprüfen. Die Benutzung des geheimen Schlüssels k erlaubt zudem eine Überprüfung der Herkunft der Nachricht.

1.2 Schlüssellose Hashfunktionen (MDCs)

In diesem Abschnitt betrachten wir verschiedene Sicherheitsanforderungen an einzelne Hashfunktionen h . Dabei nehmen wir an, dass h öffentlich bekannt ist, d.h. h ist eine schlüssellose Hashfunktion (MDC).

Sei $h: X \rightarrow Y$ eine Hashfunktion. Ein Paar $(x, y) \in X \times Y$ heißt **gültig** für h , falls $h(x) = y$ ist. Ein Paar (x, x') mit $h(x) = h(x')$ heißt **Kollisionspaar** für h . Die Anzahl $\|Y\|$ der Hashwerte bezeichnen wir mit m . Ist auch der Textraum X endlich, $\|X\| = n$, so heißt h eine (n, m) -**Hashfunktion**. In diesem Fall verlangen wir meist, dass $n \geq 2m$ ist, und wir nennen h dann eine **Kompressionsfunktion** (*compression function*).

Da h öffentlich bekannt ist, ist es sehr einfach, für einen vorgegebenen Text x ein gültiges Paar (x, y) zu erzeugen. Für bestimmte kryptografische Anwendungen ist es wichtig, dass dies nicht möglich ist, falls der Hashwert y vorgegeben wird.

Problem P1: Bestimmung eines Urbilds

Gegeben: Eine Hashfkt. $h: X \rightarrow Y$ und ein Hashwert $y \in Y$.

Gesucht: Ein Text $x \in X$ mit $h(x) = y$.

Falls es einen immensen Aufwand erfordert, für einen *vorgegebenen* Hashwert y einen Text x mit $h(x) = y$ zu finden, so heißt h **Einweg-Hashfunktion** (*one-way hash function* bzw. *preimage resistant hash function*). Diese Eigenschaft wird beispielsweise benötigt, wenn die Hashwerte der Benutzerpasswörter in einer öffentlich zugänglichen Datei abgespeichert werden, wie es bei manchen Unix-Systemen der Fall ist.

Problem P2: Bestimmung eines zweiten Urbilds

Gegeben: Eine Hashfkt. $h: X \rightarrow Y$ und ein Text $x \in X$.

Gesucht: Ein Text $x' \in X \setminus \{x\}$ mit $h(x') = h(x)$.

Falls sich für einen *vorgegebenen* Text x nur mit großem Aufwand ein weiterer Text $x' \neq x$ mit dem gleichen Hashwert $h(x') = h(x)$ finden lässt, heißt h **schwach kollisionsresistent** (*weakly collision resistant* bzw. *second preimage resistant*). Diese Eigenschaft wird in der durch Abbildung 1.2 skizzierten Anwendung benötigt. Beim Versuch, eine digitale Signatur zu fälschen (siehe unten), sieht sich der Angreifer dagegen mit folgender Problemstellung konfrontiert.

Problem P3: Bestimmung einer Kollision

Gegeben: Eine Hashfkt. $h: X \rightarrow Y$.

Gesucht: Texte $x \neq x' \in X$ mit $h(x') = h(x)$.

Falls sich dieses Problem nur mit einem immensen Aufwand lösen lässt, heißt h (**stark**) **kollisionsresistent** (*collision resistant*).

Obwohl die schwache Kollisionsresistenz eine gewisse Ähnlichkeit mit der Einweg-Eigenschaft aufweist, sind die beiden Eigenschaften im allgemeinen unvergleichbar. So muss eine schwach kollisionsresistente Funktion nicht notwendigerweise eine Einwegfunktion sein, da die Bestimmung eines Urbildes gerade für diejenigen Funktionswerte einfach sein kann, die nur ein einziges Urbild besitzen. Umgekehrt impliziert die Einweg-Eigenschaft auch nicht die schwache Kollisionsresistenz, da die Kenntnis eines Urbildes das Auffinden weiterer Urbilder sehr stark erleichtern kann.

1.2.1 Vergleich von Sicherheitsanforderungen

In diesem Abschnitt zeigen wir, dass stark kollisionsresistente Hashfunktionen sowohl schwach kollisionsresistent als auch Einweghashfunktionen sind.

Satz 1. *Sei $h: X \rightarrow Y$ eine (n, m) -Hashfunktion. Dann ist das Problem P3, ein Kollisionspaar für h zu bestimmen, auf das Problem P2, ein zweites Urbild zu bestimmen,*

```

1 wähle zufällig  $x \in X$ 
2  $x' := A(x)$ 
3 if  $x' \neq ?$  then return( $x, x'$ ) else return(?)

```

Abbildung 1.4: Reduktion des Kollisionsproblems auf das Problem, ein zweites Urbild zu bestimmen

reduzierbar. Folglich sind stark kollisionsresistente Hashfunktionen auch schwach kollisionsresistent.

Beweis. Sei A ein Las-Vegas Algorithmus, der für ein zufällig aus X gewähltes x mit Erfolgswahrscheinlichkeit ε ein zweites Urbild x' für h liefert und andernfalls $?$ ausgibt. Dann ist klar, dass der in Abbildung 1.4 dargestellte Las-Vegas Algorithmus mit Wahrscheinlichkeit ε ein Kollisionspaar findet. $\square$

Als nächstes zeigen wir, wie sich das Kollisionsproblem auf das Urbildproblem reduzieren lässt.

Satz 2. Sei $h: X \rightarrow Y$ eine (n, m) -Hashfunktion mit $n \geq 2m$. Dann ist das Problem P3, ein Kollisionspaar für h zu bestimmen, auf das Problem P1, ein Urbild zu bestimmen, reduzierbar.

Beweis. Sei A ein Invertierungsalgorithmus für h , d.h. A berechnet für jeden Hashwert y in $W(h) = \{h(x) \mid x \in X\}$ ein Urbild x mit $h(x) = y$. Betrachte den in Abbildung 1.5 dargestellten Las-Vegas Algorithmus B.

Sei $\mathcal{C} = \{h^{-1}(y) \mid y \in Y\}$. Dann hat B eine Erfolgswahrscheinlichkeit von

$$\sum_{C \in \mathcal{C}} \frac{\|C\|}{\|X\|} \cdot \frac{\|C\| - 1}{\|C\|} = \frac{1}{n} \sum_{C \in \mathcal{C}} (\|C\| - 1) = (n - m)/n \geq \frac{1}{2}.$$

$\square$

1.2.2 Das Zufallsorakelmodell (ZOM)

Das ZOM dient dazu, den Aufwand verschiedener Angriffe auf eine Hashfunktion $h: X \rightarrow Y$ nach oben abzuschätzen. Sind X und Y vorgegeben, so können wir eine Hashfunktion $h: X \rightarrow Y$ dadurch „konstruieren“, dass wir für jedes $x \in X$ zufällig ein $y \in Y$ wählen und $h(x)$ auf y setzen. Äquivalent hierzu ist, für h eine zufällige Funktion aus der Klasse $F(X, Y)$ aller m^n Funktionen von X nach Y zu wählen. Dieses Verfahren ist auf Grund des hohen Aufwands zwar nicht mehr praktikabel, wenn $n = \|X\|$ eine bestimmte Größe übersteigt. Es liefert uns aber ein theoretisches Modell für eine Hashfunktion mit „idealen“ kryptografischen Eigenschaften. Offensichtlich besteht für den Angreifer die einzige Möglichkeit, Informationen über h zu erhalten, darin, sich für eine Reihe von

```

1 wähle zufällig  $x \in X$ 
2  $y := h(x)$ 
3  $x' := A(y)$ 
4 if  $x \neq x'$  then return( $x, x'$ ) else return(?)

```

Abbildung 1.5: Reduktion des Kollisionsproblems auf das Urbildproblem

Prozedur FindPreimage(h, y, q)

```

1 wähle eine beliebige Menge  $X_0 = \{x_1, \dots, x_q\} \subseteq X$ 
2 for each  $x_i \in X_0$  do
3   if  $h(x_i) = y$  then return( $x_i$ )
4 return(?)

```

Abbildung 1.6: Bestimmung eines Urbilds für einen Hashwert

Texten die zugehörigen Hashwerte zu besorgen (was der Befragung eines funktionalen Zufallsorakels entspricht).

Eine Zufallsfunktion h eignet sich deshalb gut als kryptografische Hashfunktion, weil der Hashwert $h(x)$ für einen Text x auch dann noch schwer vorhersagbar ist, wenn der Angreifer bereits die Hashwerte einer beliebigen Zahl von anderen Texten kennt.

Proposition 3. Sei $X_0 = \{x_1, \dots, x_k\}$ eine beliebige Menge von k verschiedenen Texten aus X und seien $y_1, \dots, y_k \in Y$. Dann gilt für eine zufällig aus $F(X, Y)$ gewählte Funktion h und für jedes Paar $(x, y) \in (X - X_0) \times Y$,

$$\Pr[h(x) = y \mid h(x_i) = y_i \text{ für } i = 1, \dots, k] = 1/m.$$

Um eine obere Komplexitätsschranke für das Urbildproblem im ZOM zu erhalten, betrachten wir den in Abbildung 1.6 dargestellten Algorithmus. Hier (und bei den beiden folgenden Algorithmen) gibt der Parameter q die Anzahl der Hashwertberechnungen (also die Anzahl der gestellten Orakelfragen an das Zufallsorakel h) an. Der Zeitaufwand der Algorithmen ist dabei proportional zu q .

Satz 4. FINDPREIMAGE(h, y, q) gibt im ZOM mit Wahrscheinlichkeit $\varepsilon = 1 - (1 - 1/m)^q$ ein Urbild von y aus (unabhängig von der Wahl der Menge X_0).

Beweis. Sei $y \in Y$ fest und sei $X_0 = \{x_1, \dots, x_q\}$. Für $i = 1, \dots, q$ bezeichne E_i das Ereignis “ $h(x_i) = y$ ”. Nach Proposition 3 sind diese Ereignisse stochastisch unabhängig und ihre Wahrscheinlichkeit ist $\Pr[E_i] = 1/m$ ($i = 1, \dots, q$). Also folgt

$$\Pr[E_1 \cup \dots \cup E_q] = 1 - \Pr[\bar{E}_1 \cap \dots \cap \bar{E}_q] = 1 - (1 - 1/m)^q.$$

□

Der in Abbildung 1.7 dargestellte Algorithmus liefert uns eine obere Schranke für die Komplexität des Problems, ein zweites Urbild für $h(x)$ zu bestimmen. Die Erfolgswahrscheinlichkeit lässt sich vollkommen analog zum vorherigen Satz bestimmen.

Prozedur FindSecondPreimage(h, x, q)

```

1  $y := h(x)$ 
2 wähle eine beliebige Menge  $X_0 = \{x_1, \dots, x_{q-1}\} \subseteq X - \{x\}$ 
3 for each  $x_i \in X_0$  do
4   if  $h(x_i) = y$  then return( $x_i$ )
5 return(?)

```

Abbildung 1.7: Bestimmung eines 2. Urbilds für einen Hashwert

Prozedur Collision(h, q)

```

1 wähle eine beliebige Menge  $X_0 = \{x_1, \dots, x_q\} \subseteq X$ 
2 for each  $x_i \in X_0$  do  $y_i := h(x_i)$ 
3 if  $\exists i \neq j : y_i = y_j$  then return( $x_i, x_j$ ) else return(?)

```

Abbildung 1.8: Bestimmung eines Kollisionspaares

Satz 5. FINDSECONDPREIMAGE(h, x, q) gibt im ZOM mit Wahrscheinlichkeit $\varepsilon = 1 - (1 - 1/m)^{q-1}$ ein zweites Urbild $x_0 \neq x$ von $y = h(x)$ aus.

