

Man In The Browser 攻撃対策を実現する 人間・銀行サーバ間のセキュア通信プロトコル（その3）

向平浩貴^{†1} 神農泰圭^{†1} 土屋貴史^{†1}
大木哲史^{†1} 高橋健太^{†2} 尾形わかは^{†3} 西垣正勝^{†1}

概要：近年、インターネットバンキングにおける不正送金被害は依然として多く、その手口の中でも特に Man-in-the-Browser (MITB) 攻撃が注目を集めている。CSEC76 において、CAPTCHA を用いて MITB 攻撃に耐性のあるプロトコルを構成し、タグベース暗号の安全性をベースに CAPTCHA の安全性を定義した上で、プロトコルの安全性をタグベース CAPTCHA の安全性に帰着させ安全性証明を行った。本稿では、タグベース CAPTCHA の安全性定義および安全性証明について再検討し、さらに、新たな CAPTCHA の安全性定義として IND-C-CCA 安全性を定義し、その安全性証明を行う。IND-C-CCA 安全性は、公開鍵暗号における安全性である IND-CCA の定義を拡張したものであり、用いる CAPTCHA が IND-C-CCA 安全性を満たすならば、プロトコルは SUB-MIM 安全性を満たすことを示した。

1. はじめに

1.1 背景

近年、インターネットバンキングにおける不正送金の被害は減少傾向にあるが、その被害額は約 5 億 6,400 万円であり、依然として多くの被害が発生している[1]。不正送金にはフィッシングやなりすまし等の様々な攻撃手法が存在するが、特に Man-in-the-browser 攻撃（以下、MITB 攻撃とする）が注目を集めている。MITB 攻撃とは、ブラウザに感染したマルウェアがブラウザ・サーバ間の送受信の盗聴および改ざんを行う攻撃である。現在、多くのインターネットバンキングでは、PC・サーバ間の通信内容を SSL (TLS) で暗号化するというエンドツーエンドのセキュア通信を行うことで不正送金を対策している。しかし、MITB 攻撃では、マルウェアがクライアント側で改ざん等の不正行為を行うため、上記のような方法では対策が困難である。土屋らは、人間（ユーザ）とコンピュータ（銀行サーバ）の間に直にセキュア通信チャネルを構築するというアイデアに基づき、MITB 攻撃対策としてタグベース CAPTCHA を利用したチャレンジ&レスポンス方式のセキュア通信プロトコルを提案している。また、(1,N)-OW-CAPTCHA-CCA を満たすタグベース CAPTCHA を用いた提案プロトコルは MITB 攻撃に対し安全であることを示した[2]。

1.2 本稿の貢献

本稿では、土屋らの提案するタグベース CAPTCHA の安全性について再検討し、OW-TBC-CCA 安全性として定義し、プロトコルの安全性証明を行った。また、新たな帰着先の CAPTCHA の安全性として IND-C-CCA 安全性を提案し、プロトコルの安全性証明を行った。IND-C-CCA 安全性は一般的な CAPTCHA の安全性定義であるため、これにプロトコルの安全性を帰着させることでよりフォーマルな安全性証明を行うことができた。

1.3 本稿の構成

1 章では、本研究の背景と貢献について述べた。2 章では、本研究の対象とするインターネットバンキングにおける送金プロトコルとそれに対する MITB 攻撃について述べ、本稿で扱う暗号理論的な記法や概念について述べる。3 章では、CAPTCHA の定式化と安全性定義を行う。4 章では、提案プロトコルの定式化と安全性定義を行う。5 章で提案プロトコルの安全性証明を行い、6 章でまとめと今後の課題について述べる。

2. 準備

2.1 インターネットバンキングにおける送金プロトコル

本稿で想定するインターネットバンキングにおける送金プロトコル（図 1）について述べる。ここでは簡単のため仕組みを単純化して説明する。

2.1.1 構成要素

インターネットバンキングにおける送金プロトコルの構成要素は以下の通りである。

ユーザ：インターネットバンキングを利用する顧客である。送金処理を実行する際には、金融機関が提供する送金プロトコルに従い PC を操作し、ブラウザを利用する。人間であるユーザは低い計算能力（および記憶能力）しか有していないが、非常に高い認知能力を有するものとする。

