

Kryptologie

Johannes Köbler



Institut für Informatik
Humboldt-Universität zu Berlin

WS 2022/23

Anmeldung

- über Agnes
- und bei Moodle (wegen Punktevergabe und Bildung von Abgabegruppen)
- Mails von Agnes und von Moodle werden standardmäßig an den HU-Account gesendet (bitte regelmäßig checken oder weiterleiten)

Aktuelle Infos auf der VL-Webseite unter

- <https://hu.berlin/vlkrypto> bzw.
- <https://www.informatik.hu-berlin.de/de/forschung/gebiete/algorithmenII/Lehre/ws22/krypto>

Skript, Folien und Aufgabenblätter

- Skript und Folien werden wöchentlich ins Netz gestellt
- Übungsblätter werden in der Regel dienstags auf der VL-Seite veröffentlicht
- Die Besprechung der mündlichen Aufgaben erfolgt am Donnerstag der Folgewoche
- Lösungen dazu können bis zum Tag davor digital in Moodle hochgeladen werden, Details siehe dort
- Die schriftlichen Aufgaben sind bis Dienstag zwei Wochen nach Ausgabe um 23:59 Uhr abzugeben
- Fragen zu Übung und Vorlesung können auch im Moodle-Forum gestellt und diskutiert werden

- in Gruppen von bis zu drei Teilnehmern
- Lösungen für die schriftlichen Aufgaben sollten als PDF abgegeben werden
- die Abgabe von Lösungsvorschlägen für die mündlichen Aufgaben ist freiwillig und geht nicht in die Punktwertung ein
- besonders gut gelungene Lösungen werden mit Zustimmung der Abgebenden im Forum veröffentlicht

Scheinkriterien

- Lösen von mindestens 50% der schriftlichen Aufgaben

Prüfungsform

- mündlich (in den Prüfungszeiträumen des Wintersemesters)
- Der Übungsschein ist **nicht** Voraussetzung für die Zulassung zur Prüfung

Gibt es zum organisatorischen Ablauf noch Fragen?

- Kryptografische Verfahren schaffen Vertrauen in ungeschützten Umgebungen
- Sie ermöglichen sichere Kommunikation über unsichere Kanäle und können verhindern, dass sich ein Kommunikationspartner unfair verhält
- In unsicheren Umgebungen wie dem Internet können sie die aus direkter Interaktion gewohnte Sicherheit herstellen
- Und auch die Interaktion in sicheren Umgebungen wird um Möglichkeiten erweitert, die ohne Kryptografie nicht denkbar wären
- Im Bachelormodul **Einführung in die Kryptologie** haben wir uns mit den mathematischen Grundlagen von kryptografischen Verfahren beschäftigt, wobei (symmetrische und asymmetrische) Verschlüsselungsverfahren im Vordergrund standen
- Im aktuellen Mastermodul **Kryptologie** werden wir dagegen kryptografische Verfahren und Protokolle für andere Schutzziele betrachten wie z.B. Hashverfahren und digitale Signaturen sowie Pseudozufallsgeneratoren

- Kryptosysteme (Verschlüsselungsverfahren) dienen der Geheimhaltung von Nachrichten bzw. Daten
- Hierzu gibt es auch andere Methoden wie z.B.
 - Physikalische Maßnahmen: Tresor etc.
 - Organisatorische Maßnahmen: einsamer Waldspaziergang etc.
 - Steganografische Maßnahmen: unsichtbare Tinte etc.

Überblick weiterer Schutzziele

Andererseits können durch kryptografische Verfahren weitere **Schutzziele** realisiert werden wie z.B.

- **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

In das Umfeld der Kryptologie fallen die folgenden Begriffe

- **Kryptografie:**

Lehre von der Geheimhaltung von Informationen durch Verschlüsselung
Im weiteren Sinne: Wissenschaft von der Übermittlung, Speicherung und Verarbeitung von Daten in einer von potentiellen Gegnern bedrohten Umgebung

- **Kryptoanalysis:**

Erforschung der Methoden eines unbefugten Angriffs gegen ein Kryptoverfahren
Zweck: Vereitelung der mit seinem Einsatz verfolgten Ziele

- **Kryptoanalyse:**

Analyse eines Kryptoverfahrens zum Zweck der Bewertung seiner kryptografischen Stärken und Schwächen

- **Kryptologie:**

Wissenschaft vom Entwurf, der Anwendung und der Analyse von kryptografischen Verfahren (umfasst Kryptografie und Kryptoanalyse)

- 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

- Idee: Eine kryptografische Hashfunktion h soll zu einem vorgegebenen Text x eine zwar kompakte aber dennoch repräsentative Beschreibung $h(x)$ liefern, die unter praktischen Gesichtspunkten zur eindeutigen Identifikation von x benutzt werden kann
- Die Funktion h muss also „charakteristische Merkmale“ von x in den Hashwert $h(x)$ einfließen lassen
- Da der Fingerabdruck auch diese Eigenschaften besitzt, wird $h(x)$ oft als **digitaler Fingerabdruck** von x bezeichnet
- Gebräuchlich sind auch die Bezeichnungen **kryptografische Prüfsumme** oder **message digest** (engl. für „Nachrichtenextrakt“)

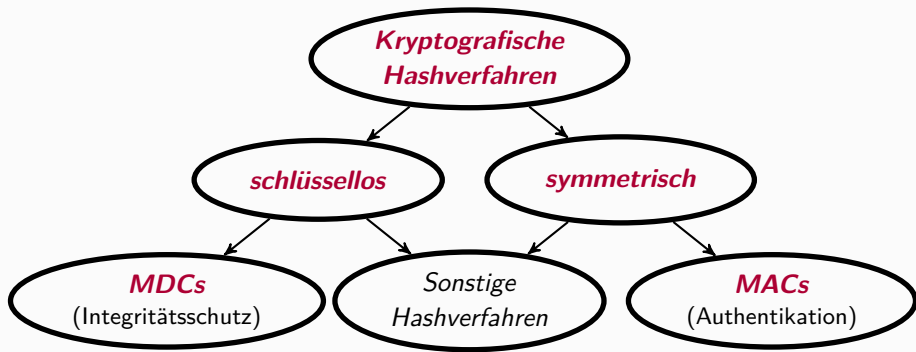
Typische Schutzziele, die sich mittels Hashfunktionen realisieren lassen:

Nachrichtenthauthentikation (message authentication)

- 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)

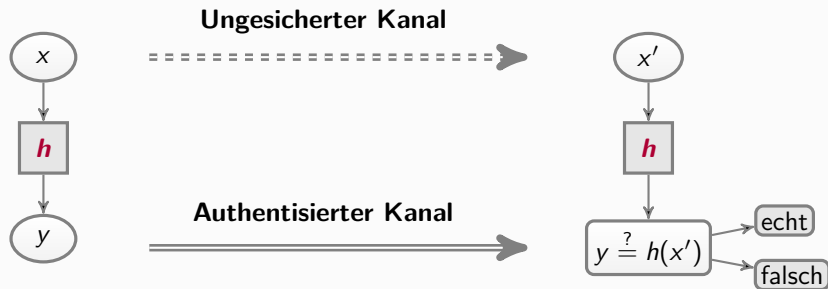
- Wie kann sich eine Person (oder ein Gerät) anderen gegenüber zweifelsfrei ausweisen?



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

- 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** (*M*anipulation *D*etection *C*ode) bezeichnet
- Zuweilen wird das Kürzel **MDC** auch als eine Abkürzung für *M*odification *D*etection *C*ode verwendet
- Seltener ist dagegen die Bezeichnung **MIC** (*m*essage *i*ntegrity *c*odes)

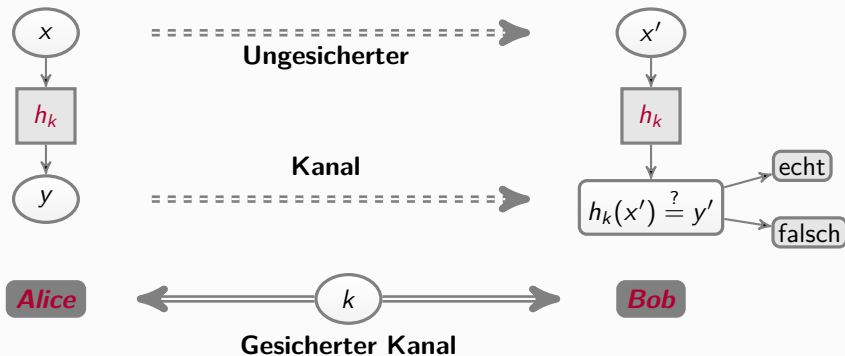
Um die Integrität eines Datensatzes x sicherzustellen, der über einen ungesicherten Kanal gesendet (bzw. auf einem vor Manipulationen nicht sicheren Webserver abgelegt) wird, kann man wie folgt verfahren:



- Der **MDC**-Hashwert $y = h(x)$ von x wird auf einem authentisierten Kanal übertragen
- Nach der Übertragung wird geprüft, ob der Datensatz noch den Hashwert y liefert

- Kryptografische Hashverfahren mit symmetrischen Schlüsseln finden hauptsächlich bei der Authentifizierung von Nachrichten Verwendung
- Diese werden daher auch als **MAC** (*message authentication code*) oder als **Authentikationscode** bezeichnet
- Daneben gibt es auch Hashverfahren mit asymmetrischen Schlüsseln
- Diese werden jedoch der Rubrik der Signaturverfahren zugeordnet, da mit ihnen ausschließlich digitale Signaturen gebildet werden

- Mit einem MAC lassen sich Nachrichten wie folgt authentisieren:



- Alice fügt der Nachricht x den Hashwert $y = h_k(x)$ hinzu und sendet beides über den unsicheren Kanal an Bob
- Bob überprüft die Echtheit einer Nachricht (x', y') , indem er den zu x' gehörigen Hashwert $h_k(x')$ berechnet und 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, hat Bob nicht nur die Möglichkeit, anhand von y die empfangene Nachricht x' auf Manipulationen, sondern auch ihre Herkunft zu überprüfen

- Wir betrachten nun verschiedene Sicherheitsanforderungen an MDCs h
- Dabei nehmen wir an, dass $h: X \rightarrow Y$ öffentlich bekannt ist
- Ein Paar $(x, y) \in X \times Y$ heißt **gültig** für h , falls $h(x) = y$ ist
- Ein Paar (x, x') mit $x \neq x'$ und $h(x) = h(x')$ 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)

Schlüssellose Hashfunktionen (MDCs)

- 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 bei vorgegebenem Hashwert y dagegen nicht möglich ist

Problem P1 (Bestimmung eines Urbilds)

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

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

- Falls es einen immensen Aufwand erfordert, bei gegebenem 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

- Für andere Anwendungen ist es dagegen wichtig, dass es für einen gegebenen Text x praktisch unmöglich ist, einen weiteren Text $x' \neq x$ mit dem gleichen Hashwert $h(x') = h(x)$ zu finden

Problem P2 (Bestimmung eines zweiten Urbilds)

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

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

- Falls Problem P2 einen immensen Aufwand erfordert, heißt h **schwach kollisionsresistent** (*weakly collision resistant bzw. second preimage resistant*)
- Diese Eigenschaft wird beim Integritätsschutz durch einen MDC benötigt

- Für bestimmte Anwendungen ist es sogar nötig, dass sich überhaupt kein Kollisionspaar finden lässt
- Diese Eigenschaft ist bspw. beim Einsatz von MDCs bei der Erstellung von digitalen Signaturen erforderlich

Problem P3 (Bestimmung einer Kollision)

Gegeben: Eine Hashfunktion $h: X \rightarrow Y$

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

- Falls Problem P3 einen immensen Aufwand erfordert, heißt h **(stark) kollisionsresistent** (*collision resistant*)

- Obwohl die schwache Kollisionsresistenz eine gewisse Ähnlichkeit mit der Einweg-Eigenschaft aufweist, sind diese beiden Eigenschaften im allgemeinen unvergleichbar:
 - Eine schwach kollisionsresistente Funktion muss 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

- Wir zeigen nun, dass stark kollisionsresistente Hashfunktionen sowohl schwach kollisionsresistent als auch Einweghashfunktionen sind
- Hierzu reduzieren wir das Kollisionsproblem P3 auf das Problem P1, ein Urbild zu bestimmen, bzw. auf das Problem P2, ein zweites Urbild zu bestimmen
- Für die Reduktion von P3 auf P1 müssen wir zusätzlich voraussetzen, dass h eine Kompressionsfunktion ist

Satz

Jede stark kollisionsresistente (n, m) -Hashfunktion $h: X \rightarrow Y$ ist auch schwach kollisionsresistent

Satz

Jede stark kollisionsresistente (n, m) -Kompressionsfunktion $h: X \rightarrow Y$ ist auch eine Einweg-Hashfunktion

Satz

Jede stark kollisionsresistente (n, m) -Hashfunktion $h: X \rightarrow Y$ ist auch schwach kollisionsresistent

Beweis.