Ist q vergleichsweise klein, so ist bei beiden bisher betrachteten Angriffen $\varepsilon \approx q/m$. Um also auf eine Erfolgswahrscheinlichkeit von $1/2$ zu kommen, ist $q \approx m/2$ zu wählen.

Geht es lediglich darum, irgendein Kollisionspaar (x, x') aufzuspüren, so bietet sich ein sogenannter **Geburtstagsangriff** an. Dieser ist deutlich zeiteffizienter zu realisieren. Wie der Name schon andeutet, basiert dieser Angriff auf dem sogenannten Geburtstagsparadoxon, welches in seiner einfachsten Form folgendes besagt.

Geburtstagsparadoxon: Bereits in einer Schulklasse mit 23 Schulkindern haben mit einer Wahrscheinlichkeit größer $1/2$ mindestens zwei Kinder am gleichen Tag Geburtstag.*

Tatsächlich zeigt der nächste Satz, dass bei q -maligem Ziehen (mit Zurücklegen) aus einer Urne mit m Kugeln mit einer Wahrscheinlichkeit von

$$1 - (m-1)(m-2) \cdots (m-q+1)/m^{q-1}$$

mindestens eine Kugel mehr als einmal gezogen wird. Für $m = 365$ und $q = 23$ ergibt dies einen Wert von ungefähr 0,507.

Zur Kollisionsbestimmung verwenden wir den in Abbildung 1.8 dargestellten Algorithmus. Bei einer naiven Implementierung würde zwar der Zeitaufwand für die Auswertung der if-Bedingung quadratisch von q abhängen. Trägt man aber jeden Text x unter dem Suchwort $h(x)$ in eine (herkömmliche) Hashtabelle der Größe q ein, so wird der Zeitaufwand für die Bearbeitung jedes einzelnen Textes x im wesentlichen durch die Berechnung von $h(x)$ bestimmt.

Satz 6. COLLISION(h, q) gibt im ZOM mit Erfolgswahrscheinlichkeit

$$\varepsilon = 1 - \frac{(m-1)(m-2) \cdots (m-q+1)}{m^{q-1}}$$

ein Kollisionspaar (x, x') für h aus.

Beweis. Sei $X_0 = \{x_1, \dots, x_q\}$. Für $i = 1, \dots, q$ bezeichne E_i das Ereignis

$$“h(x_i) \notin \{h(x_1), \dots, h(x_{i-1})\}.”$$

Dann beschreibt $E_1 \cap \dots \cap E_q$ das Ereignis “COLLISION(h, q) gibt ? aus” und für $i = 1, \dots, q$ gilt

$$\Pr[E_i | E_1 \cap \dots \cap E_{i-1}] = \frac{m-i+1}{m}.$$

*Da die Häufigkeiten der Geburtstage in Wirklichkeit nicht ganz gleichmäßig über das Jahr verteilt sind, ist die Wahrscheinlichkeit sogar noch etwas höher.

Dies führt auf die Erfolgswahrscheinlichkeit

$$\begin{aligned}\varepsilon &= 1 - \Pr[E_1 \cap \dots \cap E_q] \\ &= 1 - \Pr[E_1] \Pr[E_2 | E_1] \dots \Pr[E_q | E_1 \cap \dots \cap E_{q-1}] \\ &= 1 - \left(\frac{m-1}{m}\right) \left(\frac{m-2}{m}\right) \dots \left(\frac{m-q+1}{m}\right).\end{aligned}$$

□

Mit $1 - x \approx e^{-x}$ folgt

$$\varepsilon = 1 - \prod_{i=1}^{q-1} \left(1 - \frac{i}{m}\right) \approx 1 - \prod_{i=1}^{q-1} e^{-\frac{i}{m}} = 1 - e^{-\frac{1}{m} \sum_{i=1}^{q-1} i} = 1 - e^{-\frac{q(q-1)}{2m}} \approx 1 - e^{-\frac{q^2}{2m}} \approx q^2/2m.$$

Somit erhalten wir die Abschätzung

$$q \approx c_\varepsilon \sqrt{m}$$

mit $c_\varepsilon = \sqrt{2\varepsilon}$. Diese Abschätzung ist nur für ε -Werte nahe Null hinreichend genau. Eine bessere Abschätzung ergibt sich aus der Approximation $\varepsilon \approx 1 - e^{-\frac{q^2}{2m}}$:

$$q \approx c'_\varepsilon \sqrt{m}$$

mit $c'_\varepsilon = \sqrt{2 \ln \frac{1}{1-\varepsilon}}$. Für $\varepsilon = 1/2$ ergibt sich somit $q \approx \sqrt{(2 \ln 2)m} \approx 1,17\sqrt{m}$.

Besitzt also eine binäre Hashfunktion $h: \{0, 1\}^n \rightarrow \{0, 1\}^m$ die Hashwertlänge $m = 128$ Bit, so müssen im ZOM $q \approx 1,17 \cdot 2^{64}$ Texte gehasht werden, um mit einer Wahrscheinlichkeit von $1/2$ eine Kollision zu finden. Um einem Geburtstagsangriff widerstehen zu können, sollte eine Hashfunktion mindestens eine Hashwertlänge von 128 oder besser 160 Bit haben.

1.2.3 Iterierte Hashfunktionen

In diesem Abschnitt beschäftigen wir uns mit der Frage, wie sich aus einer kollisionsresistenten Kompressionsfunktion

$$h: \{0, 1\}^{m+t} \rightarrow \{0, 1\}^m$$

eine kollisionsresistente Hashfunktion

$$\hat{h}: \{0, 1\}^* \rightarrow \{0, 1\}^l$$

konstruieren lässt. Hierzu betrachten wir folgende kanonische Konstruktionsmethode.

Preprocessing: Transformiere $x \in \{0, 1\}^*$ mittels einer Funktion

$$y: \{0, 1\}^* \rightarrow \bigcup_{r \geq 1} \{0, 1\}^{rt}$$

zu einem String $y(x)$ mit der Eigenschaft $|y(x)| \equiv_t 0$.

Processing: Sei $IV \in \{0, 1\}^m$ ein öffentlich bekannter Initialisierungsvektor und sei $y(x) = y_1 \dots y_r$ mit $|y_i| = t$ für $i = 1, \dots, r$. Berechne eine Folge $z_0, \dots, z_r$ von Strings $z_i \in \{0, 1\}^m$ wie folgt:

$$z_i = \begin{cases} IV, & i = 0, \\ h(z_{i-1}y_i), & i = 1, \dots, r. \end{cases}$$

Optionale Ausgabetransformation: Berechne den Hashwert $\hat{h}(x) = g(z_r)$, wobei $g: \{0,1\}^m \rightarrow \{0,1\}^l$ eine öffentlich bekannte Funktion ist. (Meist wird für g die Identität verwendet.)

Um $\hat{h}(x)$ zu berechnen, muss also die Kompressionsfunktion h genau r -mal aufgerufen werden. Wir formulieren nun eine für Preprocessing-Funktionen wünschenswerte Eigenschaft.

Definition 7. Eine Funktion $y: \{0,1\}^* \rightarrow \{0,1\}^*$ heißt **suffixfrei**, falls es keine Strings $x \neq \tilde{x}$ und z in $\{0,1\}^*$ mit $y(\tilde{x}) = zy(x)$ gibt (d.h. kein Funktionswert $y(x)$ ist Suffix eines Funktionswertes $y(\tilde{x})$ an einer Stelle $\tilde{x} \neq x$).

Man beachte, dass jede suffixfreie Funktion insbesondere injektiv ist.

Satz 8. Falls die Preprocessing-Funktion y suffixfrei und die Ausgabetransformation g injektiv ist, so ist mit h auch $\hat{h}$ kollisionsresistent.

Beweis. Angenommen, es gelingt, ein Kollisionspaar $x, \tilde{x}$ für $\hat{h}$ mit $\hat{h}(x) = \hat{h}(\tilde{x})$ zu finden. Sei

$$y(x) = y_1 y_2 \dots y_{k-1} y_k \text{ und } y(\tilde{x}) = \tilde{y}_1 \tilde{y}_2 \dots \tilde{y}_{l-1} \tilde{y}_l \text{ mit } k \leq l.$$

Da y suffixfrei ist, muss ein Index $i \in \{1, \dots, k\}$ mit $y_i \neq \tilde{y}_{l-k+i}$ existieren. Weiter seien z_i ($i = 0, \dots, k$) und $\tilde{z}_j$ ($j = 0, \dots, l$) die in der Processing-Phase berechneten Hashwerte. Da g injektiv ist, muss mit $g(z_k) = \hat{h}(x) = \hat{h}(\tilde{x}) = g(\tilde{z}_l)$ auch $z_k = \tilde{z}_l$ gelten. Sei $i_{\max}$ der größte Index $i \in \{1, \dots, k\}$ mit $z_{i-1} y_i \neq \tilde{z}_{l-k+i-1} \tilde{y}_{l-k+i}$. Dann bilden $z_{i_{\max}-1} y_{i_{\max}}$ und $\tilde{z}_{l-k+i_{\max}-1} \tilde{y}_{l-k+i_{\max}}$ wegen

$$h(z_{i_{\max}-1} y_{i_{\max}}) = z_{i_{\max}} = \tilde{z}_{l-k+i_{\max}} = h(\tilde{z}_{l-k+i_{\max}-1} \tilde{y}_{l-k+i_{\max}})$$

ein Kollisionspaar für h . □

1.2.4 Die Merkle-Damgaard-Konstruktion

Merkle und Damgaard schlugen 1989 folgende konkrete Realisierung ihrer Konstruktion vor. Als Initialisierungsvektors wird der Nullvektor $IV = 0^m$ benutzt, die optionale Ausgabetransformation entfällt, und für $y(x)$ wird im Fall $t \geq 2$ die folgende Funktion verwendet. (Den Fall $t = 1$ betrachten wir später.)

Für $x = \varepsilon$ sei $y(x) = 0^t$ und für $x \in \{0,1\}^n$ mit $n > 0$ sei $k = \lceil \frac{n}{t-1} \rceil$ und $x = x_1 x_2 \dots x_{k-1} x_k$ mit $|x_1| = |x_2| = \dots = |x_{k-1}| = t-1$ sowie $|x_k| = t-1-d$, wobei $0 \leq d < t-1$. Im Fall $k = 1$ ist dann $y(x) = 0x0^d \text{bin}_{t-1}(d)$ und für $k > 1$ ist $y(x) = y_1 \dots y_{k+1}$, wobei

$$y_i = \begin{cases} 0x_1, & i = 1, \\ 1x_i, & 2 \leq i < k, \\ 1x_k 0^d, & i = k, \\ 1\text{bin}_{t-1}(d), & i = k+1, \end{cases} \quad (1.1)$$

und $\text{bin}_{t-1}(d)$ die durch führende Nullen auf die Länge $t-1$ aufgefüllte Binärdarstellung von d ist.

Satz 9. Die durch (1.1) definierte Preprocessing-Funktion y ist suffixfrei.

Beweis. Seien $x \neq \tilde{x}$ zwei Texte mit $|x| \leq |\tilde{x}|$. Wir müssen zeigen, dass $y(x) = y_1 y_2 \dots y_{k+1}$ kein Suffix von $y(\tilde{x}) = \tilde{y}_1 \tilde{y}_2 \dots \tilde{y}_{l+1}$ ist. Im Fall $x = \varepsilon$ ist dies klar. Für $x \neq \varepsilon$ machen wir folgende Fallunterscheidung.