ブラウザ：ユーザがインターネットバンキングを利用する際に用いる PC 内のブラウザである。コンピュータであるブラウザは高い計算能力（および記憶能力）を有する。

銀行サーバ：インターネットバンキングを提供する金融機関のサーバである。本稿では銀行サーバは不正を行わないものとする。銀行サーバはコンピュータであるため、高い計算能力（および記憶能力）を有する。

送金プロトコルは以下の手順に従って動作する。

- ① ユーザは送金情報 X をブラウザに入力する。
- ② ブラウザは X を銀行サーバへ送信する。
- ③ 銀行サーバは X に対する確認情報 Y をブラウザへ送信する。
- ④ ブラウザは X に対する確認情報 Y をユーザへ送信する。

^{†1} 静岡大学
Shizuoka University

^{†2} (株)日立製作所
Hitachi Ltd.

^{†3} 東京工業大学
Tokyo Institute of Technology University

- ⑤ ユーザは $Y = X$ であることを確認する。
- ⑥ ユーザはブラウザに TRUE (送金確定) を入力する
- ⑦ ブラウザは TRUE を銀行サーバへ送信する
- ⑧ 銀行サーバは TRUE を受信し, X を受理する。

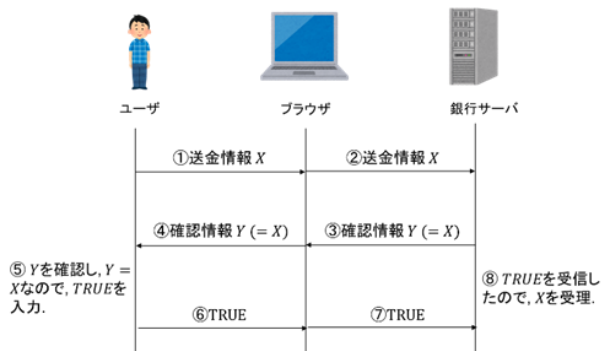


図 1 インターネットバンキングにおける送金プロトコル

2.2 MITB 攻撃

MITB 攻撃は, ブラウザに感染したマルウェアがブラウザ・サーバ間の送受信の盗聴や改ざんを行い, 不正送金を行う攻撃である。鈴木ら[3]は攻撃シナリオの違いより, MITB 攻撃を「ID 盗取型 MITB 攻撃」, 「取引内容改ざん型 MITB 攻撃」に分類している。本稿では, 取引内容改ざん型 MITB 攻撃のみを対象とする。

2.2.1 取引内容改ざん型 MITB 攻撃

一般的な送金プロトコルに対する取引内容改ざん型 MITB 攻撃を図 2 に示す。取引内容改ざん型 MITB 攻撃は以下の手順で行われる。

- ① ユーザは送金情報 X をブラウザに入力する。
- ② マルウェアは入力された X を X' に改ざんし, 銀行サーバへ送信する。
- ③ 銀行サーバは X' に対する確認情報 Y' をブラウザへ送信する。
- ④ マルウェアは受信した Y' を Y に改ざんし, ユーザへ送信する。
- ⑤ ユーザは $Y = X$ であることを確認する。
- ⑥ ユーザはブラウザに TRUE (送金確定) を入力する。
- ⑦ マルウェアは TRUE を銀行サーバへ送信する。
- ⑧ 銀行サーバは TRUE を受信し, X' を受理する。