- Falls h nicht schwach kollisionsresistent ist, ist das Problem P2, ein zweites Urbild zu bestimmen, nicht schwer
- Indem wir das Problem P3, ein Kollisionspaar für h zu finden, auf P2 reduzieren, folgt dann, dass auch P3 nicht schwer ist und somit h nicht stark kollisionsresistent sein kann
- Sei A ein Las-Vegas Algorithmus, der für ein zufällig gewähltes $x \in_R X$ mit Wahrscheinlichkeit ε ein zweites Urbild x' von $h(x)$ findet
- Dann findet folgender Algorithmus mit Wk ε ein Kollisionspaar für h :

 - 1 **guess randomly** $x \in_R X$
 - 2 $x' := A(x)$
 - 3 **if** $h(x') = h(x) \wedge x' \neq x$ **then output** (x, x')

Als nächstes reduzieren wir das Kollisionsproblem auf das Urbildproblem

Satz

Jede stark kollisionsresistente (n, m) -Kompressionsfunktion $h: X \rightarrow Y$ ist auch eine Einweg-Hashfunktion

Beweis.

- Für jedes $y \in h(X) = \{h(x) \mid x \in X\}$ sei $h^{-1}(y)$ die Urbildmenge $\{x \in X \mid h(x) = y\}$ von y
- Sei A ein Las Vegas Algorithmus, der für jedes $y \in h(X)$ mit Wahrscheinlichkeit ε ein Urbild x mit $h(x) = y$ findet

Beweis.

- Für jedes $y \in h(X) = \{h(x) \mid x \in X\}$ sei $h^{-1}(y)$ die Urbildmenge $\{x \in X \mid h(x) = y\}$ von y
- Sei A ein Las Vegas Algorithmus, der für jedes $y \in h(X)$ mit Wahrscheinlichkeit ε ein Urbild x mit $h(x) = y$ findet
- Betrachte folgenden Las-Vegas Algorithmus B :

 - 1 **guess randomly** $x \in X$
 - 2 $y := h(x)$; $x' := A(y)$
 - 3 **if** $x \neq x'$ **then return** (x, x') **else return** ?

- Sei $\mathcal{C} = \{h^{-1}(y) \mid y \in h(X)\}$ die Partition von X in die Urbildmengen
- Dann hat B eine Erfolgswahrscheinlichkeit von

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

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 $h(x) = y$ 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 kann ein Gegner nur dadurch Informationen über h erhalten, dass er für eine Reihe von Texten x_i die zugehörigen Hashwerte $h(x_i)$ berechnet (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 Gegner bereits die Hashwerte einer beliebigen Zahl von anderen Texten $x_i \neq x$ kennt

Proposition

- Sei $X_0 = \{x_1, \dots, x_k\}$ eine beliebige Menge von k verschiedenen Texten $x_i \in X$ 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 $(x, y) \in (X - X_0) \times Y$,

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

Das Zufallsorakelmodell (ZOM)

- Um eine obere Komplexitätsschranke für das Urbildproblem P1 im ZOM zu erhalten, betrachten wir folgenden Algorithmus:

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 ?
```

- Hierbei gibt der Parameter q die Anzahl der Hashwertberechnungen (also die Anzahl der gestellten Orakelfragen an das Zufallsorakel h) an
- Die Laufzeit des Algorithmus ist also proportional zu q

Satz

$\text{FindPreimage}(h, y, q)$ 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 obiger Proposition sind diese Ereignisse stochastisch unabhängig und ihre Wahrscheinlichkeit ist

$$\Pr[E_i] = 1/m \text{ für } i = 1, \dots, q$$

- Also folgt

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



Das Zufallsorakelmodell (ZOM)

- Ist q verglichen mit m relativ klein, so ist bei diesem Angriff

$$\begin{aligned}\varepsilon &= 1 - (1 - 1/m)^q \\ &= 1 - (1 - q/m + \binom{q}{2}/m^2 - \binom{q}{3}/m^3 + \dots \pm 1/m^q) \\ &\approx q/m\end{aligned}$$

- Um also auf eine Erfolgswahrscheinlichkeit von $1/2$ zu kommen, ist $q \approx m/2$ zu wählen
- Folgender Algorithmus liefert uns eine ähnliche obere Schranke für die Komplexität von Problem P2, ein zweites Urbild für $h(x)$ zu 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** ?
-

Satz

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

- Der Beweis ist analog zum Beweis des vorherigen Satzes
- Auch hier ist $q \approx m/2$ zu wählen, um auf eine Erfolgswahrscheinlichkeit von $1/2$ zu kommen

- Geht es lediglich darum, *irgendein* Kollisionspaar (x, x') aufzuspüren, so bietet sich ein sogenannter **Geburtstagsangriff** an
- Da hierfür deutlich weniger Hashwerte benötigt werden, lässt er sich entsprechend effizienter realisieren
- Dieser Angriff basiert auf dem sog. **Geburtstagsparadoxon**, welches in seiner einfachsten Form folgendes besagt

Geburtstagsparadoxon

Bereits in einer Klasse mit 23 Kindern ist die Wahrscheinlichkeit größer $1/2$, dass mindestens zwei Kinder am gleichen Tag Geburtstag haben

Der Geburtstagsangriff

- Der nächste Satz besagt, 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 mehrmals gezogen wird

- Für $m = 365$ und $q = 23$ ergibt dies einen Wert von ungefähr 0,507
- Da die Häufigkeiten der Geburtstage in einer Klasse nicht gleichverteilt sind, ist die Wahrscheinlichkeit, dass 2 Kinder am gleichen Tag Geburtstag haben, sogar noch etwas höher
- Zur Kollisionsbestimmung verwenden wir folgenden Algorithmus:

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** ?
-

Der Geburtstagsangriff

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** ?
-

- 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 Hash-tabelle der Größe q ein, so wird der Zeitaufwand für jeden einzelnen Text x im wesentlichen durch die Berechnung von $h(x)$ bestimmt

Satz

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 $X_0 = \{x_1, \dots, x_q\}$ und für $i = 1, \dots, q$ bezeichne E_i das Ereignis $h(x_i) \notin \{h(x_1), \dots, h(x_{i-1})\}$
- Dann ist $E_1 \cap \dots \cap E_q$ das Ereignis “Collision(h, q) gibt ? aus”
- Für $i = 1, \dots, q$ gilt nun

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

- Dies führt auf die Erfolgswahrscheinlichkeit

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

- Mit der Approximation $1 - x \approx e^{-x}$ erhalten wir folgende Abschätzung für ε :

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

- Für q erhalten wir daraus die Abschätzung

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

mit einer von ε abhängigen Konstante $c_\varepsilon = \sqrt{2\varepsilon}$

- Diese Abschätzung ist nur für ε -Werte nahe Null hinreichend genau

Der Geburtstagsangriff

- Aus der Abschätzung $\varepsilon \approx 1 - e^{-\frac{q^2}{2m}}$ für ε (siehe vorige Folie) erhalten wir insbesondere für größere Werte von ε eine bessere Abschätzung:

$$e^{\frac{q^2}{2m}} \approx \frac{1}{1-\varepsilon} \quad \text{bzw.} \quad q \approx c'_\varepsilon \sqrt{m} \quad \text{mit} \quad c'_\varepsilon = \sqrt{2 \ln \frac{1}{1-\varepsilon}}$$

- 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

- Im Folgenden 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, die auf Merkle und Damgaard zurückgeht

Iterierte Hashfunktionen

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

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

in einen String $y(x)$ mit der Eigenschaft $|y(x)| \equiv_t 0$

Processing: Sei $IV \in \{0, 1\}^m$ ein öffentlich bekannter

Initialisierungsvektor und sei $y(x) = y_1 \cdots y_r$ 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 Wert $\hat{h}(x) = g(z_r)$, wobei $g: \{0, 1\}^m \rightarrow \{0, 1\}^l$ eine öffentlich bekannte Funktion ist (meist wird für g die Identität verwendet)

Zur Berechnung von $\hat{h}(x)$ wird also die Funktion h genau r -mal aufgerufen

Wir formulieren nun eine für Preprocessing-Funktionen wünschenswerte Eigenschaft

Definition

- 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
- Mit anderen Worten: 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

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

Beweis.

- Wir nehmen an, dass es gelingt, ein Kollisionspaar $(x, \tilde{x})$ für $\hat{h}$ zu finden (d.h. $\hat{h}(x) = \hat{h}(\tilde{x})$ und $x \neq \tilde{x}$)
- Seien $y(x) = y_1 \dots y_r$ und $y(\tilde{x}) = \tilde{y}_1 \dots \tilde{y}_s$ mit $r \leq s$
- Da y suffixfrei ist, muss ein Index $i \in \{1, \dots, r\}$ mit $y_i \neq \tilde{y}_{s-r+i}$ existieren
- Weiter seien z_i ($i = 0, \dots, r$) und $\tilde{z}_j$ ($j = 0, \dots, s$) die in der Processing-Phase berechneten Hashwerte
- Da g injektiv ist, muss mit $g(z_r) = \hat{h}(x) = \hat{h}(\tilde{x}) = g(\tilde{z}_s)$ auch $z_r = \tilde{z}_s$ gelten

Beweis.

- Wir nehmen an, dass es gelingt, ein Kollisionspaar $(x, \tilde{x})$ für $\hat{h}$ zu finden (d.h. $\hat{h}(x) = \hat{h}(\tilde{x})$ und $x \neq \tilde{x}$)
- Seien $y(x) = y_1 \dots y_r$ und $y(\tilde{x}) = \tilde{y}_1 \dots \tilde{y}_s$ mit $r \leq s$
- Da y suffixfrei ist, muss eine Zahl $i \in \{0, \dots, r-1\}$ mit $y_{r-i} \neq \tilde{y}_{s-i}$ existieren
- Weiter seien $z_j = h(z_{j-1}y_j)$ ($j = 1, \dots, r$) und $\tilde{z}_k = h(\tilde{z}_{k-1}\tilde{y}_k)$ ($k = 1, \dots, s$) die in der Processing-Phase berechneten Hashwerte
- Da g injektiv ist, muss mit $g(z_r) = \hat{h}(x) = \hat{h}(\tilde{x}) = g(\tilde{z}_s)$ auch $z_r = \tilde{z}_s$ gelten
- Sei i_0 die kleinste Zahl $i \in \{0, \dots, r-1\}$ mit $z_{r-i_0-1}y_{r-i_0} \neq \tilde{z}_{s-i_0-1}\tilde{y}_{s-i_0}$
- Dann bilden $z_{r-i_0-1}y_{r-i_0}$ und $\tilde{z}_{s-i_0-1}\tilde{y}_{s-i_0}$ wegen

$$h(z_{r-i_0-1}y_{r-i_0}) = z_{r-i_0} = \tilde{z}_{s-i_0} = h(\tilde{z}_{s-i_0-1}\tilde{y}_{s-i_0})$$

ein Kollisionspaar für h



- Merkle und Damgaard schlugen 1989 folgende konkrete Realisierung ihrer Konstruktion vor
- Als Initialisierungsvektor 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 $y(x) = 0^t$
- Für $x \in \{0, 1\}^n$ mit $n > 0$ sei $r = \lceil \frac{n}{t-1} \rceil$ und $x = x_1 x_2 \dots x_{r-1} x_r$ mit $|x_1| = |x_2| = \dots = |x_{r-1}| = t-1$ sowie $|x_r| = t-1-d$, wobei $0 \leq d < t-1$
- Im Fall $r = 1$ ist dann $y(x) = y_1 y_2$ mit $y_1 = 0x0^d$ und $y_2 = 1\text{bin}_{t-1}(d)$
- Und für $r > 1$ ist $y(x) = y_1 \dots y_{r+1}$, wobei

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

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

Satz

Die durch (1) definierte Preprocessing-Funktion y ist suffixfrei

Beweis.

- Seien $x \neq \tilde{x}$ zwei Texte mit $|x| \leq |\tilde{x}|$
- Wir müssen zeigen, dass $y(x) = y_1 y_2 \dots y_{r+1}$ kein Suffix von $y(\tilde{x}) = \tilde{y}_1 \tilde{y}_2 \dots \tilde{y}_{s+1}$ 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 $y_{r+1} \neq \tilde{y}_{s+1}$
 2. Fall: $|x| = |\tilde{x}|$. In diesem Fall gilt $r = s$ und $|y(x)| = |y(\tilde{x})|$.
Wegen $x \neq \tilde{x}$ existiert ein Index $i \in \{1, \dots, r\}$ 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 $r < s$. Da $y(x)$ mit einer Null beginnt, aber das $(s - r + 1)$ -te Bit von $y(\tilde{x})$ eine Eins ist, kann $y(x)$ kein Suffix von $y(\tilde{x})$ sein

Nun betrachten wir den 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 $r + 1 = \lceil \frac{n}{t-1} \rceil + 1$ Aufrufe benötigt

Die MD4-Hashfunktion

- Die MD4-Hashfunktion (**M**essage **D**igest) wurde 1990 von Rivest vorgeschlagen
- Die Bitlänge von MD4 beträgt $l = 128$ Bit
- Bei einer Wortlänge von 32 Bit entspricht dies 4 Wörtern
- MD4 und die im Folgenden vorgestellten Hashfunktionen benutzen u.a. folgende Operationen auf Wörtern $X, Y \in \{0, 1\}^{32}$

Wort-Operationen	
$X \wedge Y$	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

- Die Ganzzahl-Addition wird bei MD4 und MD5 in **little endian** Architektur ausgeführt
- D.h. dass ein aus 4 Bytes, zusammengesetztes Wort $X = a_3a_2a_1a_0$, dessen Bytes $a_i \in 2^8$ die Zahlenwerte $(a_i)_2 \in [0, 255]$ haben, die Zahl $(a_0)_2 \cdot 2^{24} + (a_1)_2 \cdot 2^{16} + (a_2)_2 \cdot 2^8 + (a_3)_2$ repräsentiert
- Dagegen verwendet SHA-1 eine **big endian** Architektur
- D.h. dass $a_3a_2a_1a_0$ die Zahl $(a_3)_2 \cdot 2^{24} + (a_2)_2 \cdot 2^{16} + (a_1)_2 \cdot 2^8 + (a_0)_2$ repräsentiert
- Der MD4-Algorithmus benutzt die folgenden Funktionen f_j für $j = 0, \dots, 47$:

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

- Zudem benutzt er die folgenden Konstanten y_j, z_j, s_j für $j = 0, \dots, 47$:

	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

MD4(x)