- 1. Fall:** $|x| \not\equiv_{t-1} |\tilde{x}|$. Dann folgt $d \neq \tilde{d}$ und somit $y_{k+1} \neq \tilde{y}_{l+1}$.
- 2. Fall:** $|x| = |\tilde{x}|$. In diesem Fall ist $k = l$. Wegen $x \neq \tilde{x}$ existiert ein Index $i \in \{1, \dots, k\}$ mit $x_i \neq \tilde{x}_i$. Dies impliziert $y_i \neq \tilde{y}_i$, also ist $y(x)$ kein Suffix von $y(\tilde{x})$.
- 3. Fall:** $|x| \neq |\tilde{x}|$ und $|x| \equiv_{t-1} |\tilde{x}|$. In diesem Fall ist $k < l$. Da $y(x)$ mit einer Null beginnt, aber das $(l - k + 1)$ -te Bit von $y(\tilde{x})$ eine Eins ist, kann $y(x)$ kein Suffix von $y(\tilde{x})$ sein. $\square$

Nun kommen wir zum Fall $t = 1$. Sei y die durch $y(x) := 11f(x)$ definierte Funktion, wobei f wie folgt definiert ist:

$$f(x_1, \dots, x_n) = f(x_1) \dots f(x_n) \text{ mit } f(0) = 0 \text{ und } f(1) = 01.$$

Dann ist leicht zu sehen, dass y suffixfrei ist. Da die Kompressionsfunktion h bei der Berechnung von $\hat{h}(x)$ im Fall $t = 1$ für jedes Bit von $y(x)$ einmal aufgerufen wird, wird h genau $|y(x)| \leq 2(n+1)$ -mal aufgerufen. Im Fall $t > 1$ werden dagegen nur $k+1 = \lceil \frac{n}{t-1} \rceil + 1$ Aufrufe benötigt.

1.2.5 Die MD4-Hashfunktion

Die MD4-Hashfunktion (*Message Digest*) wurde 1990 von Rivest vorgeschlagen. Die Bitlänge von MD4 beträgt $l = 128$ Bit. Bei einer Wortlänge von 32 Bit entspricht dies 4 Wörtern. Die im Folgenden vorgestellten Hashfunktionen benutzen u.a. folgende Operationen auf Wörtern.

Operatoren auf $\{0, 1\}^{32}$	
$X \wedge Y$	bitweises „Und“ von X und Y
$X \vee Y$	bitweises „Oder“ von X und Y
$X \oplus Y$	bitweises „exklusives Oder“ von X und Y
$\neg X$	bitweises Komplement von X
$X + Y$	Ganzzahl-Addition modulo 2^{32}
$X \rightarrow s$	Rechtsshift um s Stellen
$X \leftarrow s$	zirkulärer Linksshift um s Stellen

Während die Ganzzahl-Addition bei MD4 und MD5 in *little endian* Architektur (d.h. ein aus 4 Bytes $a_3 a_2 a_1 a_0$, $0 \leq a_i \leq 255$ zusammengesetztes Wort repräsentiert die Zahl $a_0 2^{24} + a_1 2^{16} + a_2 2^8 + a_3$) ausgeführt wird, verwendet **SHA-1** eine *big endian* Architektur (d.h. $a_3 a_2 a_1 a_0$, $0 \leq a_i \leq 255$ repräsentiert die Zahl $a_3 2^{24} + a_2 2^{16} + a_1 2^8 + a_0$). Der MD4-Algorithmus benutzt die folgenden Konstanten y_j, z_j, s_j , $j = 0, \dots, 47$

	y_j (in Hexadezimaldarstellung)
$j = 0, \dots, 15$	0
$j = 16, \dots, 31$	5a827999
$j = 32, \dots, 47$	6ed9eba1

	z_j
$j = 0, \dots, 15$	0, 1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 13, 14, 15
$j = 16, \dots, 31$	0, 4, 8, 12, 1, 5, 9, 13, 2, 6, 10, 14, 3, 7, 11, 15
$j = 32, \dots, 47$	0, 8, 4, 12, 2, 10, 6, 14, 1, 9, 5, 13, 3, 11, 7, 15
	s_j
$j = 0, \dots, 15$	3, 7, 11, 19, 3, 7, 11, 19, 3, 7, 11, 19, 3, 7, 11, 19
$j = 16, \dots, 31$	3, 5, 9, 13, 3, 5, 9, 13, 3, 5, 9, 13, 3, 5, 9, 13
$j = 32, \dots, 47$	3, 9, 11, 15, 3, 9, 11, 15, 3, 9, 11, 15, 3, 9, 11, 15

und folgende Funktionen f_j , $j = 0, \dots, 47$

$$f_j(X, Y, Z) := \begin{cases} (X \wedge Y) \vee (\neg X \wedge Z), & j = 0, \dots, 15, \\ (X \wedge Y) \vee (X \wedge Z) \vee (Y \wedge Z), & j = 16, \dots, 31, \\ X \oplus Y \oplus Z, & j = 32, \dots, 47. \end{cases}$$

Für MD4 konnten nach ca. 2^{20} Hashwertberechnungen Kollisionen aufgespürt werden. Deshalb gilt MD4 nicht mehr als kollisionsresistent.

MD4(x)

```

1 input  $x \in \{0, 1\}^*, |x| = n$ 
2  $y := x10^k \mathbf{bin}_{64}(n)$ ,  $k \in \{0, 1, \dots, 511\}$  mit  $n + 1 + k + 64 \equiv 0 \pmod{512}$ 
3  $(H_1, H_2, H_3, H_4) := (67452301, efcdab89, 98badcfe, 10325476)$ 
4 sei  $y = M_1 \cdots M_r$ ,  $r = (n + 1 + k + 64)/512$ 
5 for  $i := 1$  to  $r$  do
6   sei  $M_i = X[0] \cdots X[15]$ 
7    $(A, B, C, D) := (H_1, H_2, H_3, H_4)$ 
8   for  $j := 0$  to 47 do
9      $(A, B, C, D) := (D, (A + f_j(B, C, D) + X[z_j] + y_j) \leftarrow s_j, B, C)$ 
10     $(H_1, H_2, H_3, H_4) := (H_1 + A, H_2 + B, H_3 + C, H_4 + D)$ 
11 output  $H_1 H_2 H_3 H_4$ 

```

1.2.6 Die MD5-Hashfunktion

Der MD5 ist eine 1991 von Rivest präsentierte verbesserte Version von MD4. Die Bitlänge von MD5 beträgt wie bei MD4 $l = 128$ Bit. Bei einer Wortlänge von 32 Bit entspricht dies 4 Wörtern. In MD5 werden teilweise andere Konstanten als in MD4 verwendet. Zudem besitzt MD5 eine zusätzliche 4. Runde ($j = 48, \dots, 63$), in der die Funktion $f_j(X, Y, Z) = Y \oplus (X \vee \neg Z)$ verwendet wird. Außerdem wurde die in Runde 2 von MD4 verwendete Funktion durch $f_j(X, Y, Z) := (X \wedge Z) \vee (Y \wedge \neg Z)$, $j = 16 \dots 31$, ersetzt. Die y -Konstanten sind definiert als $y_j :=$ die ersten 32 Bit der Binärdarstellung von $\text{abs}(\sin(j + 1))$, $0 \leq j \leq 63$, und für z_j und s_j werden folgende Konstanten benutzt.

	z_j
$j = 0, \dots, 15$	$z_j = j : 0, 1, 2, 3, 4, 5, 6, 7, 8, 9, 10, 11, 12, 13, 14, 15$
$j = 16, \dots, 31$	$z_j = (5j + 1) \bmod 16 : 1, 6, 11, 0, 5, 10, 15, 4, 9, 14, 3, 8, 13, 2, 7, 12$
$j = 32, \dots, 47$	$z_j = (3j + 5) \bmod 16 : 5, 8, 11, 14, 1, 4, 7, 10, 13, 0, 3, 6, 9, 12, 15, 2$
$j = 48, \dots, 63$	$z_j = 7j \bmod 16 : 0, 7, 14, 5, 12, 3, 10, 1, 8, 15, 6, 13, 4, 11, 2, 9$
	s_j
$j = 0, \dots, 15$	7, 12, 17, 22, 7, 12, 17, 22, 7, 12, 17, 22, 7, 12, 17, 22
$j = 16, \dots, 31$	5, 9, 14, 20, 5, 9, 14, 20, 5, 9, 14, 20, 5, 9, 14, 20
$j = 32, \dots, 47$	4, 11, 16, 23, 4, 11, 16, 23, 4, 11, 16, 23, 4, 11, 16, 23
$j = 48, \dots, 63$	6, 10, 15, 21, 6, 10, 15, 21, 6, 10, 15, 21, 6, 10, 15, 21

Für MD5 konnten in 2004 ebenfalls Kollisionspaare gefunden werden (für die Kompressionsfunktion von MD5 gelang dies bereits 1996).

MD5(x)

```

1  input  $x \in \{0, 1\}^*, |x| = n$ 
2   $y := x10^k \mathbf{bin}_{64}(n)$ ,  $k \in \{0, 1, \dots, 511\}$  mit  $n + 1 + k + 64 \equiv 0 \pmod{512}$ 
3   $(H_1, H_2, H_3, H_4) := (67452301, efcdab89, 98badcfe, 10325476)$ 
4  sei  $y = M_1 \cdots M_r$ ,  $r = (n + 1 + k + 64)/512$ 
5  for  $i := 1$  to  $r$  do
6    sei  $M_i = X[0] \cdots X[15]$ 
7     $(A, B, C, D) := (H_1, H_2, H_3, H_4)$ 
8    for  $j := 0$  to 63 do
9       $(A, B, C, D) := (D, B + (A + f_j(B, C, D) + X[z_j] + y_j) \leftarrow s_j, B, C)$ 
10      $(H_1, H_2, H_3, H_4) := (H_1 + A, H_2 + B, H_3 + C, H_4 + D)$ 
11 output  $H_1 H_2 H_3 H_4$ 

```

1.2.7 Die SHA-1-Hashfunktion

Der *Secure Hash Algorithm* (**SHA-1**) ist eine Weiterentwicklung des MD4 bzw. MD5 Algorithmus. Er gilt in den USA als Standard und ist Bestandteil des von der US-Behörde NIST (National Institute of Standards and Technology) im August 1991 veröffentlichten DSS (Digital Signature Standard). Die Bitlänge von **SHA-1** beträgt $l = 160$ Bit. Bei einer Wortlänge von 32 Bit entspricht dies 5 Wörtern. **SHA-1** unterscheidet sich nur geringfügig von der SHA-0 Hashfunktion, in der eine Schwachstelle dazu führt, dass nach Berechnung von ca. 2^{61} Hashwerten ein Kollisionspaar gefunden werden kann (obwohl bei einem Geburtstagsangriff auf Grund der Hashwertlänge von 160 Bit ca. 2^{80} Berechnungen erforderlich sein müssten). Diese potentielle Schwäche von SHA-0 wurde im **SHA-1** dadurch entfernt, dass **SHA-1** in Zeile 8 einen zirkulären Shift um eine Bitstelle ausführt. Der **SHA-1**-Algorithmus benutzt die folgenden Konstanten K_j , $j = 0, \dots, 79$

	K_j (in Hexadezimaldarstellung)
$j = 0, \dots, 19$	5a827999
$j = 20, \dots, 39$	6ed9eba1
$j = 40, \dots, 59$	8f1bbcdc
$j = 60, \dots, 79$	ca62c1d6

und folgende Funktionen f_j , $j = 0, \dots, 79$

$$f_j(X, Y, Z) := \begin{cases} (X \wedge Y) \vee (\neg X \wedge Z), & j = 0, \dots, 19, \\ X \oplus Y \oplus Z, & j = 20, \dots, 39, \\ (X \wedge Y) \vee (X \wedge Z) \vee (Y \wedge Z), & j = 40, \dots, 59, \\ X \oplus Y \oplus Z, & j = 60, \dots, 79. \end{cases}$$