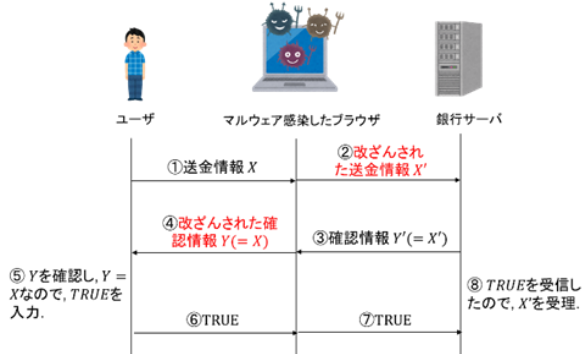


図 2 取引内容改ざん型 MITB 攻撃

2.3 記法

本稿で用いる記法を導入する。 $x \leftarrow X$ は有限集合 X から要

素を一様ランダムに選ぶことを表す。また, $y \leftarrow A(x)$ は確率的アルゴリズム A が入力 x に対し, y を出力することを表す。そして, $y \leftarrow x$ は要素 x を y に代入することを表す。 $x|y$ は2つのビット列 x と y の連結を表す。関数 $f(\lambda)$ が全ての定数 $c > 0$ に対して $\frac{1}{\lambda^c}$ よりも早く 0 に収束するとき, $f(\lambda)$ は無視可能であるといい, ある関数 f が無視可能であるということを $f(\lambda) = \varepsilon(\lambda)$ と表す。また, 確率的多項式時間を PPT と略記する。

2.4 公開鍵暗号における安全性の概念

一般に, 公開鍵暗号の安全性は達成度と攻撃方法の2つの強度を組み合わせることで捉えることができるとされている[4]。達成度は一方向性 (OW: One Wayness), 強秘匿性 (IND: Indistinguishability), 頑強性 (NM: Non Malleability) の3つに分類され, OW は暗号文 c から平文 m 全体が得られないことを保証し, IND は暗号文 c から平文 m のいかなる部分情報も得られないことを保証し, NM は平文 m に対する暗号文 $c = \text{Enc}(pk, m)$ が与えられている時に, m と関係のある別の平文 m' に対する暗号文 c' を生成できないことを保証する。攻撃方法は選択平文攻撃 (CPA: Chosen Plaintext Attack), 選択暗号文攻撃 (CCA: Chosen Ciphertext Attack) のように分類され, CPA 攻撃者は「平文をクエリすると対応する暗号文を返す暗号化オラクル」を利用することができる。上記の達成度と攻撃方法を組み合わせることで安全性は定式化される。

2.4.1 OW-CCA 安全性

公開鍵暗号 $\Pi = (\text{Gen}, \text{Enc}, \text{Dec})$ において, 攻撃者 A に対する挑戦者 B を設定し, A と B の間で実行される次のような OW-CCA ゲームを構成する。

1. B は Gen に 1^k (k はセキュリティパラメータ)を入力し, 秘密鍵 sk , 公開鍵 pk のペアを出力し, A に pk を入力する。
2. A はチャレンジ暗号文の作成を B に依頼する。
3. B は平文 m を平文空間から一様に選択し, チャレンジ暗号文 $c = \text{Enc}(pk, m)$ を作成し, A に返答する。
4. A は平文 $\hat{m}$ を出力する。このとき, $\hat{m} = m$ であれば攻撃者の勝ちとする。

上記のゲームにおいて, A は任意のタイミングで復号オラクルを利用することができる。復号オラクルは, 暗号文を送るとその復号結果を返すという動作をする。ただし, チャレンジ暗号文に関しては復号結果を返す代わりに $\perp$ (復号不可) を返す。上記の OW-CCA ゲームに対する A のアドバンテージ (優位性) を

$$\text{Adv}_A^{\text{OW-CCA}}(k) = \Pr[\hat{m} = m]$$

と定義し、いかなる PPT 攻撃者 A に対しても $Adv_A^{OW-CCA}(k) < \epsilon(k)$ が成立するとき、 Π は OW-CCA 安全であるという。

2.4.2 IND-CCA 安全性

公開鍵暗号 $\Pi = (Gen, Enc, Dec)$ において、攻撃者 A に対する挑戦者 B を設定し、 A と B の間で実行される次のような IND-CCA ゲームを構成する。

1. B は Gen に 1^k を入力し、秘密鍵 sk 、公開鍵 pk のペアを出力し、 A に pk を入力する。
2. A は 2 つの平文 m_0, m_1 を B に送る。
3. B は m_0, m_1 のうち 1 つを選択し、チャレンジ暗号文 $c = Enc(pk, m_b)$ を作成し、 A に返答する。
4. A は $\hat{b}$ を出力する。このとき、 $\hat{b} = b$ であれば攻撃者の勝ちとする。