```

1 input  $x \in \{0, 1\}^*, |x| = n$ 
2  $y := x10^k \text{bin}_{64}(n)$ ,  $k \in \{0, 1, \dots, 511\}$  mit
    $n + 1 + k + 64 \equiv 0 \pmod{512}$ 
3  $(H_1, H_2, H_3, H_4) := (67452301, \text{efcdab89}, 98\text{badcfe}, 10325476)$ 
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  $H_1 H_2 H_3 H_4$ 

```

- In Zeile 9 wird die **Kompressionsfunktion** von MD4 berechnet:

$$(A, B, C, D, X[z_j]) \mapsto (D, (A + f_j(B, C, D) + X[z_j] + y_j) \leftarrow s_j, B, C)$$

- In 1995 fand Dobbertin ein Kollisionspaar für die MD4 Hashfunktion mittels eines Angriffs, der sich in wenigen Sekunden ausführen lässt (2^{20} MD4-Berechnungen)
- Auf der Rump Session der Crypto 2004 berichtete Xiaoyun Wang, dass sich Kollisionen für MD4 mittels “hand calculation” (2^8 MD4-Berechnungen) finden lassen; veröffentlicht auf der EUROCRYPT 2005
- Im Jahr 2006 wurde ein Kollisionsangriff veröffentlicht, dessen Zeitaufwand weniger als 2 MD4-Berechnungen entspricht
- In 2008 veröffentlichte Gaëtan Leurent einen Angriff zur Berechnung eines Urbilds mit 2^{102} MD4-Berechnungen
- Dies wurde 2010 von Guo et. al. auf $2^{78,4}$ zur Urbildberechnung und $2^{69,4}$ MD4-Berechnungen zur Berechnung eines 2. Urbilds verbessert
- In 2011 wurde der MD4 als obsolet erklärt (RFC 6150)

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 Konstanten y_j , $0 \leq j \leq 63$, sind definiert als
$$y_j := \text{die ersten 32 Bit der Binärdarstellung von } \text{abs}(\sin(j + 1))$$

- Zudem benutzt der MD5 die folgenden Konstanten z_j und s_j :

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

MD5(x)

```
1 input  $x \in \{0, 1\}^*, |x| = n$ 
2  $y := x10^k \text{bin}_{64}(n)$ ,  $k \in \{0, 1, \dots, 511\}$  mit
    $n + 1 + k + 64 \equiv 0 \pmod{512}$ 
3  $(H_1, H_2, H_3, H_4) := (67452301, \text{efcdab89}, 98\text{badcfe}, 10325476)$ 
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  $H_1 H_2 H_3 H_4$ 
```

- In 1993, fanden den Boer und Bosselaers “Pseudo-Kollisionen” für die MD5 Kompressionsfunktion: zwei verschiedene Initialisierungskonstanten überführen einen Text in den denselben Hashwert
- In 1996 fand Dobbertin eine Kollision für die MD5 Kompressionsfkt.
- 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
- Wenige Tage später beschrieb Vlastimil Klima eine Methode, innerhalb von Stunden Kollisionen für MD5 auf einem Notebook zu berechnen
- Mittels der so genannten Tunneling-Methode wurde die Rechenzeit vom gleichen Autor im März 2006 auf eine Minute verkürzt

Die SHA1-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 für $j = 0, \dots, 79$:

	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 für $j = 0, \dots, 79$:

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

SHA-1(x)

```

1 input  $x \in \{0, 1\}^*, |x| = n$ 
2  $y := x10^k \text{bin}_{64}(n)$ ,  $k \in \{0, 1, \dots, 511\}$  mit
    $n + 1 + k + 64 \equiv 0 \pmod{512}$ 
3  $(H_0, \dots, H_4) := (67452301, \text{efcdab89}, 98\text{badcfe}, 10325476, \text{c3d2e1f0})$ 
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       $\text{temp} := (A \leftarrow 5) + f_j(B, C, D) + E + X[j] + K_j$ 
12       $(A, B, C, D, E) := (\text{temp}, A, B \leftarrow 30, C, D)$ 
13       $(H_0, \dots, H_4) := (H_0 + A, \dots, H_4 + E)$ 
14 output  $H_0 H_1 H_2 H_3 H_4$ 

```

Angriffe gegen SHA-0

- Auf der CRYPTO 98 stellten Chabaud und Joux einen Angriff gegen 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, deren Hashwerte sich nur an 18 von 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 2^{51} Hashwerte mit ca. 80 000 Stunden CPU-Rechenzeit auf einem 2-Prozessor 256-Itanium Supercomputer berechnet
- Der Aufwand wurde im Februar 2005 von Xiaoyun Wang, Yiqun Lisa Yin und Hongbo Yu auf 2^{39} Hashwertberechnungen reduziert

- 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
- Im Jahr 2005 stellten Xiaoyun Wang, Yiqun Lisa Yin und Hongbo Yu einen Kollisionsangriff gegen SHA-1 vor, der mit 2^{69} Hashwertberechnungen (anstelle von 2^{80} bei einem Geburtstagsangriff) auskommt
- Kurze Zeit später präsentierten Xiaoyun Wang, Andrew Yao und Frances Yao auf der Konferenz CRYPTO 2005 einen Kollisionsangriff gegen SHA-1 mit 2^{63} Hashwertberechnungen
- Auf der CRYPTO 2006 wurde ein Angriff gegen SHA-1 präsentiert, bei dem bis zu 25 Prozent der Bits in der gefälschten Nachricht frei wählbar sind
- In 2008 wurde von Stephane Manuel ein Kollisionsangriff mit einem geschätzten Aufwand von 2^{51} bis 2^{57} Berechnungen veröffentlicht

- Im Februar 2017 fanden die Google-Mitarbeiter Stevens, Bursztein, Karpman, Albertini und Markov die erste Kollision für SHA-1
- Diese bestand aus zwei verschiedenen PDF-Dateien
- Der Aufwand war jedoch enorm: eine einzelne CPU hätte dafür etwa 6500 Jahre benötigt
- Als Reaktion auf diese Angriffe wurde SHA-1 in 2011 vom NIST als veraltet erklärt und in 2013 wurde seine Verwendung für digitale Signaturen verboten

- 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, 90befffa, a4506ceb, bef9a3f7, c67178f2

- Dies sind jeweils die ersten 32 Bit der binären Nachkommastellen der dritten Wurzeln der ersten 64 Primzahlen $2, \dots, 311$