SHA-1(x)

```

1 input  $x \in \{0, 1\}^*$ ,  $|x| = n$ 
2  $y := x10^k \mathit{bin}_{64}(n)$ ,  $k \in \{0, 1, \dots, 511\}$  mit  $n + 1 + k + 64 \equiv 0 \pmod{512}$ 
3  $(H_0, H_1, H_2, H_3, H_4) := (67452301, efcdab89, 98badcfe, 10325476, c3d2e1f0)$ 
4 sei  $y = M_1 \cdots M_r$ ,  $r = (n + 1 + k + 64)/512$ 
5 for  $i := 1$  to  $r$  do
6   sei  $M_i = X[0] \cdots X[15]$ 
7   for  $t := 16$  to  $79$  do
8      $X[t] := (X[t - 3] \oplus X[t - 8] \oplus X[t - 14] \oplus X[t - 16]) \leftarrow 1$ 
9      $(A, B, C, D, E) := (H_0, H_1, H_2, H_3, H_4)$ 
10    for  $j := 0$  to  $79$  do
11       $temp := (A \leftarrow 5) + f_j(B, C, D) + E + X[j] + K_j$ 
12       $(A, B, C, D, E) := (temp, A, B \leftarrow 30, C, D)$ 
13       $(H_0, H_1, H_2, H_3, H_4) := (H_0 + A, H_1 + B, H_2 + C, H_3 + D, H_4 + E)$ 
14 output  $H_0 H_1 H_2 H_3 H_4$ 

```

1.2.8 Die SHA-2-Familie

Im Jahr 2001 veröffentlichte die US-Behörde NIST drei weitere Hashfunktionen der SHA-Familie: SHA-256, SHA-384, and SHA-512. Diese Funktionen werden auch als SHA-2 Hashfunktionen bezeichnet. In 2004 kam noch SHA-224 als vierte Variante hinzu. SHA-256 und SHA-512 haben denselben Aufbau, unterscheiden sich aber in erster Linie in der benutzten Wortlänge: 32 Bit bei SHA-256 und 64 Bit bei SHA-512. Zudem werden unterschiedliche Shift- und Summationskonstanten verwendet und auch die Rundenzahlen differieren. SHA-224 und SHA-384 sind reduzierte Varianten von SHA-256 und SHA-512. Der SHA-256-Algorithmus benutzt die folgenden Konstanten K_j , $j = 0, \dots, 63$ (in Hexadezimaldarstellung).

428a2f98, 71374491, b5c0fbcf, e9b5dba5, 3956c25b, 59f111f1, 923f82a4, ab1c5ed5,
d807aa98, 12835b01, 243185be, 550c7dc3, 72be5d74, 80deb1fe, 9bdc06a7, c19bf174,
e49b69c1, efbe4786, 0fc19dc6, 240ca1cc, 2de92c6f, 4a7484aa, 5cb0a9dc, 76f988da,
983e5152, a831c66d, b00327c8, bf597fc7, c6e00bf3, d5a79147, 06ca6351, 14292967,
27b70a85, 2e1b2138, 4d2c6dfc, 53380d13, 650a7354, 766a0abb, 81c2c92e, 92722c85,
a2bfe8a1, a81a664b, c24b8b70, c76c51a3, d192e819, d6990624, f40e3585, 106aa070,
19a4c116, 1e376c08, 2748774c, 34b0bcb5, 391c0cb3, 4ed8aa4a, 5b9cca4f, 682e6ff3,
748f82ee, 78a5636f, 84c87814, 8cc70208, 90bffffa, a4506ceb, bef9a3f7, c67178f2

Dies sind jeweils die ersten 32 Bit der binären Nachkommastellen der dritten Wurzeln der ersten 64 Primzahlen $2, \dots, 311$. SHA-256 arbeitet wie folgt.

SHA-256(x)

```

1  input  $x \in \{0, 1\}^*, |x| = n$ 
2   $y := x10^k \mathit{bin}_{64}(n)$ ,  $k \in \{0, 1, \dots, 511\}$  mit  $n + 1 + k + 64 \equiv 0 \pmod{512}$ 
3   $(H_0, H_1, H_2, H_3, H_4, H_5, H_6, H_7) := (6a09e667, bb67ae85, 3c6ef372, a54ff53a,$ 
4     $510e527f, 9b05688c, 1f83d9ab, 5be0cd19)$ 
5  sei  $y = M_1 \cdots M_r$ ,  $r = (n + 1 + k + 64)/512$ 
6  for  $i := 1$  to  $r$  do
7    sei  $M_i = X[0] \cdots X[15]$ 
8    for  $t := 16$  to  $63$  do
9       $s0 := (X[t - 15] \hookrightarrow 7) \oplus (X[t - 15] \hookrightarrow 18) \oplus (X[t - 15] \rightarrow 3)$ 
10      $s1 := (X[t - 2] \hookrightarrow 17) \oplus (X[t - 2] \hookrightarrow 19) \oplus (X[t - 2] \rightarrow 10)$ 
11      $X[t] := X[t - 16] + s0 + X[t - 7] + s1$ 
12      $(A, B, C, D, E, F, G, H) := (H_0, H_1, H_2, H_3, H_4, H_5, H_6, H_7)$ 
13     for  $j := 0$  to  $63$  do
14        $s0 := (A \hookrightarrow 2) \oplus (A \hookrightarrow 13) \oplus (A \hookrightarrow 22)$ 
15        $maj := (A \wedge B) \oplus (A \wedge C) \oplus (B \wedge C)$ 
16        $t2 := s0 + maj$ 
17        $s1 := (E \hookrightarrow 6) \oplus (E \hookrightarrow 11) \oplus (E \hookrightarrow 25)$ 
18        $ch := (E \wedge F) \oplus (\neg E \wedge G)$ 
19        $t1 := H + s1 + ch + K_j + X[j]$ 
20        $(A, B, C, D, E, F, G, H) := (t1 + t2, A, B, C, D + t1, E, F, G)$ 
21        $(H_0, H_1, H_2, H_3, H_4, H_5, H_6, H_7)$ 
22        $:= (H_0 + A, H_1 + B, H_2 + C, H_3 + D, H_4 + E, H_5 + F, H_6 + G, H_7 + H)$ 
23  output  $H_0 H_1 H_2 H_3 H_4 H_5 H_6 H_7$ 

```

Die Initialwerte von $H_0, \dots, H_7$ in den Zeilen 3 und 4 sind jeweils die ersten 32 Bit der binären Nachkommastellen der Wurzeln der Primzahlen 2, 3, 5, 7, 11, 13, 17, 19.

1.2.9 Kryptoanalyse von Hashfunktionen

Bereits 1991 wurden von Den Boer und Bosselaers Schwächen im MD4 aufgedeckt. Im August 2004 erschien ein Bericht [1] mit einer Anleitung, wie sich Kollisionen für MD4 mittels “hand calculation” finden lassen.

In 1993, fanden den Boer und Bosselaers einen Weg, so genannte “Pseudo-Kollisionen” für die MD5 Kompressionsfunktion zu generieren. In 1996, fand Dobbertin ein Kollisionspaar für die MD5 Kompressionsfunktion.

Im August 2004 wurden schließlich Kollisionen für MD5 von Xiaoyun Wang, Dengguo Feng, Xuejia Lai und Hongbo Yu berechnet. Der benötigte Aufwand wurde mit ca. 1 Stunde auf einem IBM p690 Cluster abgeschätzt.

Im März 2005 veröffentlichten Arjen Lenstra, Xiaoyun Wang und Benne de Weger zwei X.509 Zertifikate mit unterschiedlichen Public-keys, die auf denselben MD5-Hashwert führten. Nur wenige Tage später beschrieb Vlastimil Klima eine Möglichkeit, Kollisionen für MD5 innerhalb weniger Stunden auf einem Notebook zu berechnen. Mittels der so genannten Tunneling-Methode wurde die Rechenzeit vom gleichen Autor im März 2006 auf eine Minute verkürzt.

Auf der CRYPTO 98 stellten Chabaud und Joux einen Angriff auf SHA-0 vor, der ein Kollisionspaar mit nur 2^{61} Hashwertberechnungen (anstelle von 2^{80} bei einem Geburtstagsangriff) aufspürt.

In 2004 fanden Biham und Chen Beinahe-Kollisionen für den SHA-0, bei denen sich die Hashwerte nur an 18 von den 160 Bitpositionen unterschieden. Zudem legten sie volle Kollisionen für den auf 62 Runden reduzierten SHA-0 Algorithmus vor.

Schließlich wurde im August 2004 die Berechnung einer Kollision für den vollen 80-Runden SHA-0 Algorithmus von Joux, Carribault, Lemuet und Jalby bekannt gegeben. Hierzu wurden lediglich 2^{51} Hashwerte berechnet, die ca. 80 000 Stunden CPU-Rechenzeit auf einem 2-Prozessor 256-Itanium Supercomputer benötigten.

Ebenfalls im August 2004 wurde von Wang, Feng, Lai und Yu auf der CRYPTO 2004 eine Angriffsmethode für MD5, SHA-0 und andere Hashfunktionen vorgestellt, mit der sich die Anzahl der Hashwertberechnungen auf 2^{40} senken lässt. Dies wurde im Februar 2005 von Xiaoyun Wang, Yiqun Lisa Yin und Hongbo Yu geringfügig auf 2^{39} Hashwertberechnungen verbessert.

Aufgrund der erfolgreichen Angriffe auf SHA-0 rieten mehrere Experten von einer weiteren Verwendung des **SHA-1** ab. Daraufhin kündigte die amerikanische Behörde NIST an, **SHA-1** in 2010 zugunsten der SHA-2 Varianten abzulösen.

Im Jahr 2005 veröffentlichten Rijmen und Oswald einen Angriff, der mit weniger als 2^{80} Hashwertberechnungen ein Kollisionspaar für den auf 53 Runden reduzierten **SHA-1** Algorithmus findet. Nur wenig später kündigten Xiaoyun Wang, Yiqun Lisa Yin und Hongbo Yu einen Angriff auf den vollen 80-Runden **SHA-1** mit 2^{69} Hashwertberechnungen an. Im August 2005 erfuhr der benötigte Aufwand von Xiaoyun Wang, Andrew Yao und Frances Yao auf der CRYPTO 2005 eine weitere Reduktion auf 2^{63} Berechnungen. In 2008 wurde von Stephane Manuel ein Kollisionsangriff mit einem geschätzten Aufwand von 2^{51} bis 2^{57} Berechnungen veröffentlicht.

Die besten bekannten Angriffe gegen SHA-2 brechen die von 64 auf 41 Runden reduzierte Variante von SHA-256 und die von 80 auf 46 Runden reduzierte Variante von SHA-512.

Im Oktober 2012 wurde der Hash-Algorithmus Keccak als Gewinner des vom NIST ausgeschrieben Wettbewerbs für den SHA-3-Algorithmus ausgewählt. Die Intention dabei war nicht, SHA-2 als Standard durch SHA-3 abzulösen, zumal bisher keine erfolgreichen Angriffe gegen SHA-2 bekannt sind. Vielmehr ging es bei diesem Wettbewerb darum, angesichts der erfolgreichen Angriffe gegen MD5 und SHA-0, die einen ähnlichen Aufbau wie SHA-1 und SHA-2 haben, eine auf einem vollkommen anderen Entwurfsprinzip basierende Alternative zur Verfügung zu stellen.