上記のゲームにおいて、 A は任意のタイミングで復号オラクルを利用することができる。復号オラクルは、暗号文を送るとその復号結果を返すという動作をする。ただし、チャレンジ暗号文に関しては復号結果を返す代わりに $\perp$ (復号不可) を返す。上記の IND-CCA ゲームに対する A のアドバンテージ (優位性) を

$$Adv_A^{IND-CCA}(k) = \left| \Pr[\hat{b} = b] - \frac{1}{2} \right|$$

と定義し、いかなる PPT 攻撃者 A に対しても $Adv_A^{IND-CCA}(k) < \epsilon(k)$ が成立するとき、 Π は IND-CCA 安全であるという。

2.5 メッセージ認証コードにおける安全性の概念

メッセージ認証コード (MAC) に対しては Impersonation attack と Substitution attack の 2 つの攻撃手法が存在する [5]。平文 m 、 m に対する MAC を MAC_m としたとき、Impersonation attack は、攻撃者が送信者から送られてくるメッセージを見ることなく、別の正当な平文と MAC の組 $(\bar{m}, MAC_{\bar{m}})$ を出力する攻撃である。Substitution attack は、攻撃者が (m, MAC_m) を受け取り、それを別の正当な平文と MAC の組 $(\bar{m}, MAC_{\bar{m}})$ に置き換えて出力する攻撃である。

3. CAPTCHA の定式化

3.1 定義

3.1.1 CAPTCHA

CAPTCHA (Completely Automated Public Turing test to tell Computers and Human Apart) とは、人間には正解が容易であるが、機械には正解が困難 (AI-hard) な問題をユーザに出題し、正解したユーザを人間と判定する技術である。現在、多くの Web サービスでマルウェアによるサービスの不正利用を防止するために用いられている。

CAPTCHA は以下の入出力を持つ 2 つの PPT アルゴリズム (C_Enc, C_Dec) からなる。

$$\begin{aligned} c &\leftarrow C_Enc(m) \\ m &\leftarrow C_Dec(c) \end{aligned}$$

C_Enc は、平文 m を入力として受け取り、CAPTCHA 型暗号文 c を出力する。 C_Dec は、CAPTCHA 型暗号文 c を入力として受け取り、平文 m を出力する。なお、CAPTCHA は AI-hard な問題であり、 C_Dec は人間にしか実行できないものとする。

3.1.2 タグベース CAPTCHA

土屋らはタグベース暗号 (Tag-based encryption) [6] の定義をもとにタグベース CAPTCHA とその安全性を定義している。タグベース CAPTCHA は以下の入出力を持つ 2 つの PPT アルゴリズム (TBC_Enc, TBC_Dec) からなる。

$$\begin{aligned} c &\leftarrow TBC_Enc(m, t) \\ m &\leftarrow TBC_Dec(c, t) \end{aligned}$$

TBC_Enc は、平文 m とタグ t を入力として受け取り、CAPTCHA 型暗号文 c を出力する。 TBC_Dec は、CAPTCHA 型暗号文 c とタグ t を入力として受け取り、平文 m を出力する。ここで、暗号文とともにその暗号化の際に用いたタグを TBC_Dec に入力すると必ず正しい平文を返し、暗号化の際に用いたタグとは異なるタグを入力すると $\perp$ (復号不可) を返す。タグベース CAPTCHA も一般的な CAPTCHA と同様に AI-hard な問題であり、 TBC_Dec は人間にしか実行できないものとする。

3.2 安全性定義

3.2.1 CAPTCHA の安全性: IND-C-CCA 安全性

CAPTCHA の安全性として、選択暗号文攻撃に対する識別不可能性 (IND-C-CCA) を定義する。CAPTCHA における選択暗号文攻撃とは、攻撃者が CAPTCHA 型暗号文に対応する平文を入手できる条件下で、挑戦者が 2 つの平文のうちどちらを暗号化したかを攻撃者が求める攻撃である。