```
1 input  $x \in \{0, 1\}^*, |x| = n$ 
2  $y := x10^k \text{bin}_{64}(n)$ ,  $k \in \{0, \dots, 511\}$  mit  $n+1+k+64 \equiv 0 \pmod{512}$ 
3  $(H_0, \dots, H_7) := (6a09e667, \dots, 5be0cd19)$ 
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  $63$  do
8      $s0 := (X[t-15] \hookrightarrow 7) \oplus (X[t-15] \hookrightarrow 18) \oplus (X[t-15] \rightarrow 3)$ 
9      $s1 := (X[t-2] \hookrightarrow 17) \oplus (X[t-2] \hookrightarrow 19) \oplus (X[t-2] \rightarrow 10)$ 
10     $X[t] := X[t-16] + s0 + X[t-7] + s1$ 
11     $(A, B, C, D, E, F, G, H) := (H_0, H_1, H_2, H_3, H_4, H_5, H_6, H_7)$ 
12    for  $j := 0$  to  $63$  do  $\alpha$ 
13       $(H_0, H_1, \dots, H_7) := (H_0 + A, H_1 + B, \dots, H_6 + G, H_7 + H)$ 
14 output  $H_0 H_1 H_2 H_3 H_4 H_5 H_6 H_7$ 
```

Die Werte von $H_0, \dots, H_7$ in Zeile 3 sind die ersten 32 Bit der binären Nachkommastellen der Wurzeln der Primzahlen 2, 3, 5, 7, 11, 13, 17, 19

Programmstück α (SHA-256 Kompressionsfunktion)

1	$s0 := (A \hookrightarrow 2) \oplus (A \hookrightarrow 13) \oplus (A \hookrightarrow 22)$
2	$maj := (A \wedge B) \oplus (A \wedge C) \oplus (B \wedge C)$
3	$t2 := s0 + maj$
4	$s1 := (E \hookrightarrow 6) \oplus (E \hookrightarrow 11) \oplus (E \hookrightarrow 25)$
5	$ch := (E \wedge F) \oplus (\neg E \wedge G)$
6	$t1 := H + s1 + ch + K_j + X[j]$
7	$(A, B, C, D, E, F, G, H) := (t1 + t2, A, B, C, D + t1, E, F, G)$

- 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 ausgeschriebenen Wettbewerbs für den SHA3-Algorithmus ausgewählt und am 5. August 2015 als Alternative zu SHA-2 standardisiert
- 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, SHA-0 und SHA-1, die einen ähnlichen Aufbau wie SHA-2 haben, eine auf einem vollkommen anderen Entwurfsprinzip basierende Alternative zur Verfügung zu stellen

Die Sponge-Konstruktion

- Die Konstruktionsidee hinter dem SHA3-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 bereits 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
- Mittels der Funktion f wird eine Folge von **Zuständen** $s_i \in \{0, 1\}^b$ generiert, deren Länge b als **Weite** des Sponges bezeichnet wird
- Die Zustände $s_i = z_i u_i$ werden in zwei Teilblöcke z_i und u_i der Länge r bzw. c zerlegt, 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 $r = |z_i|$ der Bits des äußeren Zustands heißt **Bitrate**
- Die Anzahl $c = |u_i|$ der Bits des inneren Zustands wird als **Kapazität** des Sponges bezeichnet und ist sein wichtigster Sicherheitsparameter
- Da $r + c = b$ gelten muss, lässt sich über die Wahl der Parameter r und c ein Trade-off zwischen der Effizienz und der Sicherheit des Sponges realisieren

Die Sponge-Konstruktion

- Die Funktion f bildet den Kern der Sponge-Konstruktion
- Bevor f iteriert angewendet wird, um die Zustandsfolge s_i zu generieren, findet ein Preprocessing statt
- Wir definieren nun mögliche Anforderungen an die Preprocessing-Funktion y

Definition

- Sei y eine Funktion der Form $y: \{0, 1\}^* \rightarrow \{0, 1\}^*$
- y heißt **Paddingfunktion**, falls für alle $x \in \{0, 1\}^*$ ein $z \in \{0, 1\}^*$ mit $y(x) = xz$ existiert
- Eine solche Funktion heißt **Paddingfunktion für eine Bitrate $r \geq 1$** , falls $|y(x)| \equiv_r 0$ für alle $x \in \{0, 1\}^*$ gilt
- y heißt **sponge-konform für eine Bitrate $r \geq 1$** , falls gilt:
 - $\forall i = 0, \dots, r-1 \exists z \in \{0, 1\}^* \forall x \in \{0, 1\}^*$ mit $|x| \equiv_r i : y(x) = xz$
 - $\forall k \geq 0 \forall x \neq x' : y(x) \neq y(x')0^{kr}$

Beispiel

- Es ist leicht zu sehen, dass die Funktion

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

eine sponge-konforme Paddingfunktion für die Bitrate r ist

- Ohne die 1 am Ende von $\text{pad10}^*1_r(x) = x10^d1$ wäre dies nicht der Fall

Definition

- Sei y eine Paddingfunktion für $r \geq 1$ und sei $f : \{0, 1\}^b \rightarrow \{0, 1\}^b$
- Für $x \in \{0, 1\}^*$ sei $y(x) = y_1 \dots y_k$ mit $|y_i| = r$ für $i = 1, \dots, k$
- Sei $s_0 = 0^b$ und für $i \geq 1$ sei $s_i = z_i u_i$ mit $z_i \in \{0, 1\}^r$, $u_i \in \{0, 1\}^c$ und

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

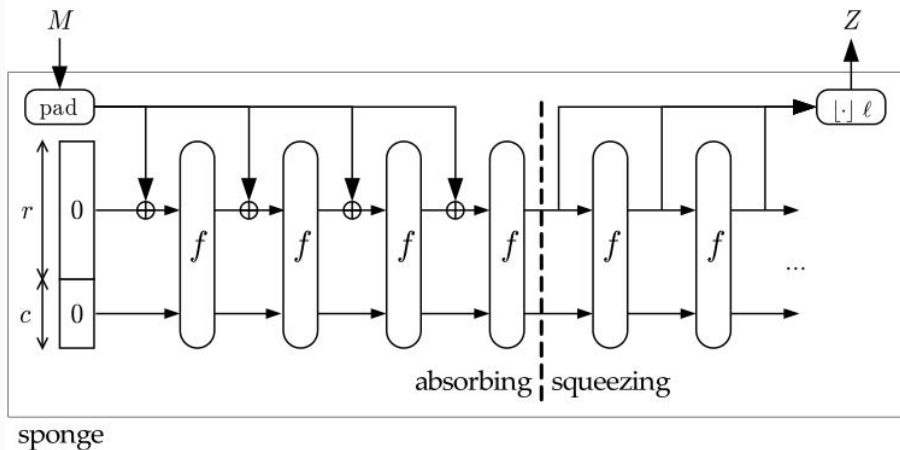
- Dann ist $\text{Sponge}_{f,y,r} : \mathbb{N} \times \{0, 1\}^* \rightarrow \{0, 1\}^*$ die Funktion mit

$$\text{Sponge}_{f,y,r}(l, x) = z_k \dots z_{k+m-1} z'_{k+m}$$

wobei $m = \lfloor \frac{l}{r} \rfloor$ und z'_{k+m} das Präfix von z_{k+m} der Länge $l - mr$ ist

- Für die Analyse definieren wir noch die beiden Funktionen

$$\text{Absorb}_{f,r}(y_1 \dots y_k) = s_k \quad \text{und} \quad \text{Squeeze}_{f,r}(l, s_k) = z_k \dots z_{k+m-1} z'_{k+m}$$



Die Sponge-Konstruktion

- Den Aufwand, ein Paar $x \neq x'$ mit $Sponge_{f,y,r}(l, x) = Sponge_{f,y,r}(l, x')$ zu finden, können wir nach oben durch den Aufwand abschätzen, ein Paar $y \neq y'$ in $\bigcup_{k \geq 1} \{0, 1\}^{kr}$ mit $y = y_1 \dots y_k$, $y' = y'_1 \dots y'_{k'}$ und $Absorb_{f,r}(y) = s_k = s'_{k'} = Absorb_{f,r}(y')$ zu finden
- Es genügt sogar, dass die inneren Zustände u_k von $s_k = z_k u_k$ und $u'_{k'}$ von $s'_{k'} = z'_{k'} u'_{k'}$ gleich sind, d.h. $y \neq y'$ ist ein **inneres Kollisionspaar** für die Funktion $Absorb_{f,r}$
- Setzen wir nämlich y_{k+1} und $y'_{k'+1}$ auf die äußeren Zustände z_k von s_k und $z'_{k'}$ von $s'_{k'}$, so folgt für die Texte $x = yy_{k+1}$ und $x' = y'y'_{k'+1}$:

$$\begin{aligned} Absorb_{f,r}(x) &= f(s_k \oplus (y_{k+1} 0^c)) = f(0^r u_k) = f(0^r u'_{k'}) \\ &= f(s'_{k'} \oplus (y'_{k'+1} 0^c)) = Absorb_{f,r}(x') \end{aligned}$$

- Falls das Suffix z von $y(x) = xz$ nur von $|x| \bmod r$ abhängt, gilt wegen $|x| \equiv_r |x'| \equiv_r 0$ dann auch $y(x') = x'z$ und somit folgt

$$Absorb_{f,r}(y(x)) = Absorb_{f,r}(xz) = Absorb_{f,r}(x'z) = Absorb_{f,r}(y(x'))$$

was $Sponge_{f,y,r}(l, x) = Sponge_{f,y,r}(l, x')$ für alle $l \geq 1$ impliziert

- Um eine innere Kollision $y \neq y'$ zu finden, hilft es, sich die 2^c inneren Zustände $u \in \{0, 1\}^c$ als Knoten eines gerichteten Multigraphen G vorzustellen, der für jeden Zustand $s = wv$ eine Kante $v \xrightarrow{w,z} u$ von v nach u mit dem Label w, z enthält, wobei $f(wv) = zu$ ist
- Dann lässt sich jedem String $y = y_1 \dots y_k \in \{0, 1\}^{kr}$ ein Pfad

$$P_y : u_0 \xrightarrow{y_1, z_1} u_1 \xrightarrow{z_1 \oplus y_2, z_2} u_2 \xrightarrow{z_2 \oplus y_3, z_3} u_3 \dots \xrightarrow{z_{k-1} \oplus y_k, z_k} u_k$$

der Länge k in G mit dem Startknoten $u_0 = 0^c$ zuordnen

- Umgekehrt ex. zu jedem Pfad $P : u_0 \xrightarrow{w_1, z_1} u_1 \xrightarrow{w_2, z_2} u_2 \xrightarrow{w_3, z_3} \dots \xrightarrow{w_k, z_k} u_k$ der Länge k in G mit Startknoten $u_0 = 0^c$ genau ein y mit $P_y = P$, nämlich $y = y_1 \dots y_k$ mit $y_1 = w_1$ und $y_i = z_{i-1} \oplus w_i$ für $i = 2, \dots, k$
- Daher ist $y \neq y'$ genau dann eine innere Kollision für $\text{Absorb}_{f,r}$, wenn P_y und $P_{y'}$ zwei verschiedene Pfade vom Startknoten 0^c zum gleichen Zielknoten $u_k = u'_k$ sind

- Ist f eine Zufallsfunktion (ZOM), so lassen bereits bekannte Werte von f keine Rückschlüsse auf die Werte von f für andere Argumente zu
- Anders als beim ZOM für eine Hashfunktion kann es sich dennoch für den Gegner lohnen, die Argumente von f adaptiv nach einer Strategie S zu wählen, um möglichst schnell eine innere Kollision $y \neq y'$ mit bestimmten Eigenschaften zu finden
- Der Algorithmus **InnerCollision** auf der nächsten Folie fasst dieses Vorgehen zusammen
- Die benutzte Strategie S wird hierbei in Form einer Routine als Teil der Eingabe übergeben und der Parameter q beschränkt die Anzahl der Aufrufe der Transformation f

Bestimmung eines inneren Kollisionspaares

Prozedur InnerCollision($f, r, q, \mathcal{S}$) mit $f : \{0, 1\}^b \rightarrow \{0, 1\}^b$