1.2.10 Die Sponge-Konstruktion

Die Konstruktionsidee hinter dem SHA-3-Gewinner Keccak wird von den Autoren als *Sponge* (Schwamm) bezeichnet. Auf der Basis dieser Entwurfsmethode lassen sich außer Hashfunktionen bspw. auch Pseudozufallsgeneratoren gewinnen. Der Aufbau eines Sponges ähnelt oberflächlich betrachtet der in 1.2.3 vorgestellten Konstruktion von iterierten Hashfunktionen, weist aber einige Unterschiede auf. So basiert ein Sponge statt auf einer Kompressionsfunktion h auf einer Permutation (oder allgemeiner Transformation) $f : \{0, 1\}^b \rightarrow \{0, 1\}^b$, die wie h iteriert angewendet wird. Dabei wird der aktuelle b -Bitblock in zwei Teilblöcke der Länge r und c unterteilt, die als äußerer bzw. innerer Zustand bezeichnet werden. Wie der Name schon sagt, verbleiben die Bits des inneren Zustands im Sponge, d.h. sie dienen nur zur Berechnung des nächsten Zustands und werden im Gegensatz zu den Bits des äußeren Zustands nicht unmittelbar für die Gewinnung der Ausgabe genutzt. Die Anzahl c der Bits des inneren Zustands wird als **Kapazität**

des Sponges bezeichnet und ist sein wichtigster Sicherheitsparameter. Die Anzahl r der Bits des äußeren Zustands heißt **Bitrate**, wobei $r + c = b$ gelten muss.

Bevor die Funktion f im Kern des Algorithmus iteriert angewendet wird, um eine Zustandsfolge zu generieren, wird ein Preprocessing ausgeführt. Die Anforderungen an diese Funktion definieren wir vorab.

Definition 10. Sei $r \geq 1$. Eine Funktion $y: \{0, 1\}^* \rightarrow \bigcup_{k \geq 1} \{0, 1\}^{kr}$ heißt **sponge-konforme Paddingfunktion** für Bitrate r , falls gilt:

- $\forall n \geq 0 \exists z \forall x \in \{0, 1\}^n : y(x) = xz$,
- $\forall k \geq 0 \forall x \neq x' : y(x) \neq y(x')0^{kr}$.

Es ist leicht zu sehen, dass die Funktion $\text{pad10}^*\mathbf{1}_r: \{0, 1\}^* \rightarrow \{0, 1\}^*$ definiert durch

$$\text{pad10}^*\mathbf{1}_r(x) = x10^d1 \text{ mit } d = \min \{i \geq 0 \mid |x| + 2 + i \equiv_r 0\}$$

sponge-konform für die Bitrate r ist. Tatsächlich ist $\text{pad10}^*\mathbf{1}_r$ sogar für jede Bitrate $r' \geq 1$ sponge-konform. Ohne die abschließende 1 wäre dies nicht der Fall.

Definition 11. Seien $r \geq 1$, y ein sponge-konformes Padding für r und $f: \{0, 1\}^b \rightarrow \{0, 1\}^b$. Die Funktion $\text{Sponge}_{f,y,r}: \mathbb{N} \times \{0, 1\}^* \rightarrow \{0, 1\}^*$ ist wie folgt definiert: Für $x \in \{0, 1\}^*$ sei $y_1 \dots y_k := y(x)$ mit $|y_i| = r$ ($1 \leq i \leq k$). Wir definieren die Zustände $s_i, i \geq 0$:

$$s_i = \begin{cases} 0^b & i = 0 \\ f(s_{i-1} \oplus (y_i 0^c)) & 1 \leq i \leq k \quad (\text{Absorptionsphase}) \\ f(s_{i-1}) & i > k \quad (\text{Squeezing-Phase}) \end{cases}.$$

Weiter bezeichne z_i die ersten r Bits von s_{k+i-1} , $i \geq 1$, es sei $m = \lfloor \frac{l}{r} \rfloor$ und z'_{m+1} bezeichne die ersten $l - mr$ Bits von z_{m+1} . Dann ist

$$\text{Sponge}_{f,y,r}(l, x) = z_1 \dots z_m z'_{m+1}.$$

Für die Analyse definieren wir

$$\text{Absorb}_{f,y,r}(x) = s_k \text{ und } \text{Squeeze}_{f,r}(l, s_k) = z_1 \dots z_m z'_{m+1}.$$

Den Aufwand, für festes l ein Kollisionspaar $x \neq x'$ mit $\text{Sponge}_{f,y,r}(l, x) = \text{Sponge}_{f,y,r}(l, x')$ zu finden, können wir nach oben durch den Aufwand abschätzen, ein Paar $x \neq x'$ zu finden, so dass $\text{Absorb}_{f,y,r}(x) = \text{Absorb}_{f,y,r}(x')$. Da in der Absorptionsphase der äußere Zustand (d.h. die Folge der ersten r Bits) beliebig und somit auch identisch gesetzt werden kann, genügt es, ein *inneres Kollisionspaar* zu finden, d.h. solche $x \neq x'$ so dass $\text{Absorb}_{f,y,r}^i(x) = \text{Absorb}_{f,y,r}^i(x')$, wobei $\text{Absorb}_{f,y,r}^i(x)$ die Folge der letzten c Bits von $\text{Absorb}_{f,y,r}(x)$ bezeichnet.

Um eine solche innere Kollision zu finden, hilft es, sich die 2^c inneren Zustände als Knoten eines gerichteten Multigraphen G vorzustellen, wobei jeder Knoten 2^r ausgehende Kanten mit Label 0^r bis 1^r hat. Ziel ist es dann, zwei verschiedene Pfade von 0^c zu demselben Knoten v zu finden, wobei zwei Pfade auch dann verschieden sind, wenn sich die Kanten nur in den Labeln unterscheiden. Anders als beim ZOM für eine Hashfunktion lohnt es sich hier für den Angreifer, die Argumente adaptiv nach einer Strategie $\mathcal{S}$ zu wählen. Der Algorithmus in Abb. 1.9 fasst dieses Vorgehen zusammen. Der Einfachheit halber gibt er ein Kollisionspaar nach dem Padding aus; für $\text{pad10}^*\mathbf{1}_r$ und alle y , deren Padding nur von $|x| \bmod r$ abhängt, lässt sich dieses aber leicht auf ein Paar vor dem Preprocessing erweitern.

Prozedur InnerCollision($f, r, q, \mathcal{S}$)

```

1   $c := b - r$ , wobei  $f : \{0, 1\}^b \rightarrow \{0, 1\}^b$ 
2  initialisiere den gerichteten Multigraphen  $G = (V, A) := (\{0, 1\}^c, \emptyset)$ 
3  for  $i := 1$  to  $q$  do
4    wähle  $v \in V$  und  $x \in \{0, 1\}^r$  nach Strategie  $\mathcal{S}$ 
5     $x'v' := f(xv)$ 
6     $A := A \cup \{(v, v', x, x')\}$ 
7  if  $\exists$  verschiedene Pfade  $(0^c, u_1, x_1, x'_1), \dots, (u_{k-1}, u_k, x_k, x'_k)$  und
8     $(0^c, v_1, y_1, y'_1), \dots, (v_{l-1}, v_l, y_l, y'_l)$  mit  $u_k = v_l$  in  $G$ 
9    return  $(x_1(x_2 \oplus x'_1) \dots (x_k \oplus x'_{k-1}), y_1(y_2 \oplus y'_1) \dots (y_l \oplus y'_{l-1}))$ 
10 else
11 return (?)

```

Abbildung 1.9: Bestimmung eines inneren Kollisionspaares

Satz 12. Für jede Strategie $\mathcal{S}$ gibt INNERCOLLISION($f, r, q, \mathcal{S}$) im ZOM mit Erfolgswahrscheinlichkeit höchstens

$$\varepsilon = 1 - \prod_{i=1}^q \left(1 - \frac{i}{2^c}\right)$$

ein Kollisionspaar (x, x') für $\text{Absorb}_{f, \text{id}, r}^i(x)$ aus. Wählt $\mathcal{S}$ nur von 0^c erreichbare Knoten v und kein Paar (v, x) mehrmals, so ist die Erfolgswahrscheinlichkeit exakt ε .

Beweis. Sei E_i das Ereignis “ G enthält nach i Durchläufen keine zwei verschiedenen Pfade von 0^c zu einem Knoten v ”. Da nur durch eine Kante zwischen zwei von 0^c aus erreichbaren Knoten ein zweiter Pfad von 0^c aus geschlossen werden kann und nach $i - 1$ Durchläufen höchstens i von 2^c Knoten erreichbar sind, gilt (unabhängig von $\mathcal{S}$):

$$\Pr[E_i | E_1 \cap \dots \cap E_{i-1}] \geq 1 - \frac{i}{2^c}.$$

Wählt $\mathcal{S}$ nur erreichbare Knoten und keine (v, x) mehrfach, so sind unter Annahme von $E_1 \cap \dots \cap E_{i-1}$ auch i Knoten erreichbar (sonst gäbe es bereits zwei Pfade von 0^c zu einem Knoten in G) und es gilt Gleichheit. Analog zum Beweis vom Satz 6 folgt der behauptete Wert ε , mit Gleichheit im Fall der Wahl erreichbarer Knoten durch $\mathcal{S}$. $\square$

Auch hier lässt sich q in Abhängigkeit von ε mittels $1 - x \approx e^{-x}$ abschätzen und es folgt:

$$q \approx c_\varepsilon 2^{\frac{c}{2}}, \quad c_\varepsilon = \sqrt{2 \ln \frac{1}{1 - \varepsilon}}.$$

1.2.11 SHA-3

Der Standard SHA-3 definiert die oben beschriebene Sponge-Konstruktion, 7 verschiedene bijektive Funktionen $f_w, w = 2^i, i \in \{0, \dots, 6\}$ als Kern des Sponges $\text{Sponge}_{f_w, \text{pad}10^*1_r, r}$, sowie verschiedene Kombinationen von Bitraten r und Ausgabelängen l (c ist durch $25w - r$ bestimmt).

Jede Funktion $f_w : \{0, 1\}^{5 \times 5 \times w} \rightarrow \{0, 1\}^{5 \times 5 \times w}$ bildet ein zweidimensionales Feld A aus w -Bit-Wörtern auf ein ebensolches Feld $f_w(A)$ ab. Dabei wird $(12 + \log_2 w)$ -mal eine

Rundenfunktion $f'_w : \{0, 1\}^{5 \times 5 \times w} \times \{0, 1\}^w \rightarrow \{0, 1\}^{5 \times 5 \times w}$ aufgerufen, die A und eine Rundenkonstante RC_i auf A' abbildet.

Es gilt

$$f'_w(A, RC) = \iota_{RC}(\chi(\pi(\rho(\theta(A))))),$$

wobei θ , ρ , π , χ und ι_{RC} Bijektionen von $\{0, 1\}^{5 \times 5 \times w}$ nach $\{0, 1\}^{5 \times 5 \times w}$ sind. Die Funktion θ besteht aus $\oplus$ -Operationen und ist so gewählt, dass sich $\theta^{-1}(A)$ an möglichst vielen Bits ändert, falls eines in A geflippt wird. Danach permutieren die Funktionen ρ und π die Bits von A innerhalb und zwischen den Wörtern. Ähnlich einer S-Box im SPN ist χ eine nichtlineare Funktion (die einzige solche in der Definition von f'_w), die nur auf 5-Bit-Blöcken arbeitet (jedes Bit hängt sogar nur von 2 anderen ab). Schlussendlich setzt ι_{RC} das Wort $A_{0,0}$ auf $A_{0,0} \oplus RC$.

Für die Werte $l \in \{224, 256, 384, 512\}$ definiert der Standard FIPS 202:

$$\text{SHA3-}l(x) = \text{Sponge}_{f_{1600}, \text{pad}10^*1_{r,r}}(l, x01), \quad \text{wobei } r = 1600 - 2l.$$

Das zusätzliche Padding 01 soll dabei **SHA-3** von anderen Anwendungen von Keccak mit denselben Werten w, l, r unterscheiden.

1.3 Nachrichten-Authentikationscodes (MACs)

Definition 13. Eine **Hashfamilie** $\mathcal{H} = (X, Y, K, H)$ wird durch folgende Komponenten beschrieben:

- X , eine endliche oder unendliche Menge von Texten,
- Y , endliche Menge aller möglichen **Hashwerte**, $\|Y\| \leq \|X\|$,
- K , endlicher **Schlüsselraum** (key space), wobei jeder Schlüssel $k \in K$ eine Hashfunktion $h_k: X \rightarrow Y$ in H spezifiziert, d.h. $H = \{h_k \mid k \in K\}$.