攻撃者 A に対する挑戦者 B を設定し、 A と B の間で実行される図 3 のような IND-C-CCA ゲームを構成する。

1. A は 2 つの平文 m_0, m_1 を B に送る。
2. B は m_0, m_1 のうち 1 つを選択し、CAPTCHA 型暗号文 $c = Enc(m_b)$ を作成し、これをチャレンジとして A に返答する。
3. A は $\hat{b}$ を出力する。このとき、 $\hat{b} = b$ であれば攻撃者の勝ちとする。

上記のゲームにおいて、 A は任意のタイミングでヒューマンオラクル H を利用することができる。ヒューマンオラクルは J. Blocki ら [7] により定義されたオラクルで、CAPTCHA を解く能力を持つ人間のみがアクセスすることができ、「CAPTCHA 型暗号文をクエリするとその平文を返す」という動作をする。ただし、 A が H に対しチャレンジ c をクエリすることは禁止されている。上記の IND-C-CCA ゲームに対する A のアドバンテージ (優位性) を

$$Adv_A^{IND-C-CCA} = \left| \Pr[\hat{b} = b] - \frac{1}{2} \right|$$

と定義し、いかなる A に対し $Adv_A^{IND-C-CCA} < \epsilon$ が成立するとき、CAPTCHA は IND-C-CCA 安全であるという。

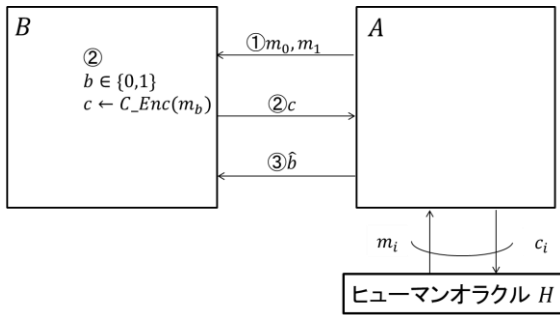


図 3 IND-C-CCA ゲーム

3.2.2 タグベース CAPTCHA の安全性:OW-TBC-CCA 安全

タグベース CAPTCHA の安全性として、タグ選択暗号文攻撃に対する一方向性 (OW-TBC-CCA) を定義する。タグベース CAPTCHA におけるタグ選択暗号文攻撃とは、攻撃者が CAPTCHA 型暗号文に対応する平文を入手できる条件下で、挑戦者から提示された CAPTCHA 型暗号文(ただし、暗号文のタグは攻撃者が指定できる)の平文を攻撃者が求める攻撃である。

攻撃者Aに対する挑戦者Bを設定し、AとBの間で実行される図4のようなOW-TBC-CCAゲームを構成する。

1. Aはターゲットタグ t^* を生成し、Bに t^* を送信する。
2. Bは平文 m を平文空間から一様に選択し、CAPTCHA 型暗号文 $c^* = C_Enc(t^*, m)$ を作成し、これをチャレンジとしてAに返答する。
3. Aは平文 $\hat{m}$ を出力する。このとき、 $\hat{m} = m$ であれば攻撃者の勝ちとする。

上記のゲームにおいて、Aは任意のタイミングでタグヒューマンオラクル tH を利用することができる。タグヒューマンオラクルは、ヒューマンオラクルと同様、CAPTCHA を解く能力を持つ人間のみがアクセスすることができるオラクルである。 tH は、「タグと CAPTCHA 型暗号文からなるクエリに対し、そのタグがその暗号文を暗号化した際に用いられたタグである場合には平文を、それ以外の場合には $\perp$ (復号不可) を返す」という動作をする。Aが tH に対し、ターゲットタグ t^* を含むクエリをすることは禁止されている。上記の OW-TBC-CCA ゲームに対するAのアドバンテージ (優位性) を

$$Adv_A^{OW-TBC-CCA} = \Pr[\hat{m} = m]$$

と定義し、いかなるAに対しても $Adv_A^{OW-TBC-CCA} < \epsilon$ が成立するとき、タグベース CAPTCHA は OW-TBC-CCA 安全であるという。