```

1   $c := b - r$ 
2  initialisiere den Multi-Digraphen  $G = (V, E) := (\{0, 1\}^c, \emptyset)$ 
3  for  $i := 1$  to  $q$  do
4    wähle  $v \in V$  und  $w \in \{0, 1\}^r$  nach Strategie  $\mathcal{S}$ 
5     $E := E \cup \{v \xrightarrow{w, z} u\}$ , wobei  $f(wv) = zu$  ist
6  if  $\exists$  2 Pfade  $0^c \xrightarrow{w_1, z_1} u_1 \xrightarrow{w_2, z_2} u_2 \xrightarrow{w_3, z_3} \dots \xrightarrow{w_k, z_k} u_k$  und
7       $0^c \xrightarrow{w'_1, z'_1} u'_1 \xrightarrow{w'_2, z'_2} u'_2 \xrightarrow{w'_3, z'_3} \dots \xrightarrow{w'_{k'}, z'_{k'}} u'_{k'}$  mit  $u_k = u'_{k'}$  then
8     $y := y_1 \dots y_k$  mit  $y_1 = w_1$  und  $y_i = z_{i-1} \oplus w_i$  für  $i = 2, \dots, k$ 
9     $y' := y'_1 \dots y'_{k'}$  mit  $y'_1 = w'_1$  und  $y'_i = z'_{i-1} \oplus w'_i$  für  $i = 2, \dots, k'$ 
10   return( $y, y'$ )
11 else
12   return(?)

```

Satz

- Für jede Strategie $\mathcal{S}$ gibt $\text{InnerCollision}(f, r, q, \mathcal{S})$ im ZOM mit Erfolgswahrscheinlichkeit höchstens

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

ein inneres Kollisionspaar (y, y') für die Funktion $\text{Absorb}_{f,r}$ aus

- Wählt $\mathcal{S}$ nur von 0^c aus erreichbare Knoten v und kein Argument wv mehrmals, so ist die Erfolgswahrscheinlichkeit exakt ε

Die Sponge-Konstruktion

Beweis.

- Sei E_i das Ereignis “ G enthält nach dem i -ten Durchlauf der for-Schleife zwei verschiedene Pfade von 0^c zu einem Knoten u ”
- 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

$$\Pr[E_i | \bar{E}_1 \cap \dots \cap \bar{E}_{i-1}] \leq \frac{i}{2^c}$$

- Wählt $\mathcal{S}$ nur erreichbare Knoten v und kein Argument wv mehrfach, so impliziert das Ereignis $\bar{E}_1 \cap \dots \cap \bar{E}_{i-1}$, dass nach dem i -ten Durchlauf **genau** i Knoten erreichbar sind (sonst gäbe es bereits zwei Pfade von 0^c zu einem Knoten in G) und es gilt sogar Gleichheit
 - Daher ist die Erfolgswahrscheinlichkeit von $\text{InnerCollision}(f, r, q, \mathcal{S})$
- $$1 - \Pr[\bar{E}_1 \cap \dots \cap \bar{E}_q] = 1 - \Pr[\bar{E}_1] \Pr[\bar{E}_2 | \bar{E}_1] \cdots \Pr[\bar{E}_q | \bar{E}_1 \cap \dots \cap \bar{E}_{q-1}]$$

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

- Mit der Approximation $1 - x \approx e^{-x}$ erhalten wir die Abschätzung

$$\begin{aligned}\varepsilon &= 1 - \prod_{i=1}^q \left(1 - \frac{i}{2^c}\right) \approx 1 - \prod_{i=1}^q e^{-\frac{i}{2^c}} = 1 - e^{-\frac{1}{2^c} \sum_{i=1}^q i} \\ &= 1 - e^{-\frac{q(q+1)}{2 \cdot 2^c}} \approx 1 - e^{-\frac{q^2}{2 \cdot 2^c}} \approx q^2 / 2 \cdot 2^c\end{aligned}$$

- Für q ergibt sich daraus die Abschätzung

$$q \approx c_\varepsilon \sqrt{2^c}$$

mit einer von ε abhängigen Konstanten $c_\varepsilon = \sqrt{2\varepsilon}$

- Da im Fall $l < c$ ein Geburtstagsangriff gegen die Funktion $x \mapsto \text{Sponge}_{f,y,r}(l, x)$ effizienter ist, erhalten wir die obere Schranke $\min(2^{c/2}, 2^{l/2})$ für den Aufwand eines mit Wahrscheinlichkeit $1/2$ erfolgreichen Angriffs gegen die Funktion $\text{Sponge}_{f,y,r}$ im ZOM

- Um geeignete Kandidaten für SHA-3 zu finden, veranstaltete das NIST ähnlich wie beim AES von Nov. 2007 bis Okt. 2012 einen Wettbewerb
- Die Minimalanforderungen an die Hashfunktionen waren, dass sie Nachrichten bis zu einer Obergrenze von mindestens $2^{64} - 1$ Bit hashen und mindestens die Hash-Längen 224, 256, 384 und 512 Bit unterstützen
- Von den eingereichten 64 Hashfunktionen erfüllten 51 die Teilnahmebedingungen und wurden für Runde 1 akzeptiert
- In Runde 1 wurde die Sicherheit und Performanz dieser 51 Kandidaten von Kryptologen aus aller Welt in einem offenen Bewertungsprozess analysiert und 14 der 51 Kandidaten erreichten Runde 2
- Nach jeder Runde konnten Veränderungen an den Algorithmen vorgenommen werden, um auf die Ergebnisse der Analysen zu reagieren
- Im Fall von Keccak wurde die Rundenzahl der Permutationsfunktion von 18 auf 24 erhöht, um eine größere Sicherheitsreserve zu erreichen
- Zudem wurde die Bitrate bei einigen Varianten vergrößert, um die Effizienz zu verbessern

- Im Dezember 2010 wurden die fünf Finalisten BLAKE, Grøstl, JH, Keccak und Skein bekanntgegeben
- Der Gewinner Keccak wies eine deutlich bessere Hardware-Performanz als seine Konkurrenten auf
- Dagegen hatten BLAKE und Skein die beste Software-Performanz und waren in den meisten Fällen schneller als SHA-2
- Keccak erlaubt für die Sponge-Weite die Wahl folgender 7 Werte $b_w \in \{25 \cdot 2^w \mid w = 0, \dots, 6\} = \{25, 50, 100, 200, 400, 800, 1600\}$
- Entsprechend stehen 7 Permutationen $f_w : \{0, 1\}^{b_w} \mapsto \{0, 1\}^{b_w}$, $w = 0, \dots, 6$ zur Verfügung
- Im SHA3-Standard ist $w = 6$, also $b = b_6 = 1600$ und SHA-3 stellt für $l = 224, 256, 384, 512$ jeweils eine Hashfunktion SHA3- l mit der festen Hash-Länge l und der Kapazität $c = 2l$ zur Verfügung
- Zudem enthält SHA-3 die Hashfunktionen SHAKE128 und SHAKE256 mit flexibler Hash-Länge und den Kapazitäten $c = 256$ bzw. $c = 512$

- Die SHA3-Familie besteht aus folgenden Hashfunktionen, die sich in der Kombination von Bitrate r und Ausgabelänge l unterscheiden:

Name	Hashlänge l	Rate r	Kap. c	Padding
SHA3-224	224	1152	448	0110*1
SHA3-256	256	1088	512	- " -
SHA3-384	384	832	768	- " -
SHA3-512	512	576	1024	- " -
SHAKE128	variabel	1344	256	111110*1
SHAKE256	- " -	1088	512	- " -

- Für die Hash-Längen $l \in \{224, 256, 384, 512\}$ gilt also

$$\text{SHA3-}l(x) = \text{Sponge}_{f_6, \text{pad}10*1, r}(l, x01) \text{ mit } c = 2l \text{ und } r = 1600 - 2l$$

- Das zusätzliche Padding 01 soll dabei SHA3- l von anderen Anwendungen von Keccak mit denselben Werten b, l, r unterscheiden

- Wir betrachten nun die Funktion $f = f_w$ für $w = 6$ im Detail
- Die Funktion $f : \{0, 1\}^{5 \times 5 \times 64} \rightarrow \{0, 1\}^{5 \times 5 \times 64}$ bildet einen $(5 \times 5 \times 64)$ -Block X auf einen ebensolchen Block $X' = f(X)$ ab
- Hierzu wird $12 + 2w = 24$ -mal eine Rundenfunktion

$$f^r : \{0, 1\}^{5 \times 5 \times 64} \rightarrow \{0, 1\}^{5 \times 5 \times 64}, \quad r = 0, \dots, 23$$

aufgerufen, die unter Verwendung einer Rundenkonstanten RC_r den Block $X_{r+1} = f^r(X_r)$ für $r = 0, \dots, 23$ berechnet, wobei $X_0 = X$ und $X' = X_{24}$ ist

- Es gilt

$$f^r(X_r) = \iota_r(\chi(\pi(\rho(\theta(X_r))))) ,$$

wobei θ , ρ , π , χ und ι_r Permutationen auf der Menge $\{0, 1\}^{5 \times 5 \times 64}$ sind

- Zur Beschreibung der Permutationen θ , ρ , π , χ und ι_r bezeichnen wir die einzelnen Bits eines Blocks $Z \in \{0, 1\}^{5 \times 5 \times 64}$ mit $Z[x, y, z]$ für $x, y = 0, \dots, 4$ und $z = 0, \dots, 63$
- Die Funktion θ ist so gewählt, dass sich $\theta^{-1}(Z)$ an möglichst vielen Bits ändert, falls eines in Z geflippt wird
- Hierzu führt θ auf den **Wörtern** $Z[x, y] = Z[x, y, 0] \dots Z[x, y, 63]$ des Zustandsvektors Z folgende Operationen aus:
 - $S(x) := Z[x, 0] \oplus \dots \oplus Z[x, 4]$
 - $R(x) := S(x) \leftarrow 1$
 - $Z[x, y] := Z[x, y] \oplus S(x - 1) \oplus R(x + 1)$,
 wobei die Indexoperationen modulo 5 erfolgen
- Die Transposition ρ rotiert die Wörter von Z wie folgt:
 - $Z[x, y] := Z[x, y] \leftarrow w_{x,y} \bmod 64$,
 wobei $w_{0,0} = 0$ und $w_{x,y} = (t + 1)(t + 2)/2$ für $(x, y) = (0, 1) \begin{pmatrix} 3 & 1 \\ 2 & 0 \end{pmatrix}^t$ und $t = 0, \dots, 23$ gilt

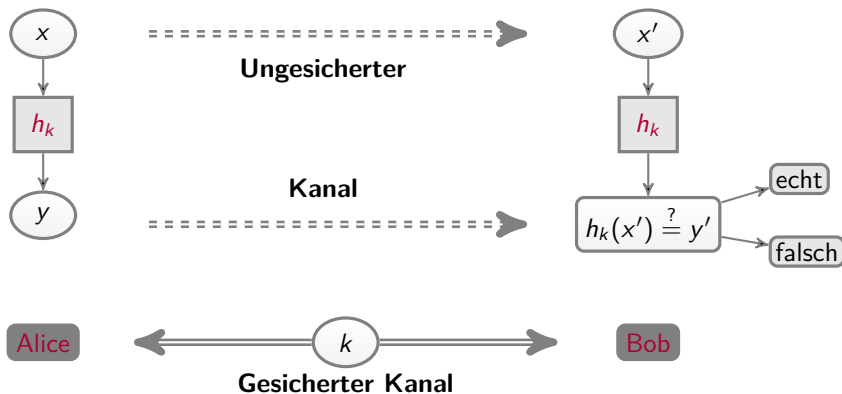
- Die Transposition π permutiert die Wörter von Z wie folgt:
 - $Z[3x + 2y, x] := Z[x, y]$
- Die Funktion χ ist die einzige nichtlineare Funktion und operiert innerhalb der 5-Bit-Spalten $Z[x, z]$ des Zustandsvektors, wobei jedes Bit nur von 2 anderen in dieser Spalte abhängt:
 - $Z[x, y] := Z[x, y] \oplus (\neg Z[x, y + 1] \wedge Z[x, y + 2])$
- Schlussendlich setzt ι_r das Wort $Z[0, 0]$ auf das Wort
 - $Z[0, 0] := Z[0, 0] \oplus RC_r,$

wobei das Bit an Position $2^m - 1$ in RC_r für $m = 0, \dots, 6$ das Bit an Position $7r + m$ eines LFSR mit dem erzeugenden Polynom $x^8 + x^6 + x^5 + x^4 + 1$ ist und alle übrigen Bits in RC_r Null sind

Definition

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**



- Hierbei ist k der symmetrische Authentifikationsschlüssel und $y = h_k(x)$ der **MAC**-Wert für x unter k
- Indem Alice ihre Nachricht x um den Hashwert $y = h_k(x)$ ergänzt, hat Bob nicht nur die Möglichkeit, anhand von y die empfangene Nachricht x' auf Manipulationen, sondern auch ihre Herkunft zu überprüfen

Damit ein geheimer Schlüssel k für die Authentifizierung mehrerer Nachrichten benutzt werden kann, ohne dass dies einem potentiellen Gegner 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 eines berechnungsresistenten MAC ist es einem Gegner praktisch nicht möglich, an Bob eine Nachricht x zu schicken, die Bob als von Alice stammend anerkennt
- Zu beachten ist allerdings, dass die Berechnungsresistenz nichts für den Fall aussagt, dass der Schlüssel k bekannt ist
- So kann nicht davon ausgegangen werden, dass die Funktion $h_k(x)$ bei bekanntem k die Einweg-Eigenschaft besitzt oder schwach (beziehungsweise stark) kollisionsresistent ist
- Es ist jedoch leicht zu sehen, dass es die Berechnungsresistenz erfordert, dass $h_k(x)$ bei geheimgehaltenem k zumindest schwach kollisionsresistent ist
- Dies ist etwa der Fall, wenn k im Speicher eines ausforschungssicheren Chips abgelegt wird

- Mithilfe eines berechnungsresistenten MAC 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

- Ein Angriff gegen einen MAC hat die unbefugte Berechnung von MAC-Werten zum Ziel
- Das heißt, der Gegner 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 Gegner kennt nur den benutzten MAC und versucht ein Paar (x, y) mit $h_k(x) = y$ zu generieren, wobei k der (dem Gegner unbekannte) Schlüssel ist

Substitution:

Der Gegner 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 Gegner 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 Gegner kann die Texte x_i selbst wählen

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

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

Auch wenn die Anwender den Schlüssel nach jeder MAC-Berechnung wechseln, sollte $\mathcal{H}$ zumindest einem **Substitutionsangriff** widerstehen, da der Gegner sonst ein Paar (x, y) abfangen und durch ein anderes ersetzen könnte

Modell: Schlüssel k und Nachrichten x werden unabhängig gemäß einer Wahrscheinlichkeitsverteilung $p(k, x) = p(k)p(x)$ generiert, welche dem Gegner bekannt ist

- Dabei nehmen wir an, dass $p(x) > 0$ und $p(k) > 0$ für alle $x \in X$ und $k \in K$ gilt
- Sei α die Wahrscheinlichkeit, mit der einem Gegner mit optimaler Strategie ein Impersonationsangriff gelingt
- Zudem sei β die Wahrscheinlichkeit, mit der einem Gegner mit optimaler Strategie ein Substitutionsangriff gelingt

Erfolgswahrscheinlichkeit für Impersonation

- Für ein Paar $(x, y) \in X \times Y$ sei $\alpha(x, y)$ die Wahrscheinlichkeit, dass dem Gegner mit der Wahl des Paares (x, y) eine Impersonation gelingt