Im folgenden werden wir die Größe $\|X\|$ des Textraumes mit n , die des Hashwertbereiches Y mit m und die des Schlüsselraumes K mit l bezeichnen. Wir nennen dann $\mathcal{H}$ auch eine **(n, m, l)-Hashfamilie** oder einen **(n, m, l)-MAC**.

Damit ein geheimer Schlüssel k für die Authentifizierung mehrerer Nachrichten benutzt werden kann, ohne dass dies einem potentiellen Angreifer zur nichtautorisierten Berechnung von gültigen MAC-Werten verhilft, sollte folgende Bedingung erfüllt sein.

Berechnungsresistenz: Auch wenn eine Reihe von unter einem Schlüssel k generierten Text-Hashwert-Paaren $(x_1, h_k(x_1)), \dots, (x_n, h_k(x_n))$ bekannt ist, erfordert es einen immensen Aufwand, ohne Kenntnis von k ein weiteres Paar (x, y) mit $y = h_k(x)$ zu finden.

Bei Verwendung einer berechnungsresistenten Hashfunktion ist es einem Angreifer nicht möglich, an Bob eine Nachricht x zu schicken, die Bob als von Alice stammend anerkennt.

Verwendung eines MAC zur Versiegelung von Software

Mithilfe einer berechnungsresistenten Hashfunktion kann der Integritätsschutz für mehrere Datensätze auf die Geheimhaltung eines Schlüssels k zurückgeführt werden.

Um die Datensätze $x_1, \dots, x_n$ gegen unbefugt vorgenommene Veränderungen zu schützen, legt man sie zusammen mit ihren MAC-Werten $y_1 =$

$h_k(x_1), \dots, y_n = h_k(x_n)$ auf einem unsicheren Speichermedium ab und bewahrt den geheimen Schlüssel k an einem sicheren Ort auf. Bei einem späteren Zugriff auf einen Datensatz x_i lässt sich dessen Unversehrtheit durch einen Vergleich von y_i mit dem Ergebnis $h_k(x_i)$ einer erneuten MAC-Berechnung überprüfen.

Da auf diese Weise ein wirksamer Schutz der Datensätze gegen Viren und andere Manipulationen erreicht wird, spricht man von einer Versiegelung der gespeicherten Datensätze.

1.3.1 Angriffe gegen symmetrische Hashfunktionen

Ein Angriff gegen einen MAC hat die unbefugte Berechnung von MAC-Werten zum Ziel. Das heißt, der Angreifer versucht, MAC-Werte $h_k(x)$ ohne Kenntnis des geheimen Schlüssels k zu berechnen. Entsprechend der Art des zur Verfügung stehenden Textmaterials lassen sich die Angriffe gegen einen MAC wie folgt klassifizieren.

Impersonation

Der Angreifer kennt nur den benutzten MAC und versucht ein Paar (x, y) mit $h_k(x) = y$ zu generieren, wobei k der (dem Angreifer unbekannte) Schlüssel ist.

Substitution

Der Angreifer versucht in Kenntnis eines Paares $(x, h_k(x))$ ein Paar (x', y') mit $x' \neq x$ und $h_k(x') = y'$ zu generieren.

Angriff bei bekanntem Text (*known-text attack*)

Der Angreifer kennt für eine Reihe von Texten $x_1, \dots, x_r$ (die er nicht selbst wählen konnte) die zugehörigen MAC-Werte $h_k(x_1), \dots, h_k(x_r)$ und versucht, ein Paar (x', y') mit $h_k(x') = y'$ und $x' \notin \{x_1, \dots, x_r\}$ zu generieren.

Angriff bei frei wählbarem Text (*chosen-text attack*)

Der Angreifer kann die Texte x_i selbst wählen.

Angriff bei adaptiv wählbarem Text (*adaptive chosen-text attack*)

Der Angreifer kann die Wahl des Textes x_i von den zuvor erhaltenen MAC-Werten $h_k(x_j)$, $j < i$, abhängig machen.

Wechseln die Anwender nach jeder MAC-Wertberechnung den Schlüssel, so genügt es, dass $\mathcal{H}$ einem Impersonationsangriff widersteht.

1.3.2 Informationstheoretische Sicherheit von MACs

Modell: Schlüssel k und Nachrichten x werden unabhängig gemäß einer Wahrscheinlichkeitsverteilung $p(k, x) = p(k)p(x)$ generiert, welche dem Angreifer bekannt ist. Wir nehmen o.B.d.A. an, dass $p(x) > 0$ und $p(k) > 0$ für alle $x \in X$ und $k \in K$ gilt.

Erfolgswahrscheinlichkeit für Impersonation

Sei α die Wahrscheinlichkeit, mit der sich ein Angreifer bei optimaler Strategie als Alice ausgeben kann, ohne dass Bob dies bemerkt.

Für ein Paar (x, y) sei $p(x \mapsto y)$ die Wahrscheinlichkeit, dass ein zufällig gewählter Schlüssel den Text x auf den MAC-Wert y abbildet:

$$p(x \mapsto y) = p(y|x) = \sum_{k \in K(x,y)} p(k).$$

wobei $K(x, y) = \{k \in K \mid h_k(x) = y\}$ alle Schlüssel enthält, die x auf y abbilden. D.h. $p(x \mapsto y)$ ist die Wahrscheinlichkeit, dass Bob das Paar (x, y) als echt akzeptiert. Somit gibt $p(x \mapsto y)$ die Wahrscheinlichkeit an, mit der einem Angreifer bei Wahl des Paares (x, y) eine Impersonation gelingt, weshalb wir diese Wahrscheinlichkeit auch mit $\alpha(x, y)$ bezeichnen. Schließlich ist $\alpha(x) = \max\{\alpha(x, y) \mid y \in Y\}$ die Wahrscheinlichkeit, mit der einem Angreifer bei optimaler Strategie eine Impersonation mit dem Text x gelingt, und es gilt $\alpha = \max\{\alpha(x) \mid x \in X\}$.

Beispiel 14. Sei $K = \{1, 2, 3\}$, $X = \{a, b, c, d\}$ und $Y = \{0, 1\}$. Wir beschreiben H durch die zugehörige **Authentikationsmatrix**. Die Zeilen und Spalten dieser Matrix werden mit den Schlüsseln $k \in K$ und den Texten $x \in X$ indiziert und ihr Eintrag in Zeile k und Spalte x ist der Wert $h_k(x)$.

		0,1	0,2	0,3	0,4
		a	b	c	d
0,25	1	0	0	0	1
0,30	2	1	1	0	1
0,45	3	0	1	1	0

Die umrahmten Zahlen geben die Wahrscheinlichkeiten $p(x)$ bzw. $p(k)$ an. Dann hat der Angreifer folgende Erfolgsaussichten $\alpha(x)$, falls er an Bob den Text x senden möchte.

x	a	b	c	d
$p(x \mapsto 0)$	0,7	0,25	0,55	0,45
$p(x \mapsto 1)$	0,3	0,75	0,45	0,55
$\alpha(x)$	0,7	0,75	0,55	0,55

Folglich ist $\alpha = 0,75$. ◁

Satz 15. Für alle $x \in X$ ist $\alpha(x) \geq \frac{1}{m}$ und daher gilt $\alpha \geq \frac{1}{m}$.

Beweis. Sei $x \in X$ beliebig. Dann gilt

$$\sum_{y \in Y} p(x \mapsto y) = \sum_{y \in Y} \sum_{k \in K(x, y)} p(k) = \sum_{k \in K} p(k) = 1.$$

Somit existiert für jedes $x \in X$ ein $y \in Y$ mit $p(x \mapsto y) \geq \frac{1}{m}$ und dies impliziert

$$\alpha(x) = \max_{y \in Y} p(x \mapsto y) \geq \frac{1}{m}.$$

◻

Bemerkung 16. Wie der Beweis zeigt, gilt $\alpha = \frac{1}{m}$ genau dann, wenn für alle Paare $(x, y) \in X \times Y$ gilt,

$$\sum_{k \in K(x, y)} p(k) = \frac{1}{m}.$$

D.h. bei Gleichverteilung der Schlüssel muss in jeder Spalte der Authentikationsmatrix jeder MAC-Wert gleich oft vorkommen. Dies lässt sich am einfachsten dadurch erreichen, dass man $K = Y$ setzt und für h_k die konstante Funktion $h_k(x) = k$ wählt.

Das folgende Lemma benötigen wir für den Beweis des nächsten Satzes.

Lemma 17. *Sei $\mathcal{X}$ eine Zufallsvariable mit endlichem Wertebereich $W(\mathcal{X}) \subseteq \mathbb{R}^+$. Dann gilt $\log E(\mathcal{X}) \geq E(\log \mathcal{X})$.*

Beweis. Sei $W(\mathcal{X}) = \{x_1, \dots, x_n\}$ und für $i = 1, \dots, n$ sei $p_i = \Pr[\mathcal{X} = x_i]$. Da die Funktion $x \mapsto \log_2 x$ konkav ist, folgt mit der Jensenschen Ungleichung

$$\log E(\mathcal{X}) = \log_2(\sum p_i x_i) \geq \sum p_i \log_2 x_i = E(\log \mathcal{X}).$$

□

Satz 18. *Für jeden MAC (X, Y, K, H) gilt:*

$$\alpha \geq \frac{1}{2^{H(\mathcal{K}) - H(\mathcal{K}|\mathcal{X}, \mathcal{Y})}} (\geq 1/l).$$

Hierbei sind $\mathcal{X}, \mathcal{Y}, \mathcal{K}$ Zufallsvariablen, die die Verteilungen der Nachrichten, der MAC-Werte und der Schlüssel beschreiben.

Der Wert von α kann also um so kleiner werden, je gleichmäßiger die Schlüsselverteilung ist und je mehr Information die Beobachtung eines gültigen Paares (x, y) über den Schlüssel liefert.

Beweis. Da $\alpha = \max_{x,y} \alpha(x, y)$ ist, folgt $E(\alpha(\mathcal{X}, \mathcal{Y})) = \sum_{x,y} p(x, y) \alpha(x, y) \leq \alpha$, wobei $E(\alpha(\mathcal{X}, \mathcal{Y}))$ die Erfolgswahrscheinlichkeit eines (probabilistischen) Angreifers ist, der das Paar (x, y) gemäß der Verteilung $(\mathcal{X}, \mathcal{Y})$ wählt. Somit folgt unter Anwendung von Lemma 17,

$$\log \alpha \geq \log E(\alpha(\mathcal{X}, \mathcal{Y})) \geq E(\log \alpha(\mathcal{X}, \mathcal{Y})) = \sum_{x,y} \underbrace{p(x, y)}_{p(x)p(y|x)} \underbrace{\log p(y|x)}_{-\log \frac{1}{p(y|x)}} = -H(\mathcal{Y}|\mathcal{X}).$$

Wegen

$$H(\mathcal{K}, \mathcal{Y}, \mathcal{X}) = H(\mathcal{X}) + H(\mathcal{Y}|\mathcal{X}) + H(\mathcal{K}|\mathcal{X}, \mathcal{Y})$$

und

$$H(\mathcal{K}, \mathcal{Y}, \mathcal{X}) = \underbrace{H(\mathcal{K}, \mathcal{X})}_{=H(\mathcal{K})+H(\mathcal{X})} + \underbrace{H(\mathcal{Y}|\mathcal{K}, \mathcal{X})}_{=0}.$$

gilt zudem $H(\mathcal{Y}|\mathcal{X}) = H(\mathcal{K}) - H(\mathcal{K}|\mathcal{X}, \mathcal{Y})$ und somit $\log \alpha \geq H(\mathcal{K}|\mathcal{X}, \mathcal{Y}) - H(\mathcal{K})$. □

Beispiel 19 (Fortsetzung von Beispiel 14). *Es gilt*

$$H(\mathcal{K}) = \sum_k p(k) \log \frac{1}{p(k)} = 0,45 \cdot 1,152 + 0,3 \cdot 1,737 + 0,25 \cdot 2,0 = 1,54.$$

Um $H(\mathcal{K}|\mathcal{X}, \mathcal{Y})$ zu bestimmen, benötigen wir die bedingten Verteilungen $\mathcal{K}_{x,y}$ für alle Paare $(x, y) \in X \times Y$.