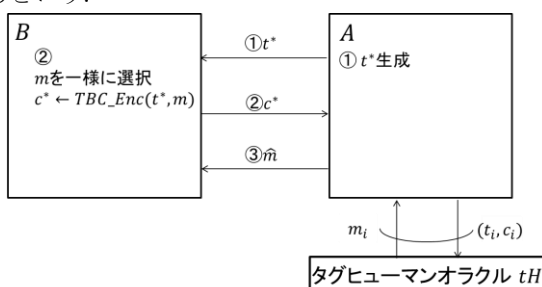


図 4 OW-TBC-CCA ゲーム

4. 提案方式

4.1 提案プロトコルの概要

本稿で提案するユーザ・銀行サーバ間のセキュア通信プロトコルの概要を図5に示す。本プロトコルは以下の手順に従って動作する。なお、図5の③において、一般的なCAPTCHAを用いた場合は $C = C_Enc(Y|R)$ 、タグベースCAPTCHAを用いた場合は $C = TBC_Enc(Y, R)$ となる。

- ① ユーザは送金情報 X をブラウザに入力する。
- ② ブラウザは X を銀行サーバへ送信する。
- ③ 銀行サーバは X を Y に代入する。乱数 R を生成し、 Y と R からCAPTCHA型暗号文 C を生成する。
- ④ 銀行サーバは C をブラウザに送信する。
- ⑤ ブラウザは C をユーザに提示する。
- ⑥ ユーザは C を解き、 Y と R を得る。 $Y = X$ ならば、 $Q(=R)$ 、 $Y \neq X$ ならば、 $Q(=\perp)$ をブラウザに入力する。
- ⑦ ユーザは Q をブラウザに入力する。
- ⑧ ブラウザは Q を銀行サーバへ送信する。
- ⑨ 銀行サーバは、 $Q = R$ ならば、 X を受理、 $Q = \perp$ ならば送金中止の処理を行う。

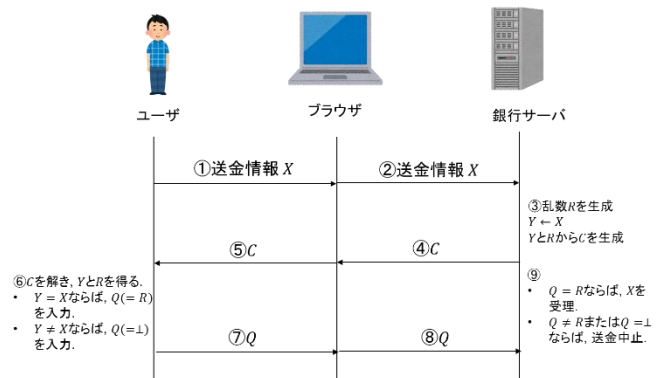


図 5 提案プロトコルの概要

4.2 提案プロトコルの定式化

提案プロトコルの本質は、送金情報 X をユーザ (人間) から銀行サーバに正しく送信することである。そこで、ユーザを送信者 S 、銀行サーバを受信者 R として、提案プロトコルを定式化すると

$$Protocol = \langle S^{H(\cdot)}, R \rangle$$

という形となる (図6)。ここで、人間であるユーザは「ヒューマンオラクルに任意にアクセスできる送信者 $S^{H(\cdot)}$ 」として表現されている。 x_S と x_R は S と R への入力であり、 y_S と y_R は S と R からの出力である。提案プロトコルでは、 $x_S = X$ 、 $x_R = \phi$ であり、 $y_S = TRUE$ (MITB 攻撃が発生していない場合) あるいは $\perp$ (MITB 攻撃が発生した場合)、 $y_R = X$ (MITB 攻撃が発生していない場合) あるいは $\perp$ (MITB 攻撃が発生した場合) である。 $y_R = \perp$ の際には銀行サーバは送金処理を中止する。