- Dann ist $\alpha(x, y) = p(y|x) = p(x \mapsto y)$ die Wahrscheinlichkeit, dass der zufällig gewählte Schlüssel den Text x auf den MAC-Wert y abbildet:

$$\alpha(x, y) = \sum_{k \in K(x, y)} p(k)$$

wobei $K(x, y) = \{k \in K \mid h_k(x) = y\}$ alle Schlüssel enthält, die x auf y abbilden

- Folglich ist $\alpha(x) = \max\{\alpha(x, y) \mid y \in Y\}$ die maximale Wahrscheinlichkeit, mit der eine Impersonation gelingt, falls der Text x gewählt wird
- Also ist $\alpha = \max\{\alpha(x) \mid x \in X\}$

Beispiel

- Sei $K = \{1, 2, 3\}$, $X = \{a, b, c, d\}$ und $Y = \{0, 1\}$
- Wir beschreiben H durch die zugehörige Authentifikationsmatrix
- 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 $h_k(x)$:

		0,1	0,2	0,3	0,4
		<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>
0,25	1	0	0	0	1
0,30	2	1	1	0	1
0,45	3	0	1	1	0

- Die umrahmten Zahlen sind die Wahrscheinlichkeiten $p(x)$ bzw. $p(k)$

Beispiel (Fortsetzung)

		0,1	0,2	0,3	0,4
		<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>
0,25	1	0	0	0	1
0,30	2	1	1	0	1
0,45	3	0	1	1	0

- Dann hat der Gegner folgende Erfolgsaussichten $\alpha(x)$, falls er an *Bob* den Text x senden möchte:

x	<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>
$p(x \mapsto 0)$	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

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

Beweis.

- Für beliebiges $x \in X$ 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

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

Ein Maß für den Informationsgehalt

- In der Informationstheorie wird die Unsicherheit über eine Nachrichtenquelle X nach ihrer Entropie bemessen
- Dabei entspricht die Unsicherheit über X genau dem Informationsgewinn, der sich aus der Beobachtung der Quelle X ergibt
- Intuitiv ist die in einer einzelnen Nachricht x steckende Information umso größer, desto unwahrscheinlicher sie ist
- Tritt eine Nachricht x mit der Wahrscheinlichkeit $p(x) = \Pr[X=x] > 0$ auf, dann ist ihr Informationsgehalt definiert als

$$\text{Inf}_X(x) = \log_2(1/p(x)) = -\log_2 p(x)$$

- Im Fall $p(x) = 0$ sei $\text{Inf}_X(x) = 0$

Ein Maß für den Informationsgehalt

- Diese Definition des Informationsgehalts ergibt sich zwangsläufig aus den beiden folgenden Axiomen:
 - Der (gemeinsame) Informationsgehalt $\text{Inf}_{X,Y}(x,y)$ von zwei Nachrichten x und y , die aus unabhängigen Quellen X und Y stammen, ist $\text{Inf}_X(x) + \text{Inf}_Y(y)$
 - Eine Nachricht x , die mit Wahrscheinlichkeit $\Pr[X = x] = 1/2$ auftritt, hat den Informationsgehalt $\text{Inf}_X(x) = 1$
- Die Einheit des Informationsgehalts ist **bit** (basic information unit)
- Die Entropie von X ist nun der erwartete Informationsgehalt einer von X generierten Nachricht

Definition

- Sei X eine Zufallsvariable mit Wertebereich $W(X) = \{x_1, \dots, x_n\}$ und sei $p_i = \Pr[X = x_i]$
- Dann ist die **Entropie** von X definiert als

$$\mathcal{H}(X) = \sum_{i=1}^n p_i \ln f_X(x_i) = \sum_{i=1}^n p_i \log_2(1/p_i) = - \sum_{i=1}^n p_i \log_2(p_i)$$

Beispiel

- Sei X eine Zufallsvariable mit der Verteilung

x_i	sonnig	leicht bewölkt	bewölkt	stark bewölkt	Regen	Schnee	Nebel
p_i	1/4	1/4	1/8	1/8	1/8	1/16	1/16

- Dann ergibt sich die Entropie von X zu

$$\mathcal{H}(X) = 1/4 \cdot (2 + 2) + 1/8 \cdot (3 + 3 + 3) + 1/16 \cdot (4 + 4) = 2,625 \quad \triangleleft$$

- Die Entropie nimmt im Fall der Gleichverteilung $p_1 = \dots = p_n = 1/n$ den Wert $\log_2(n)$ an, während sie für jede andere Verteilung auf einer n -elementigen Menge einen Wert $\mathcal{H}(X) < \log_2(n)$ hat (siehe unten)
- Die Unsicherheit über eine Zufallsvariable X ist um so größer, je größer der Wertebereich und je gleichmäßiger die Verteilung von X ist
- Bringt X zum Beispiel nur einen einzigen Wert mit positiver Wahrscheinlichkeit hervor, dann (und nur dann) hat $\mathcal{H}(X)$ den Wert 0
- Für den Nachweis von oberen Schranken für die Entropie benutzen wir folgende Hilfsmittel aus der Analysis

Definition

- Sei $I \subseteq \mathbb{R}$ ein Intervall. Eine Funktion $f : I \rightarrow \mathbb{R}$ heißt **konkav** auf I , falls für alle $x \neq y \in I$ und $0 < t < 1$ gilt:

$$f(tx + (1-t)y) \geq tf(x) + (1-t)f(y)$$

- Gilt sogar „ $>$ “ anstelle von „ $\geq$ “, so heißt f **streng konkav** auf I
- Im Falle von „ $<$ “ bzw. „ $\leq$ “ heißt f **(streng) konvex** auf I

Beispiel

Die Funktion $f(x) = \log_2(x)$ ist streng konkav auf $(0, \infty)$



Für den Beweis des nächsten Satzes benötigen wir die **Jensensche Ungleichung**, die wir ohne Beweis angeben

Jensensche Ungleichung

- Sei f eine streng konkave Funktion auf I und seien $0 < a_1, \dots, a_n < 1$ reelle Zahlen mit $\sum_{i=1}^n a_i = 1$
- Dann gilt für alle $x_1, \dots, x_n \in I$,

$$f\left(\sum_{i=1}^n a_i x_i\right) \geq \sum_{i=1}^n a_i f(x_i)$$

- Im Falle einer streng konvexen Funktion f gilt $\leq$ anstelle von $\geq$
- Dabei gilt Gleichheit genau dann, wenn alle x_i den gleichen Wert haben

Satz

- Sei X eine Zufallsvariable mit Wertebereich $W(X) = \{x_1, \dots, x_n\}$ und Verteilung $p_i = \Pr[X=x_i]$ für $i = 1, \dots, n$
- Dann gilt $H(X) \leq \log_2(n)$, wobei Gleichheit genau im Fall $p_i = 1/n$ für $i = 1, \dots, n$ eintritt

Beweis.

- Aufgrund der Jensenschen Ungleichung gilt

$$\mathcal{H}(X) = \sum_{i=1}^n p_i \log_2(1/p_i) \leq \log_2 \sum_{i=1}^n (p_i/p_i) = \log_2 n,$$

wobei Gleichheit genau im Fall $1/p_1 = \dots = 1/p_n$ eintritt

- Letzteres ist mit der Bedingung $p_i = 1/n$ für $i = 1, \dots, n$ gleichbedeutend



Das folgende Lemma benötigen wir für den Beweis des nächsten Satzes

Lemma

- 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

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



Zudem benötigen wir noch den Begriff der bedingten Entropie

Definition. Seien X, Y Zufallsvariablen

Die **bedingte Entropie** von X unter Y ist definiert als

$$\mathcal{H}(X|Y) = \sum_{y \in W(Y)} p(y) \mathcal{H}(X|y) = \sum_y p(y) \sum_x p(x|y) \log_2(1/p(x|y)),$$

wobei $X|y$ die Zufallsvariable mit der Verteilung $p_y(x) = p(x|y)$ ist

Satz

- Für jeden MAC (X, Y, K, H) gilt:

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

- 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

Satz

Für jeden MAC (X, Y, K, H) gilt $\alpha \geq 1/2^{H(K)-H(K|X,Y)} \geq 1/l$

Beweis.

- Wegen $\alpha = \max_{x,y} \alpha(x,y)$ ist $E(\alpha(X,Y)) = \sum_{x,y} p(x,y) \alpha(x,y) \leq \alpha$
- Dabei ist $E(\alpha(X,Y))$ die Erfolgswahrscheinlichkeit eines (probabilistischen) Gegners, der das Paar (x,y) gemäß der Verteilung (X,Y) wählt
- Somit folgt unter Anwendung von obigem Lemma

$$\begin{aligned} \log \alpha &\geq \log E(\alpha(X,Y)) \geq E(\log \alpha(X,Y)) \\ &= \sum_{x,y} \underbrace{p(x,y)}_{p(x)p(y|x)} \underbrace{\log \alpha(x,y)}_{\log p(y|x)} \\ &= - \sum_x p(x) \sum_y p(y|x) \log \frac{1}{p(y|x)} = -H(Y|X) \end{aligned}$$

Beweis.

- Somit folgt unter Anwendung von obigem Lemma

$$\log \alpha \geq -H(\mathcal{Y} | \mathcal{X})$$

- Zudem gilt

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

und

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

- Daher folgt $H(\mathcal{Y} | \mathcal{X}) = H(\mathcal{K}) - H(\mathcal{K} | \mathcal{X}, \mathcal{Y})$ und somit $\log \alpha \geq H(\mathcal{K} | \mathcal{X}, \mathcal{Y}) - H(\mathcal{K})$
- Dies ist äquivalent zu $\alpha \geq 1/2^{H(\mathcal{K}) - H(\mathcal{K} | \mathcal{X}, \mathcal{Y})}$



Beispiel (Fortsetzung)

- Sei $K = \{1, 2, 3\}$, $X = \{a, b, c, d\}$ und $Y = \{0, 1\}$ mit

		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
	$p(x \mapsto 0)$	0,7	0,25	0,55	0,45
	$p(x \mapsto 1)$	0,3	0,75	0,45	0,55

- Dann gilt

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

Beispiel (Schluss)

- Um $H(\mathcal{K}|\mathcal{X}, \mathcal{Y})$ zu berechnen, benötigen wir die gemeinsame Verteilung $p(x, y) = p(x)p(y|x) = p(x)p(x \mapsto y)$ von $\mathcal{X}, \mathcal{Y}$ sowie die bedingten Verteilungen $\mathcal{K}_{x,y}$ für alle Paare $(x, y) \in X \times Y$:

(x, y)	$(a, 0)$	$(a, 1)$	$(b, 0)$	$(b, 1)$	$(c, 0)$	$(c, 1)$	$(d, 0)$	$(d, 1)$
$p(x \mapsto y)$	0,7	0,3	0,25	0,75	0,55	0,45	0,45	0,55
$p(x, y)$	0,07	0,03	0,05	0,15	0,165	0,135	0,18	0,22
$p(1 x, y)$	5/14	0	1	0	5/11	0	0	5/11
$p(2 x, y)$	0	1	0	2/5	6/11	0	0	6/11
$p(3 x, y)$	9/14	0	0	3/5	0	1	1	0
$H(\mathcal{K} x, y)$	$\approx 0,94$	0	0	$\approx 0,97$	$\approx 0,99$	0	0	$\approx 0,99$

- 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

- Als Nächstes schätzen wir die Wahrscheinlichkeit β ab, mit der ein Gegner mit optimaler Strategie eine von Alice gesendete Nachricht durch eine andere Nachricht ersetzen kann, ohne dass Bob dies bemerkt
- Dabei gehen wir davon aus, dass der Gegner keinen Einfluss auf die Wahl der von Alice gesendeten Nachricht x hat
- Falls der Gegner 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

- Hierbei ist $K(x, y, x', y') = \{k \in K \mid h_k(x) = y \wedge h_k(x') = y'\}$

Erfolgswahrscheinlichkeit für Substitution

- Im Fall, dass Alice das Paar (x, y) sendet, ist die maximale Erfolgswahrscheinlichkeit des Gegners

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

- Falls $p(x, y) = 0$ ist, setzen wir $\beta(x, y) = 0$
- Da wir davon ausgehen, dass der Gegner 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 erhalten wir die Berechnungsformel

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

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

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

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

Beispiel (Fortsetzung)

		<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>
0,25	1	0	0	0	1
0,30	2	1	1	0	1
0,45	3	0	1	1	0

(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 Strategien für den Gegner, ein von Alice gesendetes Paar (x,y) durch ein anderes Paar (x',y') zu ersetzen, sind fett gedruckt

Beispiel (Schluss)

		0,1	0,2	0,3	0,4
		<i>a</i>	<i>b</i>	<i>c</i>	<i>d</i>
0,25	1	0	0	0	1
0,30	2	1	1	0	1
0,45	3	0	1	1	0

- Für β erhalten wir somit den Wert

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

Erfolgswahrscheinlichkeit für Substitution

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

Satz

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

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 $p(x' \mapsto y' | x \mapsto y) \geq \frac{1}{m}$
- Dies impliziert

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

- Also 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

- 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

Beispiel (Fortsetzung)

- Ersetzt der Gegner ein Paar (x, y) durch ein Paar (x', y') mit $x' \neq x$, so wird dieses Paar von genau einem der drei infrage kommenden Schlüssel akzeptiert
- Dies liegt daran, dass in je 2 Spalten der Authentifikationsmatrix 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