(x, y)	$(a, 0)$	$(a, 1)$	$(b, 0)$	$(b, 1)$	$(c, 0)$	$(c, 1)$	$(d, 0)$	$(d, 1)$
$p(1 x, y)$	$\frac{5}{14}$	0	1	0	$\frac{5}{11}$	0	0	$\frac{5}{11}$
$p(2 x, y)$	0	1	0	$\frac{2}{5}$	$\frac{6}{11}$	0	0	$\frac{6}{11}$
$p(3 x, y)$	$\frac{9}{14}$	0	0	$\frac{3}{5}$	0	1	1	0
$H(\mathcal{K} x, y)$	$\approx 0,94$	0	0	$\approx 0,97$	$\approx 0,99$	0	0	$\approx 0,99$
$p(x, y)$	0,07	0,03	0,05	0,15	0,165	0,135	0,18	0,22

Hierbei gilt $p(x, y) = p(x)p(y|x) = p(x)p(x \mapsto y)$. Somit ist

$$H(\mathcal{K}|\mathcal{X}, \mathcal{Y}) = \sum_{x,y} p(x, y) H(\mathcal{K}|x, y) \approx 0,52$$

und wir erhalten die untere Schranke $\alpha \geq \frac{1}{2^{H(\mathcal{K}) - H(\mathcal{K}|\mathcal{X}, \mathcal{Y})}} \approx \frac{1}{2^{1,54 - 0,52}} = \frac{1}{2^{1,02}} \approx 0,493$.

Erfolgswahrscheinlichkeit für Substitution

Bezeichne β die Wahrscheinlichkeit, mit der ein Angreifer bei optimaler Strategie eine von Alice gesendete Nachricht x durch eine andere Nachricht x' ersetzen kann, ohne dass Bob dies bemerkt. Dabei gehen wir davon aus, dass der Angreifer keinen Einfluss auf die Wahl der von Alice gesendeten Nachricht x hat.

Falls der Angreifer ein von Alice gesendetes Paar (x, y) durch das Paar (x', y') ersetzt, ist seine Erfolgswahrscheinlichkeit gleich der bedingten Wahrscheinlichkeit

$$p(x' \mapsto y' | x \mapsto y) = \frac{p(x \mapsto y, x' \mapsto y')}{p(x \mapsto y)} = \frac{\sum_{k \in K(x, y, x', y')} p(k)}{\sum_{k \in K(x, y)} p(k)},$$

dass ein zufällig gewählter Schlüssel k den Text x' auf y' abbildet, wenn bereits bekannt ist, dass $h_k(x) = y$ ist. Falls Alice also das Paar (x, y) sendet, so ist die maximale Erfolgswahrscheinlichkeit des Angreifers

$$\beta(x, y) := \max_{x' \neq x, y'} p(x' \mapsto y' | x \mapsto y).$$

Man beachte, dass $\beta(x, y)$ nur im Fall $p(x, y) > 0$ definiert ist. Da der Angreifer keinen Einfluss auf die Wahl von (x, y) hat, ist β gleich dem Erwartungswert von $\beta(x, y)$ unter der Verteilung $p(x, y)$, mit der Alice diese Paare generiert. Somit ergibt sich β zu

$$\beta = E(\beta(\mathcal{X}, \mathcal{Y})) = \sum_{x \in X, y \in Y} p(x, y) \beta(x, y).$$

Wegen $p(x, y) = p(x)p(x \mapsto y)$ können wir β unter Verwendung der Funktion

$$\beta'(x, y) = \beta(x, y)p(x \mapsto y) = \max_{x' \neq x, y'} p(x' \mapsto y', x \mapsto y)$$

auch einfacher mittels der Formel $\beta = \sum_{x \in X} p(x) \sum_{y \in Y} \beta'(x, y)$ berechnen.

Beispiel 20 (Fortsetzung von Beispiel 14).

(x, y)	$p(x' \mapsto y', x \mapsto y)$								$\beta'(x, y)$	$p(x \mapsto y)$	$\beta(x, y)$
	$(a, 0)$	$(a, 1)$	$(b, 0)$	$(b, 1)$	$(c, 0)$	$(c, 1)$	$(d, 0)$	$(d, 1)$			
$(a, 0)$			0,25	0,45	0,25	0,45	0,45	0,25	0,45	0,7	0,643
$(a, 1)$			0	0,3	0,3	0	0	0,3	0,3	0,3	1
$(b, 0)$	0,25	0			0,25	0	0	0,25	0,25	0,25	1
$(b, 1)$	0,45	0,3			0,3	0,45	0,45	0,3	0,45	0,75	0,6
$(c, 0)$	0,25	0,3	0,25	0,3			0	0,55	0,55	0,55	1
$(c, 1)$	0,45	0	0	0,45			0,45	0	0,45	0,45	1
$(d, 0)$	0,45	0	0	0,45	0	0,45			0,45	0,45	1
$(d, 1)$	0,25	0,3	0,25	0,3	0,55	0			0,55	0,55	1

Die optimalen Wahlmöglichkeiten des Angreifers, ein Paar (x, y) durch ein anderes Paar (x', y') zu ersetzen, sind in der Tabelle fett gedruckt. Für β erhalten wir somit den Wert

$$\begin{aligned} \beta &= \sum_{x \in X} p(x) \sum_{y \in Y} \beta'(x, y) \\ &= 0,1(0,45 + 0,3) + 0,2(0,25 + 0,45) + 0,3(0,55 + 0,45) + 0,4(0,45 + 0,55) \\ &= 0,915. \end{aligned}$$

Als nächstes zeigen wir für β die gleiche untere Schranke wie für α .

Satz 21. Für alle $(x, y) \in X \times Y$ mit $p(x, y) > 0$ ist $\beta(x, y) \geq \frac{1}{m}$ und daher gilt $\beta \geq \frac{1}{m}$.

Beweis. Sei $(x, y) \in X \times Y$ ein Paar mit $p(x, y) > 0$. Dann gilt für beliebige $x' \in X - \{x\}$,

$$\sum_{y' \in Y} p(x' \mapsto y' | x \mapsto y) = \frac{\sum_{y' \in Y} \sum_{k \in K(x', y'; x, y)} p(k)}{\sum_{k \in K(x, y)} p(k)} = 1.$$

Somit existiert ein $y' \in Y$ mit $p(x' \mapsto y' | x \mapsto y) \geq \frac{1}{m}$ und dies impliziert

$$\beta(x, y) = \max_{x' \neq x, y'} p(x' \mapsto y' | x \mapsto y) \geq \frac{1}{m}.$$

Folglich ist

$$\beta = \sum_{x \in X, y \in Y} p(x, y) \beta(x, y) \geq \frac{1}{m} \sum_{x \in X, y \in Y} p(x, y) = \frac{1}{m}.$$

□

Beispiel 22. Sei $X = Y = \{0, 1, 2\} = \mathbb{Z}_3$ und sei $K = \mathbb{Z}_3 \times \mathbb{Z}_3$. Für $k = (a, b) \in K$ und $x \in X$ sei

$$h_k(x) = ax + b \bmod 3.$$

Die zugehörige **Authentikationsmatrix** ist

	0	1	2
(0, 0)	0	0	0
(0, 1)	1	1	1
(0, 2)	2	2	2
(1, 0)	0	1	2
(1, 1)	1	2	0
(1, 2)	2	0	1
(2, 0)	0	2	1
(2, 1)	1	0	2
(2, 2)	2	1	0

Wir nehmen an, dass der Schlüssel unter Gleichverteilung gewählt wird. Ersetzt der Angreifer ein Paar (x, y) durch ein Paar (x', y') mit $x' \neq x$, so wird dieses Paar von genau einem der 3 infrage kommenden Schlüssel akzeptiert. Dies liegt daran, dass in je 2 Spalten der Authentikationsmatrix jedes MAC-Wertepaar genau einmal vorkommt. Folglich ist $p(x' \mapsto y' | x \mapsto y) = 1/3$ und somit hat β den optimalen Wert $\beta = 1/3$. ◁

Lemma 23. Sei (X, Y, K, H) ein MAC mit $\beta = \frac{1}{m}$. Dann gilt

$$p(x' \mapsto y' | x \mapsto y) = 1/m$$

für alle Doppelpaare (x, y, x', y') mit $x \neq x'$.

Beweis. Wir zeigen zuerst, dass die Behauptung unter der Voraussetzung $p(x \mapsto y) > 0$ gilt. Wäre nämlich

$$p(x' \mapsto y' | x \mapsto y) > 1/m,$$

dann wäre auch

$$\beta(x, y) = \max_{x' \neq x, y'} p(x' \mapsto y' | x \mapsto y) > 1/m.$$

Da für alle Paare (u, v) mit $p(u \mapsto v) > 0$ nach Satz 21 die Ungleichung $\beta(u, v) \geq 1/m$ gilt und zudem $p(x, y) = p(x)p(x \mapsto y) > 0$ ist, folgt hieraus

$$\beta = \sum_{x \in X, y \in Y} p(x, y)\beta(x, y) > 1/m,$$

was im Widerspruch zur Voraussetzung des Satzes steht. Ist andererseits

$$p(x' \mapsto y' | x \mapsto y) < 1/m,$$

muss wegen

$$\sum_{y'' \in Y} p(x' \mapsto y'' | x \mapsto y) = 1$$

auch ein MAC-Wert y'' mit $p(x' \mapsto y'' | x \mapsto y) > 1/m$ existieren, woraus sich wie bereits gezeigt ein Widerspruch ergibt.

Es bleibt zu zeigen, dass $p(x \mapsto y) > 0$ für alle Paare (x, y) gilt. Wäre $p(x \mapsto y) = 0$, so würde für ein beliebiges Paar (u, v) mit $p(u \mapsto v) > 0$ auch $p(x \mapsto y | u \mapsto v) = 0$ sein. Wie bereits gezeigt, steht dies jedoch im Widerspruch zur Voraussetzung $\beta = 1/m$. $\square$

Satz 24. Ein MAC (X, Y, K, H) erfüllt $\beta = \frac{1}{m}$ genau dann, wenn

$$p(x \mapsto y, x' \mapsto y') = 1/m^2$$

für alle Doppelpaare (x, y, x', y') mit $x \neq x'$ gilt.

Beweis. Sei (X, Y, K, H) ein MAC mit $\beta = \frac{1}{m}$. Nach obigem Lemma impliziert dies, dass

$$p(x' \mapsto y' | x \mapsto y) = 1/m$$

für alle Doppelpaare (x, y, x', y') mit $x \neq x'$ gilt. Dies impliziert nun

$$p(x' \mapsto y') = \sum_y p(x \mapsto y)p(x' \mapsto y' | x \mapsto y) = 1/m$$

und daher

$$p(x \mapsto y, x' \mapsto y') = p(x' \mapsto y')p(x \mapsto y | x' \mapsto y') = 1/m^2.$$

Umgekehrt rechnet man leicht nach, dass die Bedingung $\beta = \frac{1}{m}$ erfüllt ist, wenn für alle Doppelpaare (x, y, x', y') mit $x \neq x'$ die Gleichheit $p(x \mapsto y, x' \mapsto y') = 1/m^2$ gilt. $\square$

Bemerkung 25. Nach obigem Satz gilt $\beta = \frac{1}{m}$ genau dann, wenn für alle Doppelpaare (x, y, x', y') mit $x \neq x'$ gilt,

$$p(x \mapsto y, x' \mapsto y') = \sum_{k \in K(x, y, x', y')} p(k) = \frac{1}{m^2}.$$

D.h. bei Gleichverteilung der Schlüssel gilt $\beta = \frac{1}{m}$ genau dann, wenn in je zwei Spalten der Authentikationsmatrix jedes MAC-Wertepaar gleich oft vorkommt.