提案プロトコルの要件として、次節の形で完全性 (Completeness) と健全性 (Soundness) を定義する。

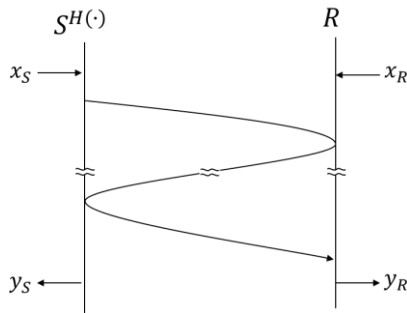


図 6 提案プロトコルの定式化

4.2.1 完全性 (Completeness)

CAPTCHA を解く能力を持つ人間 $S^{H(\cdot)}$ とサーバ R がプロトコルを実行したとき、圧倒的確率で以下を満たす場合、提案プロトコルは完全性を満たすという。

$$\Pr[x_S = y_R \mid \langle S^{H(\cdot)}(x_S), R(x_R) \rangle = (y_S, y_R)]$$

4.2.2 健全性 (Soundness)

CAPTCHA を解く能力を持たない機械 (攻撃者) A に対し、以下で定義される $Adv_{A,R}^{SND}$ が無視できる場合、提案プロトコルは健全性を満たすという。

$$Adv_{A,R}^{SND} = \Pr[x_S \neq y_R \wedge y_R \neq \perp \mid \langle A(x_S), R(x_R) \rangle = (y_A, y_R)]$$

4.3 提案プロトコルの安全性定義

本稿で対象としている「取引内容改ざん型 MITB 攻撃」は、メッセージ (送金情報) のすり替えという観点から考えると、メッセージ認証コード (MAC) に対する Impersonation attack ならびに Substitution attack に相当する攻撃であると捉えることができる。そこで、これらの定義をもとに、提案プロトコルの取引内容改ざん型 MITB 攻撃に対する安全性として、IMP-MIM 安全性、SUB-MIM 安全性の2つを定義する。このうち、本稿ではより強い安全性定義である SUB-MIM 安全性について安全性証明を行う。

4.3.1 IMP-MIM 安全性

提案プロトコルに対する IMP-MIM 攻撃者を A とする。 A は「CAPTCHA を解く能力を持つ人間 $S^{H(\cdot)}$ を学習時にのみ使役することができる PPT アルゴリズム」である。攻撃者 A 、銀行サーバ R 間の以下の IMP-MIM ゲームを定義する。

学習フェーズ: $A^{S^{H(\cdot)}}$, R 間でプロトコル $\langle A^{S^{H(\cdot)}}, R \rangle$ を任意の入力 x_{S0} を用いて任意の回数実行し、通信系列 π を得る。 A は学習フェーズにのみ CAPTCHA を解く能力を持つ人間 $S^{H(\cdot)}$ にアクセスすることができる。ここでは、これを $A^{S^{H(\cdot)}}$ と表記している。

攻撃フェーズ: A , R 間でプロトコル $\langle A, R \rangle$ を実行する。ここで、 A は学習フェーズの際に入力した x_{S0} を入力することはできない。

上記のゲームにおける A のアドバンテージを

$$Adv_A^{IMP-MIM} =$$

$$\Pr[x_S \neq y_R \wedge y_R \neq \perp \mid \langle A^{S^{H(\cdot)}}(x_S), R(x_R) \rangle = (y_A, y_R)]$$

と定義し、いかなるアルゴリズム A に対してもアドバンテージが無視できるとき、提案プロトコルは IMP-MIM 安全

を満たすという。

4.3.2 SUB-MIM 安全性

提案プロトコルに対する SUB-MIM 攻撃者を A とする。 A は「CAPTCHA を解く能力を持つ人間 $S^{H(\cdot)}$ を学習時および攻撃時に使役することができる PPT アルゴリズム」である。攻撃者 A 、銀行サーバ R 間の以下の SUB-MIM ゲームによって定義される。

学習フェーズ: $A^{S^{H(\cdot)}}$, R 間でプロトコル $\langle A^{S^{H(\cdot)}}, R \rangle$ を任意の入力 x_{S0} を用いて任意の回数実行し、通信系列 π を得る。 A は学習フェーズに CAPTCHA を解く能力を持つ人間 $S^{H(\cdot)}$ にアクセスすることができる。ここでは、これを $A^{S^{H(\cdot)}}$ と表記している。