In einem MAC (X, Y, K, H) mit $\beta = \frac{1}{m}$ gilt für alle Doppelpaare (x, y, x', y') mit $x \neq x'$ die Gleichheit

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

Beweis.

- Wir setzen zunächst voraus, dass $p(x \mapsto y) > 0$ für alle Paare $(x, y) \in X \times Y$ gilt
- Würde nun für ein Doppelpaar (x, y, x', y') mit $x \neq x'$

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

gelten, dann wäre auch

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

Beweis (Fortsetzung)

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

$$\beta = \sum_{u \in X, v \in Y} p(u, v) \beta(u, v) = p(x, y) \underbrace{\beta(x, y)}_{> 1/m} + \sum_{(u, v) \neq (x, y)} p(u, v) \underbrace{\beta(u, v)}_{\geq 1/m} > 1/m$$

folgen, 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, was wir bereits widerlegt haben

Beweis (Schluss)

- Es bleibt zu zeigen, dass $p(x \mapsto y) > 0$ für alle Paare $(x, y) \in X \times Y$ gilt
- Wäre $p(x \mapsto y) = 0$, so würde für ein beliebiges Paar (u, v) mit $p(u \mapsto v) > 0$ auch $p(x \mapsto y | u \mapsto v) = 0 < 1/m$ sein, was wir bereits widerlegt haben □

Satz

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

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

- Dies impliziert nun

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

und daher

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

Satz

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 (Schluss)

- 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 □

Es gilt also $\beta = \frac{1}{m}$ genau dann, wenn für alle (x, y, x', y') mit $x \neq x'$ gilt:

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

Sind die Schlüssel gleichverteilt, so bedeutet dies, dass in je zwei Spalten $x \neq x'$ der Authentifikationsmatrix jedes Paar (y, y') 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

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:

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

Ein MAC (X, Y, K, H) ist also genau dann 2-universal, wenn für alle Textpaare $x, x' \in X$ mit $x \neq x'$ jedes MAC-Wertpaar $y, y' \in Y$ mit Wk $1/m^2$ auftritt

Bemerkung

- 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 obiger Bemerkung ist klar, dass ein MAC bei gleichverteilten Schlüsseln genau dann die Bedingung $\beta = \frac{1}{m}$ erfüllt, wenn er 2-universal ist
- In diesem Fall nimmt 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

- 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 $K(x, y, x', y')$ für jedes Doppelpaar (x, y, x', y') mit $x \neq x'$ genau einen Schlüssel enthält
- Ein Schlüssel (a, b) gehört genau dann zu dieser Menge, wenn er die beiden Kongruenzen

$$ax + b \equiv_p y,$$

$$ax' + b \equiv_p y'$$

erfüllt

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

$$a = (y - y')(x - x')^{-1} \bmod p,$$

$$b = y - x(y - y')(x - x')^{-1} \bmod p$$

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



- Die Hashfunktionen des vorigen Satzes erfüllen wegen $n = m = p$ nicht die Kompressionseigenschaft
- Zwar lässt sich n noch geringfügig von p auf $p + 1$ vergrößern (und somit der Quotient m/n von 1 auf $\frac{p}{p+1}$ verkleinern), ohne K und Y zu verändern (siehe Übungen)
- Wie der nächste Satz zeigt, lässt sich eine stärkere Kompression mit dem Parameterwert $\lambda = 1$ jedoch nicht realisieren

Satz

Für einen $(n, m, l, 1)$ -MAC gilt

$$m \geq n - 1$$

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

Satz

Für einen $(n, m, l, 1)$ -MAC gilt

$$n \leq m + 1$$

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

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 Authentifikationsmatrix A wieder auf einen 2-universalen MAC führt
- Also können wir annehmen, dass die erste Zeile der Authentifikationsmatrix A nur Einsen enthält

Beweis (Schluss)

- 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



- Der nächste Satz liefert 2-universale MACs mit beliebig kleinem Kompressionsquotienten m/n
- Für den Beweis benötigen wir das folgende Lemma

Lemma

- 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



Satz

- Sei p prim und für $x = (x_1, \dots, x_d) \in \mathbb{Z}_2^d$ und $k = (k_1, \dots, k_d) \in \mathbb{Z}_p^d$ sei

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

- Dann ist (X, Y, K, H) mit $X = \mathbb{Z}_2^d - \{0^d\}$, $Y = \mathbb{Z}_p$ 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

Beweis (Fortsetzung)

- 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

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

- 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 □

Bemerkung

- 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 kleinen Kompressionsquotienten

$$\frac{m}{n} = \frac{p(p-1)}{p^d - 1} \approx p^{2-d}$$

bei einem λ -Wert von $\lambda = p^{d-2}$

- Wie der nächste Satz zeigt, lässt sich dies nicht mit einem kleineren λ -Wert (bzw. nicht mit einer kleineren Schlüssellänge) erreichen

Konstruktion von 2-universalen MACs

Für den Beweis des nächsten Satzes benötigen wir folgendes Lemma

Lemma

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

Beweis.

- Da die Funktion $x \mapsto x^2$ konvex ist, folgt mit der Jensenschen Ungleichung $(\sum b_i/m)^2 \leq \sum b_i^2/m$
- Folglich ist $(\sum b_i)^2 = m^2 \underbrace{(\sum b_i/m)^2}_{\leq \sum b_i^2/m} \leq m \sum b_i^2$

□

Jensensche Ungleichung

Für eine konkave Funktion f und $a_1, \dots, a_n \in (0, 1)$ mit $\sum_{i=1}^n a_i = 1$ ist

$$f\left(\sum_{i=1}^n a_i x_i\right) \geq \sum_{i=1}^n a_i f(x_i)$$

Satz

Für jeden (n, m, l, λ) -MAC gilt

$$l = \lambda m^2 \geq n(m-1) + 1 \quad \text{und somit} \quad m/n \geq \underbrace{(m-1)/m(\lambda - 1/m^2)}_{\approx 1/\lambda}$$

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 Authentifikationsmatrix 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

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

Beweis (Fortsetzung)

- 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

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

Beweis (Schluss)

- Somit erhalten wir

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

und daher folgt

$$\begin{aligned} ((\lambda - 1)n(n - 1) + n(\lambda m - 1))(\lambda m^2 - 1) &\geq (n(\lambda m - 1))^2 \\ \Rightarrow (\lambda n - n - \lambda + \lambda m)(\lambda m^2 - 1) &\geq n(\lambda m - 1)^2 \\ \Rightarrow -\lambda^2 m^2 + \lambda^2 m^3 &\geq \lambda n m^2 + \lambda n - \lambda + \lambda m - 2\lambda n m \\ \Rightarrow \lambda^2(m^3 - m^2) &\geq \lambda(n(m - 1)^2 + m - 1) \\ \Rightarrow \lambda m^2 &\geq n(m - 1) + 1 \end{aligned}$$



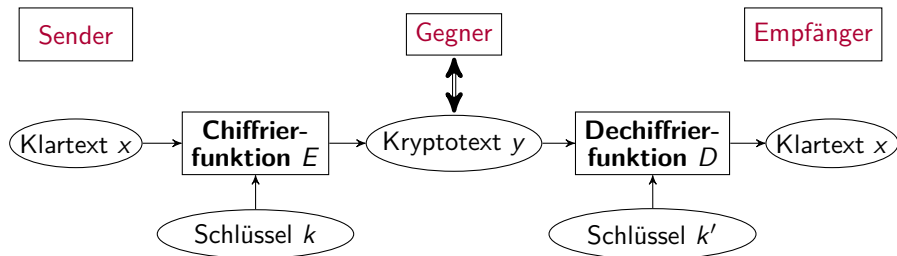
Definition

Ein **Kryptosystem** (M, C, E, D, K, S) wird durch folgende Komponenten beschrieben:

- $M \subseteq A^*$, der **Klartextraum** (message space), wobei A das **Klartextalphabet** ist
- $C \subseteq B^*$, der **Kryptotextraum** (ciphertext space), wobei B das **Kryptotextalphabet** ist
- K , der **Schlüsselraum** (key space)
- $E : K \times M \rightarrow C$, die **Verschlüsselungsfunktion** (encryption function)
- $D : K \times C \rightarrow M$, die **Entschlüsselungsfunktion** (decryption function)
- $S \subseteq K \times K$, eine Menge von Schlüsselpaaren (k, k') , so dass für jeden Klartext $x \in M$ folgende Beziehung gilt:

$$D(k', E(k, x)) = x \quad (*)$$

- Bei **symmetrischen Kryptosystemen** ist $S = \{(k, k) \mid k \in K\}$, weshalb wir in diesem Fall auf die Angabe von S verzichten können



- Zu jedem Schlüssel $k \in K$ korrespondiert eine
 - **Chiffrierfunktion** $E_k : x \mapsto E(k, x)$ und eine
 - **Dechiffrierfunktion** $D_k : y \mapsto D(k, y)$
- Die Gesamtheit dieser Abbildungen wird auch **Chiffre** (englisch **cipher**) genannt
- Daneben wird der Begriff „Chiffre“ auch als Bezeichnung für einzelne Kryptotextzeichen oder kleinere Kryptotextsequenzen verwendet

Lemma

Für jedes Paar $(k, k') \in S$ ist die Chiffrierfunktion E_k injektiv

Beweis

- Angenommen, für zwei Klartexte x_1 und x_2 gilt $E(k, x_1) = E(k, x_2)$
- Dann folgt

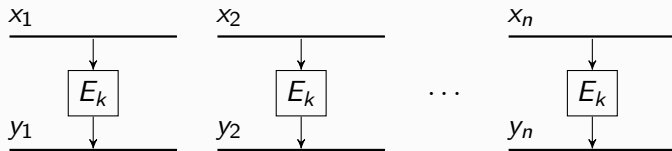
$$x_1 \stackrel{(*)}{=} D(k', \underbrace{E(k, x_1)}_{E(k, x_2)}) = D(k', E(k, x_2)) \stackrel{(*)}{=} x_2$$



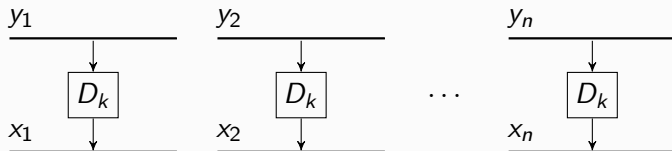
- Für den DES wurden vier verschiedene Betriebsarten vorgeschlagen, in denen grundsätzlich jede Blockchiffre E mit beliebiger Blocklänge ℓ betrieben werden kann
- Bei den ersten beiden Betriebsarten (ECB und CBC) werden Kryptotextblöcke der Länge ℓ übertragen
- Mit einer Blockchiffre kann aber auch ein **Stromsystem** realisiert werden, mit dem sich Kryptotextblöcke einer beliebigen Länge t , $1 \leq t \leq \ell$, übertragen lassen
- Der ECB-Modus ist die naheliegendste Betriebsart, wird aber in der Praxis so gut wie nie benutzt, da sie eine Reihe von Angriffsmöglichkeiten bietet

ECB-Modus (electronic codebook; elektron. Codebuch) ¹⁴⁶

- Die Binärnachricht x wird in Klartextblöcke x_i der Länge ℓ zerlegt
- Der letzte Block x_n wird ggf. nach einer vereinbarten Regel aufgefüllt
- Die Blöcke werden nacheinander mit demselben Schlüssel K einzeln verschlüsselt, übertragen und auf Empfängerseite wieder entschlüsselt:



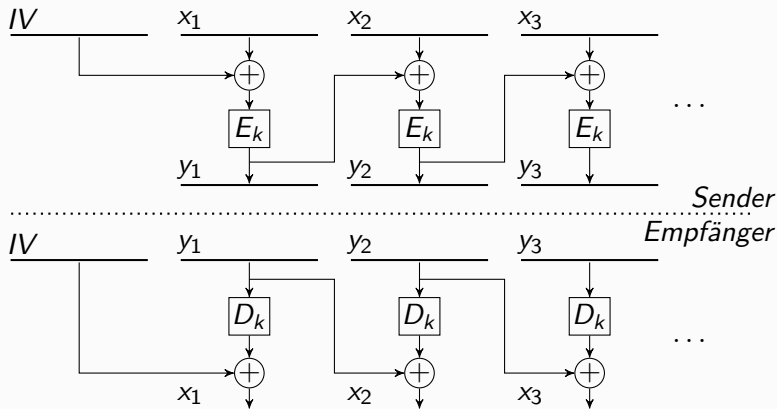
Sender



Empfänger

CBC-Modus (cipher block chaining)

- Um unbemerkte Manipulationen des Kryptotextes zu verhindern wird jeder Kryptotextblock nicht nur in Abhängigkeit des aktuellen Klartextblocks, sondern aller vorausgehenden Blöcken verschlüsselt
- Als Folge hiervon werden gleiche Klartextblöcke auf unterschiedliche Kryptotextblöcke abgebildet:



CBC-Modus (cipher block chaining)

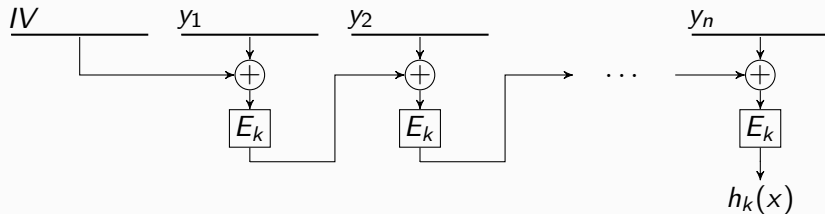
- Jeder Klartextblock x_i wird mit dem Kryptotextblock $E_k(x_{i-1})$ bitweise (modulo 2) addiert, bevor er verschlüsselt wird
- Zur Verschlüsselung von x_1 wird ein **Initialisierungsvektor** IV verwendet
- Dieser wird üblicherweise unverschlüsselt übertragen, sollte aber nicht mehrmals verwendet werden
- Er wird meist durch einen Pseudozufallsgenerator erzeugt
- Der CBC-Modus verhindert auch, dass sich bestimmte Klartextmuster direkt auf den Kryptotext übertragen (was im ECB-Modus der Fall ist)
- Ein extremes Beispiel ist die Verschlüsselung einer Graphikdatei x , deren Pixel durch ℓ -Bit Blöcke kodiert sind
- Zwar werden dann im ECB-Modus die Pixel substituiert, aber ansonsten stellt der Kryptotext y exakt dieselbe Graphik wie der Klartext x dar
- Dadurch wird eine „visuelle Entschlüsselung“ durch Betrachten der verschlüsselten Graphikdatei y ermöglicht

- Als Basis für die Konstruktion eines MAC kann auch ein symmetrisches Kryptosystem dienen
- Sei (M, C, E, D, K) ein symmetrisches Kryptosystem mit $M = C = \{0, 1\}^t$
- Zudem sei $IV := 0^t$ und sei $k \in K$ ein geheimer Schlüssel
- Sei y eine Funktion für den Preprocessing-Schritt, die für jeden Text $x \in \{0, 1\}^*$ einen nichtleeren Bitstring $y(x) \in \bigcup_{n \geq 1} \{0, 1\}^{tn}$ liefert
- Berechnung von $h_k(x)$:

```
1  $y := y(x) = y_1 \dots y_n, n \geq 1, y_i \in \{0, 1\}^t$   
2  $z_0 := IV$   
3 for  $i = 1$  to  $n$  do  
4    $z_i := E(k, z_{i-1} \oplus y_i)$   
5 output  $h_k(x) = z_n$ 
```

- Die MAC-Wert-Länge beträgt also t Bit

```
1  $y := y(x) = y_1 \dots y_n, n \geq 1, y_i \in \{0, 1\}^t$   
2  $z_0 := IV$   
3 for  $i = 1$  to  $n$  do  
4    $z_i := E_k(z_{i-1} \oplus y_i)$   
5 output  $h_k(x) = z_n$ 
```



- Wird auf den Preprocessing-Schritt verzichtet, so lässt sich leicht ein Angriff mit 2 adaptiven Fragen ausführen
- Um den MAC-Wert $h_k(x)$ für einen Text

$$x = x_1 \cdots x_n x_{n+1} \cdots x_{n+m}$$

mit $n, m \geq 1$ und $|x_i| = t$ für $i = 1, \dots, n + m$ zu ermitteln, erfragt der Gegner die MAC-Werte $z' = h_k(x')$ und $z'' = h_k(x'')$ für die Texte

$$x' = x_1 \cdots x_n \text{ und } x'' = (x_{n+1} \oplus IV \oplus z')x_{n+2} \cdots x_{n+m}$$

und erhält den gesuchten MAC-Wert $h_k(x) = z''$

- Diesen Angriff kann man zwar ausschließen, indem man die Anzahl n der Blöcke fixiert
- Allerdings würde dadurch die Anwendbarkeit des CBC-MAC erheblich eingeschränkt
- Der folgende Geburtstagsangriff ist aber auch bei fester Textlänge möglich

- Bei diesem Angriff kann der Gegner den Wert $h_k(x)$ für einen Text $x = x_1 \dots x_n \in \{0, 1\}^{tn}$ durch $q + 1$ MAC-Fragen x^i ermitteln, wobei er die letzten $n - 1$ Blöcke $x_2, \dots, x_n \in \{0, 1\}^t$ von x beliebig wählen kann
- Um die Fragen $x^i = x_1^i x_2^i x_3 \dots x_n \in \{0, 1\}^{tn}$, für $i = 1, \dots, q$ zu generieren, wählt der Gegner
 - q beliebige paarweise verschiedene Blöcke $x_1^1, \dots, x_1^q$ sowie
 - q zufällige Blöcke $x_2^1, \dots, x_2^q$ aus der Menge $\{0, 1\}^t$
- Dann erfragt er die MAC-Werte $z^i = h_k(x^i)$ für $i = 1, \dots, q$
- Die Texte $x^1, \dots, x^q$ sind paarweise verschieden, da sie sich im ersten Block unterscheiden

- Seien $z_1^i = E_k(IV \oplus x_1^i)$, $i \in [q]$, die im ersten Iterationsschritt des CBC-MAC berechneten Ausgabeblöcke
- Da die Blöcke x_2^i zufällig gewählt werden, sind die Eingaben $z_1^i \oplus x_2^i$ und Ausgaben $E_k(z_1^i \oplus x_2^i)$ im zweiten Iterationsschritt ebenfalls zufällig
- Ist $q \approx 1,17 \cdot 2^{t/2}$, so gilt also

$$\Pr[\exists i \neq j : E_k(z_1^i \oplus x_2^i) = E_k(z_1^j \oplus x_2^j)] = \Pr[\exists i \neq j : x_2^i = x_2^j] \approx 1/2$$

- Da ab dem dritten Iterationsschritt bei den q Texten $x^1, \dots, x^q$ die gleichen Textblöcke $x_3, \dots, x_n$ verarbeitet werden, sind die beiden Blöcke $z_1^i \oplus x_2^i$ und $z_1^j \oplus x_2^j$ genau dann gleich, wenn dies auf die MAC-Werte $z_n^i = h_k(x^i)$ und $z_n^j = h_k(x^j)$ zutrifft
- Daher kann der Gegner das Indexpaar (i, j) mit $z_1^i \oplus x_2^i = z_1^j \oplus x_2^j$ auch leicht finden, sofern es existiert (was wir im Folgenden annehmen)

- Wegen $x_1^i \neq x_1^j$ folgt $z_1^i = E_k(IV \oplus x_1^i) \neq E_k(IV \oplus x_1^j) = z_1^j$
- Daher impliziert die Kollision $z_1^i \oplus x_2^i = z_1^j \oplus x_2^j$, dass $x_2^i \neq x_2^j$ ist
- O.B.d.A. gelte $x_2^j \neq x_2$ (sonst vertauschen wir die Indizes i und j)
- Ersetzen wir nun die Blöcke x_2^i und x_2^j der beiden Texte x^i und x^j durch die Blöcke

$$\tilde{x}_2^i = x_2^i \oplus u \text{ und } \tilde{x}_2^j = x_2^j \oplus u \text{ für } u = x_2^j \oplus x_2 (\neq 0^t),$$

so liefert h_k für die resultierenden Texte

$$\tilde{x}^i = x_1^i \tilde{x}_2^i x_3 \cdots x_n \text{ und } \tilde{x}^j = x_1^j \tilde{x}_2^j x_3 \cdots x_n = x_1^j x_2 x_3 \cdots x_n$$

wegen $z_1^i \oplus \tilde{x}_2^i = z_1^i \oplus x_2^i \oplus u = z_1^j \oplus x_2^j \oplus u = z_1^j \oplus \tilde{x}_2^j$ denselben MAC-Wert $h_k(\tilde{x}^i) = h_k(\tilde{x}^j)$

- Daher kann der Gegner den MAC-Wert von $x^{q+1} = \tilde{x}^i (\neq x^i)$ erfragen, um den MAC-Wert von $\tilde{x}^j \notin \{x^1, \dots, x^{q+1}\}$ zu ermitteln

Definition

- Sei $0 \leq \varepsilon \leq 1$ und sei $q \in \mathbb{N}$
- Ein (ε, q) -Fälscher für einen MAC $\mathcal{H}$ ist ein probabilistischer Algorithmus $\mathcal{A}$, der q Fragen $x_1, \dots, x_q$ stellt und aus den Antworten $z_i = h_k(x_i)$ mit Wahrscheinlichkeit mindestens ε (bei zufällig gewähltem Schlüssel k) ein Paar (x, z) mit $h_k(x) = z$ und $x \notin \{x_1, \dots, x_q\}$ ausgibt
- Wir unterscheiden zwischen adaptiven und nicht-adaptiven Fragen
- Im ersten Fall darf der Text x_i von den MAC-Werten der Texte $x_1, \dots, x_{i-1}$ abhängen, im zweiten nicht
- Zudem unterscheiden wir zwischen selektiven und existentiellen Fälschungen
- Im ersten Fall ist der Gegner in der Lage, den MAC-Wert für einen beliebig vorgegebenen Text x zu finden, im zweiten Fall gelingt ihm dies nur für bestimmte Texte

Beispiel

- Der oben beschriebene Geburtstagsangriff auf einen CBC-MAC führt auf einen adaptiven existentiellen $(\frac{1}{2}, q+1)$ -Fälscher für $q \approx 1,17 \cdot 2^{\frac{t}{2}}$
- Dabei ist nur die letzte MAC-Frage x^{q+1} adaptiv und der Text x kann abgesehen vom ersten t -Bitblock beliebig vorgegeben werden ◀
- Eine Variante dieses Angriffs ist auch bei Verwendung einer Preprocessing-Funktion möglich
- Meist wird hierzu die Funktion $y : x \mapsto y(x) = y_0 \dots y_n$ mit $y_0 = \text{bin}_t(|x|)$ und $y_1 \dots y_n = x 0^{nt-|x|}$ verwendet, wobei $n = \lceil |x|/t \rceil$ ist
- Der erste Block $y_0 = \text{bin}_t(|x|)$ kodiert also die Länge von x als Binärzahl, die mit führenden Nullen auf die Länge t erweitert wird, und der letzte Block wird ebenfalls mit Nullen auf die Länge t aufgefüllt

Falls der Textraum eines MAC den Werteraum eines anderen MAC enthält, lassen sich diese leicht komponieren (Nested-MAC oder NMAC)

Definition

- Seien $\mathcal{H}_1 = (X, Y, K_1, F)$ mit $F = \{f_k \mid k \in K_1\}$ und $\mathcal{H}_2 = (Y, Z, K_2, G)$ mit $G = \{g_k \mid k \in K_2\}$ MACs
- Dann ist $\mathcal{H}_1 \circ \mathcal{H}_2 = (X, Z, K, H)$ die Komposition von $\mathcal{H}_1$ und $\mathcal{H}_2$, wobei $K = K_1 \times K_2$ und $H = \{g_{k_2} \circ f_{k_1} \mid (k_1, k_2) \in K\}$ ist

Beispiel

Wählt man für $\mathcal{H}_2$ einen MAC mit fester Textlänge und für $\mathcal{H}_1$ eine (schlüssellose) Hashfunktion (etwa SHA-1), so erhält man einen so genannten HMAC (Hash-MAC)

- Eine Variante hiervon ist der auf SHA-1 basierende HMAC
- Hierbei werden zwei Varianten von SHA-1 mit symmetrischen Schlüsseln komponiert, wobei jedoch beidesmal derselbe Schlüssel benutzt wird
- Seien

$$ipad = \underbrace{36 \dots 36}_{64mal} \text{ und } opad = \underbrace{5C \dots 5C}_{64mal}$$

512 Bit Konstanten

- Dann berechnet sich HMAC wie folgt:

$$\text{HMAC}_k(x) = \text{SHA-1}((k \oplus opad)\text{SHA-1}((k \oplus ipad)x))$$

- Hierbei fungiert die Funktion $f_k(x) = \text{SHA-1}((k \oplus ipad)x)$ als Hashfunktion mit Schlüssel, die beliebig lange Texte hasht, und der MAC $g_k(y) = \text{SHA-1}((k \oplus opad)y)$ wird nur auf Bitstrings der Länge $512+160=672$ angewendet

Wie der nächste Satz zeigt, genügt es, wenn f_k kollisionsresistent und g_k berechnungsresistent ist, um einen berechnungsresistenten HMAC mit beliebiger Eingabelänge zu erhalten

Definition

- Ein (ε, q) -Kollisionsangreifer für einen MAC $\mathcal{H} = (X, Y, K, H)$ ist ein probabilistischer Algorithmus $\mathcal{A}$, der q Fragen $x_1, \dots, x_q$ stellt und aus den Antworten $y_i = h_k(x_i)$ mit Wahrscheinlichkeit mindestens ε ein Paar (x, x') mit $h_k(x) = h_k(x')$ berechnet
- Hierbei ist k der dem Gegner unbekannte (und zufällig gewählte) Schlüssel
- Da der Gegner den Schlüssel k nicht kennt, ist ein Kollisionsangriff gegen einen MAC $\mathcal{H}$ meist schwieriger zu realisieren als ein Kollisionsangriff gegen eine schlüssellose Hashfunktion
- Andererseits bringt die Kenntnis des Schlüssels bei einem Geburtstagsangriff keinen Vorteil

Satz. Seien $\mathcal{H}_1 = (X, Y, K_1, F)$ und $\mathcal{H}_2 = (Y, Z, K_2, G)$ MACs

Falls

- für $\mathcal{H}_1$ kein adaptiver $(\varepsilon_1, q + 1)$ -Kollisionsangreifer und
- für $\mathcal{H}_2$ kein adaptiver (ε_2, q) -Fälscher existieren,

dann existiert für $\mathcal{H} = \mathcal{H}_1 \circ \mathcal{H}_2$ kein adaptiver $(\varepsilon_1 + \varepsilon_2, q)$ -Fälscher

Beweis.

- Sei A ein adaptiver (ε, q) -Fälscher für $\mathcal{H}$
- Seien $z_i = g_{k_2}(f_{k_1}(x_i))$ die Antworten, die A auf seine Fragen $x_1, \dots, x_q$ an den MAC $\mathcal{H}$ erhält, und sei (x, z) die Ausgabe von A
- Wir müssen zeigen, dass die Erfolgswk von A

$$\Pr[\underbrace{x \notin \{x_1, \dots, x_q\}}_B \wedge \underbrace{g_{k_2}(f_{k_1}(x)) = z}_C] < \varepsilon_1 + \varepsilon_2$$

ist, wobei (k_1, k_2) zufällig aus $K = K_1 \times K_2$ gewählt wird

- Hierzu zeigen wir folgende zwei Behauptungen

Behauptung 1

$$\Pr[\underbrace{x \notin \{x_1, \dots, x_q\}}_B \wedge \underbrace{f_{k_1}(x) \in \{f_{k_1}(x_1), \dots, f_{k_1}(x_q)\}}_D] < \varepsilon_1$$

Behauptung 2

$$\Pr[\underbrace{f_{k_1}(x) \notin \{f_{k_1}(x_1), \dots, f_{k_1}(x_q)\}}_{\bar{D}} \wedge \underbrace{g_{k_2}(f_{k_1}(x)) = z}_C] < \varepsilon_2$$

Beweis (Schluss)

- Nun folgt für die Erfolgswk von A

$$\begin{aligned} & \Pr[\underbrace{x \notin \{x_1, \dots, x_q\}}_B \wedge \underbrace{g_{k_2}(f_{k_1}(x)) = z}_C] \\ &= \underbrace{\Pr(B \cap D \cap C)}_{< \varepsilon_1} + \underbrace{\Pr(B \cap \bar{D} \cap C)}_{< \varepsilon_2} < \varepsilon_1 + \varepsilon_2 \end{aligned}$$



Beweis von Behauptung 1

$$\Pr[\underbrace{x \notin \{x_1, \dots, x_q\}}_B \wedge \underbrace{f_{k_1}(x) \in \{f_{k_1}(x_1), \dots, f_{k_1}(x_q)\}}_D] < \varepsilon_1$$

- Betrachte den adaptiven Kollisionsangreifer A' gegen $\mathcal{H}_1$, der zufällig einen Schlüssel $k_2 \in K_2$ wählt und A wie folgt simuliert:
 - Für jede Frage x_i von A erfragt A' den MAC-Wert $y_i = f_{k_1}(x_i)$ und gibt an A die Antwort $z_i = g_{k_2}(y_i)$ zurück
 - Sobald A ein Paar (x, z) ausgibt, erfragt A' den MAC-Wert $y = f_{k_1}(x)$ und gibt im Fall $x \notin \{x_1, \dots, x_q\} \wedge y \in \{y_1, \dots, y_q\}$ das Paar (x, x_i) für einen beliebigen Index i mit $y = y_i$ aus
- Da A' genau dann Erfolg hat, wenn das Ereignis $B \cap D$ eintritt, folgt die Behauptung □

Beweis von Behauptung 2

$$\Pr[\underbrace{f_{k_1}(x) \notin \{f_{k_1}(x_1), \dots, f_{k_1}(x_q)\}}_{\overline{D}} \wedge \underbrace{g_{k_2}(f_{k_1}(x)) = z}_C] < \varepsilon_2$$

- Betrachte den adaptiven Fälscher A'' gegen $\mathcal{H}_2$, der zufällig einen Schlüssel $k_1 \in K_1$ wählt und A wie folgt simuliert:
 - A'' gibt bei jeder Anfrage x_i von A die Antwort des Orakels g_{k_2} auf die Frage $y_i = f_{k_1}(x_i)$ zurück und
 - sobald A ein Paar (x, z) ausgibt, gibt A'' das Paar $(f_{k_1}(x), z)$ aus
- Da A'' genau dann Erfolg hat, wenn das Ereignis $\overline{D} \cap C$ eintritt, folgt die Behauptung □