Ab jetzt setzen wir voraus, dass der Schlüssel unter Gleichverteilung gewählt wird, d.h. es gilt $p(k) = \frac{1}{|K|}$ für alle $k \in K$.

Definition 26. Ein MAC (X, Y, K, H) heißt **2-universal**, falls für alle $x, x' \in X$ mit $x \neq x'$ und alle $y, y' \in Y$ gilt:

$$\|K(x, y, x', y')\| = \frac{\|K\|}{m^2}.$$

Bemerkung 27. Bei der Konstruktion von 2-universalen MACs spielt der Parameter $\lambda = \frac{\|K\|}{m^2}$ eine wichtige Rolle. Da λ notwendigerweise positiv und ganzzahlig ist, muss insbesondere $\|K\| \geq m^2$ gelten.

Im Folgenden nennen wir einen 2-universalen (n, m, l) -MAC mit $\lambda = l/m^2$ kurz einen (n, m, l, λ) -MAC.

Auf Grund von Bemerkung 25 ist klar, dass ein MAC bei gleichverteilten Schlüsseln genau dann die Bedingung $\beta = \frac{1}{m}$ erfüllt, wenn er 2-universal ist. Auf Grund von Bemerkung 16 nimmt in diesem Fall auch α den optimalen Wert $\frac{1}{m}$ an.

Der nächste Satz zeigt eine einfache Konstruktionsmöglichkeit von 2-universalen MACs mit dem Parameterwert $\lambda = 1$.

Satz 28. Sei p prim und für $a, b, x \in \mathbb{Z}_p$ sei

$$h_{a,b}(x) = ax + b \bmod p.$$

Dann ist (X, Y, K, H) mit $X = Y = \mathbb{Z}_p$ und $K = \mathbb{Z}_p \times \mathbb{Z}_p$ ein $(p, p, p^2, 1)$ -MAC.

Beweis. Wir müssen zeigen, dass die Größe von $K(x, y, x', y')$ für alle Doppelpaare (x, y, x', y') mit $x \neq x'$ konstant ist. Ein Schlüssel (a, b) gehört genau dann zu dieser Menge, wenn er die beiden Kongruenzen

$$\begin{aligned} ax + b &\equiv_p y, \\ ax' + b &\equiv_p y' \end{aligned}$$

erfüllt. Da dies jedoch nur auf den Schlüssel (a, b) mit

$$\begin{aligned} a &= (y' - y)(x' - x)^{-1} \bmod p, \\ b &= y - x(y' - y)(x' - x)^{-1} \bmod p \end{aligned}$$

zutrifft, folgt $\|K(x', y', x, y)\| = 1$. □

Die Hashfunktionen des vorigen Satzes erfüllen wegen $n = m = p$ nicht die Kompressions-eigenschaft. Zwar lässt sich n noch geringfügig von p auf $p + 1$ vergrößern, ohne K und Y (und damit λ) zu verändern (siehe Übungen). Wie der nächste Satz zeigt, lässt sich eine stärkere Kompression mit dem Parameterwert $\lambda = 1$ jedoch nicht realisieren.

Satz 29. Für einen $(n, m, l, 1)$ -MAC gilt

$$m \geq n - 1$$

und somit $l = m^2 \geq (n - 1)^2$ sowie $n/m \leq \frac{m+1}{m} (\approx 1)$.

Beweis. O.B.d.A. sei $\|K\| = \{1, \dots, l\}$ und $Y = \{1, \dots, m\}$. Es ist leicht zu sehen, dass eine (bijektive) Umbenennung $\pi: Y \rightarrow Y$ der MAC-Werte in einer einzelnen Spalte der Authentikationsmatrix A wieder auf einen 2-universalen MAC führt. Also können wir zudem annehmen, dass die erste Zeile der Authentikationsmatrix A nur Einsen enthält. Da A 2-universal ist, gilt:

- In jeder Zeile $i = 2, \dots, m^2$ kommt höchstens eine Eins vor.
- Jede Spalte j enthält eine Eins in Zeile 1 und $m - 1$ Einsen in den übrigen Zeilen.

Da in den Zeilen $i = 2, \dots, m^2$ insgesamt genau $n(m - 1)$ Einsen vorkommen, folgt

$$\underbrace{\text{Anzahl der Zeilen}}_{m^2} \geq \underbrace{\text{Anzahl der Zeilen mit einer Eins}}_{1+n(m-1)},$$

was $m^2 - 1 \geq n(m - 1)$ bzw. $n \leq m + 1$ impliziert. $\square$

Der nächste Satz liefert 2-universale MACs mit beliebig großem Kompressionsfaktor. Für den Beweis benötigen wir das folgende Lemma.

Lemma 30. *Sei A eine $(k \times \ell)$ -Matrix über einem endlichen Körper $\mathbb{F}$, deren k Zeilen linear unabhängig sind. Dann besitzt das lineare Gleichungssystem*

$$Ax = y$$

für jedes $y \in \mathbb{F}^k$ genau $|\mathbb{F}|^{\ell-k}$ Lösungen $x \in \mathbb{F}^\ell$.

Beweis. Siehe Übungen. $\square$

Satz 31. *Sei p prim und für $x = (x_1, \dots, x_d) \in \{0, 1\}^d$ und $k = (k_1, \dots, k_d) \in \mathbb{Z}_p^d$ sei*

$$h_k(x) = kx = \sum_{i=1}^d k_i x_i \pmod{p}.$$

Dann ist (X, Y, K, H) mit $X = \{0, 1\}^d - \{0^d\}$, $Y = \mathbb{Z}_p$ und $K = \mathbb{Z}_p^d$ ein $(2^d - 1, p, p^d, p^{d-2})$ -MAC.

Beweis. Wir müssen zeigen, dass die Größe von $K(x, y, x', y')$ für alle Doppelpaare (x, y, x', y') mit $x \neq x'$ konstant ist. Es gilt

$$\begin{aligned} k \in K(x, y, x', y') &\Leftrightarrow h_k(x) = y \wedge h_k(x') = y' \\ &\Leftrightarrow k \cdot x = y \wedge k \cdot x' = y'. \end{aligned}$$

Fassen wir $x = x_1 \cdots x_d$ und $x' = x'_1 \cdots x'_d$ zu einer Matrix A zusammen, so ist dies äquivalent zu

$$\begin{pmatrix} x_1 & \cdots & x_d \\ x'_1 & \cdots & x'_d \end{pmatrix} \cdot \begin{pmatrix} k_1 \\ \vdots \\ k_d \end{pmatrix} = \begin{pmatrix} y \\ y' \end{pmatrix}.$$

Da die beiden Zeilen von A verschieden und damit linear unabhängig sind, folgt mit obigem Lemma, dass genau $|K(x, y, x', y')| = p^{d-2}$ Schlüssel $k = (k_1, \dots, k_d)$ mit dieser Eigenschaft existieren. $\square$

Bemerkung 32. *Obige Konstruktion liefert einen λ -Wert von $\frac{|K|}{m^2} = p^{d-2}$. Durch Erweiterung von X auf eine geeignete Teilmenge $X' \subseteq \mathbb{Z}_p^d$ lässt sich der Textraum von $2^d - 1$ auf $\frac{p^d - 1}{p - 1}$ vergrößern (siehe Übungen). Dies führt auf einen beliebig groß wählbaren Kompressionsfaktor von $\frac{p^d - 1}{p(p - 1)}$ bei einem λ -Wert von $\lambda = p^{d-2}$. Wie der nächste Satz zeigt, lässt sich dies nicht mit einem kleineren λ -Wert erreichen.*

Im Beweis des nächsten Satzes benötigen wir folgendes Lemma.

Lemma 33. Für beliebige reelle Zahlen $b_1, \dots, b_m \in \mathbb{R}$ gilt $\left(\sum_{i=1}^m b_i\right)^2 \leq m \sum_{i=1}^m b_i^2$.

Beweis. Da die Funktion $x \mapsto x^2$ konvex ist, folgt mit der Jensenschen Ungleichung $(\sum b_i/m)^2 \leq \sum b_i^2/m$ und somit

$$\left(\sum b_i\right)^2 = m^2 \underbrace{\left(\sum b_i/m\right)^2}_{\leq \sum b_i^2/m} \leq m \sum b_i^2.$$

□

Satz 34. Für einen (n, m, l, λ) -MAC gilt

$$l = \lambda m^2 \geq n(m-1) + 1$$

und somit $n/m \leq (\lambda - 1/m^2) \frac{m}{m-1} (\approx \lambda)$.

Beweis. O.B.d.A. können wir wieder $\|K\| = \{k_1, \dots, k_l\}$ und $Y = \{1, \dots, m\}$ annehmen, und dass die erste Zeile der Authentikationsmatrix nur aus Einsen besteht. Für jede Zeile $i = 1, \dots, l$ bezeichne e_i die Anzahl der Einsen in dieser Zeile (also $e_1 = n$). Da in jeder Spalte jeder MAC-Wert genau λm -mal vorkommt, gilt

$$\sum_{i=1}^l e_i = \lambda n m \quad \text{und} \quad \sum_{i=2}^l e_i = \lambda n m - n = n(\lambda m - 1).$$

Sei $z = \sum_{i=2}^l z_i$, wobei z_i die Anzahl von Spaltenpaaren (j, j') mit $j \neq j'$ und $h_{k_i}(x_j) = h_{k_i}(x_{j'}) = 1$ ist. Dann folgt

$$z = \sum_{i=2}^l z_i = \sum_{i=2}^l e_i(e_i - 1) = \sum_{i=2}^l e_i^2 - \sum_{i=2}^l e_i = \sum_{i=2}^l e_i^2 - n(\lambda m - 1).$$

Mit obigem Lemma ergibt sich

$$\sum_{i=2}^l e_i^2 \geq \frac{\left(\sum_{i=2}^l e_i\right)^2}{l-1} = \frac{(n(\lambda m - 1))^2}{l-1}.$$

Da andererseits in jedem Spaltenpaar das MAC-Wertepaar $(1, 1)$ in genau λ Zeilen vorkommt (genauer: einmal in Zeile 1 und $(\lambda - 1)$ -mal in den Zeilen $i = 2, \dots, l$), und da $n(n-1)$ solche Spaltenpaare existieren, ergibt sich andererseits die Gleichung

$$z = (\lambda - 1)n(n-1).$$

Somit erhalten wir

$$\begin{aligned} (\lambda - 1)n(n-1) &= z = \sum_{i=2}^l e_i^2 - n(\lambda m - 1) \geq \frac{(n(\lambda m - 1))^2}{l-1} - n(\lambda m - 1) \\ &\Rightarrow ((\lambda - 1)n(n-1) + n(\lambda m - 1))(\lambda m^2 - 1) \geq (n(\lambda m - 1))^2 \\ &\Rightarrow (\lambda n - n - \lambda + \lambda m)(\lambda m^2 - 1) \geq n(\lambda m - 1)^2 \\ &\Rightarrow -\lambda^2 m^2 + \lambda^2 m^3 \geq \lambda n m^2 + \lambda n - \lambda + \lambda m - 2\lambda n m \\ &\Rightarrow \lambda^2(m^3 - m^2) \geq \lambda(n(m-1)^2 + m - 1) \\ &\Rightarrow \lambda m^2 \geq n(m-1) + 1. \end{aligned}$$

□