攻撃フェーズ: $A^{S^{H(\cdot)}}$, R 間でプロトコル $\langle A^{S^{H(\cdot)}}, R \rangle$ を実行する。 A は攻撃フェーズにも CAPTCHA を解く能力を持つ人間 $S^{H(\cdot)}$ にアクセスすることができる。ここでは、これを $A^{S^{H(\cdot)}}$ と表記している。 A は学習フェーズの際に入力した x_{S0} を入力することはできない。

上記のゲームにおける A のアドバンテージを

$$Adv_A^{SUB-MIM} =$$

$$\Pr[x_S \neq y_R \wedge y_R \neq \perp \mid \langle A^{S^{H(\cdot)}, S^{H(\cdot)}}(x_S), R(x_R) \rangle = (y_A, y_R)]$$

と定義し、いかなるアルゴリズム A に対してもアドバンテージが無視できるとき、提案プロトコルは SUB-MIM 安全を満たすという。

5. 提案プロトコルの安全性証明

5.1 OW-TBC-CCA 安全性への帰着

OW-TBC-CCA 安全なタグベース CAPTCHA を用いる提案プロトコルは SUB-MIM 安全を満たすことを証明する。

定理 1

タグベース CAPTCHA が OW-TBC-CCA 安全ならば、そのタグベース CAPTCHA を用いる提案プロトコルは SUB-MIM 安全を満たす。

定理 1 の証明

定理 1 の対偶をとり、以下の (1) を証明する。

- (1) 提案プロトコルの SUB-MIM 安全性を無視できない確率で破る攻撃者 A が存在するならば、タグベース CAPTCHA の OW-TBC-CCA 安全性を無視できない確率で破る攻撃者 B^A が存在する。

B^A を図 7 のように構成する。① B は A に通信系列 π を入力する。② B は送金情報 X を A に入力する。③ A は $X' (\neq X)$ を出力する。④ B は X' をターゲットタグ t^* として OW-TBC-CCA ゲームの挑戦者に送る。⑤ 挑戦者は平文 m を一様選択し、 $c^* \leftarrow TBC_Enc(t^*, m)$ を計算し、 c^* を B に入力する。 B は c^* を A に入力する。⑥ A は c' を出力する。⑦ B は (X, c') をタグヒューマンオラクルに送る。ここで、 $X \neq t^*$ であるためクエリを送信することができる。⑧ タグヒューマンオラクルは、 $TBC_Dec(X, c')$ を実行した結果 (m' あるい

図 7 定理 1 の証明

図 8 定理 2 の証明

- [1] “平成 29 年上半期におけるサイバー空間をめぐる脅威の情勢等について”,
http://www.npa.go.jp/publications/statistics/cybersecurity/data/H29_kami_cyber_jousei.pdf (参照 2018/06/07).
- [2] 土屋 貴史. 他. ”Man In The Browser 攻撃対策を実現する人間・銀行サーバ間のセキュア通信プロトコル(その2)”. CSEC. 2017, 2017-CSEC-76 (6), 1-7 (2017-02-23), 2188-8655.
- [3] “MITB 対策 | 不正送金対策 | トランザクション署名 | OCRA 仕様 OTP トークン | OATH 準拠 | 飛天ジャパン”,
https://ftsafe.co.jp/solutions/ocra_mitb/ (参照 2018/06/07).
- [4] M. Bellare, A. Desai, D. Pointcheval and P. Rogaway, Relations Among Notions of Security for Public-Key Encryption Scheme, Advances in Cryptology– CRYPTO ’98, volume 1462 of LNCS, pages 26-45.
- [5] G.Simmons, Authentication theory/coding theory. Advances in Cryptology, Springer, Santa Barbara, California, USA, 1985; 411–431.
- [6] E. Kiltz, Chosen-ciphertext security from tag-based encryption. TCC 2006, LNCS 3876, pp. 581–600.
- [7] J. Blocki and H. –S. Zhou, Designing Proof of Human-Work Puzzles for Cryptocurrency and Beyond. Berlin, Heidelberg: Springer Berlin Heidelberg, 2016, pp. 517–546.