

Ad hoc útvonalválasztó protokollok biztonsága

ÁCS GERGELY, BUTTYÁN LEVENTE, VAJDA ISTVÁN

Budapesti Műszaki és Gazdaságtudományi Egyetem, Híradástechnikai Tanszék, CrySyS laboratórium
{acs, buttyan, vajda}@crys.sys.hu

Kulcsszavak: ad hoc hálózatok, forrás alapú ad hoc útvonalválasztás, bizonyított biztonság, szimulációs paradigma

A cikkben egy olyan formális módszert mutatunk be, amellyel az ad hoc hálózatok számára javasolt, igény szerinti, forrás-alapú útvonalválasztó protokollokat (on-demand source routing) lehet biztonsági szempontból elemezni. A módszer alapját a szimulációs paradigma adja, mely egy jól ismert, általános eljárás kriptográfiai protokollok biztonságának bizonyítására. Formálisan megfogalmazzuk, hogy mit értünk biztonságos útvonalválasztás alatt, melyhez felhasználjuk a statisztikai megkülönböztethetlenség fogalmát. A gyakorlati alkalmazást egy példán keresztül szemléltetjük, melyben ismertetjük az endairA útvonalválasztó protokoll működését, és bebizonyítjuk, hogy a protokoll biztonságos az általunk definiált modellben.

1. Bevezetés

Az ad hoc hálózat egyenrangú csomópontokból álló, előre telepített infrastruktúrával nem rendelkező, vezeték nélküli hálózat, melyben a csomópontok terminál- és hálózati funkciókat egyaránt ellátnak. A csomópontok gyakran korlátozott energia ellátással (elemmel, akkumulátorral) rendelkeznek. Ezért, és a rádió kapcsolatok miatt fellépő interferencia csökkentése érdekében, az ad hoc hálózat csomópontjai multihop kommunikációt használnak. Infrastruktúra hiányában, a multihop kapcsolathoz szükséges útvonalválasztó és -karbantartó feladatokat maguk a csomópontok végzik.

Az útvonalválasztásnak alapvetően két formáját lehet megkülönböztetni: a pro-aktív és a reaktív (vagy igény szerinti) eljárásokat. Jelen cikkben az utóbbi típussal foglalkozunk. Reaktív útvonalválasztás során egy kapcsolatot kezdeményező csomópont csak igény esetén próbál útvonalat találni a célcsomópont felé. Ekkor a kezdeményező csomópont egy útvonalkérésrel (route request – *rreq*) árasztja el a hálózatot. A kérést minden csomópont megkapja, majd ahhoz saját azonosítóját hozzáadva, többesküldéssel (broadcast) továbbítja. A célcsomóponthoz érkező kérésüzenetekre, a célcsomópont egy (vagy több) útvonalválasszal (route reply – *rrep*) válaszol, mely tartalmazza a kérésben rögzített csomópontazonosítók sorozatát, azaz a felfedezett útvonalat. A válasz a kérés által követett útvonal fordítottján jut vissza a kezdeményező csomóponthoz.

A biztonságos útvonalválasztás feladata ezen mechanizmus helyes működésének megvalósítása támadó jelenlétében. Ez elsődleges fontosságú, mivel az útvonalválasztó mechanizmus manipulálásával egy támadó viszonylag kevés erőforrás felhasználásával az egész hálózatot működésképtelenné teheti.

Az irodalomban számos biztonságosnak mondott útvonalválasztó protokoll jelent meg ([3] jó áttekintést ad a területről), ám ezen protokollok szerzői formálisan

nem igazolták állításaikat. Tudomásunk szerint [2] az első olyan munka, melyben ad hoc útvonalválasztó protokollok elemzésére precíz matematikai modell jelent meg. [2]-ben a szerzők megmutatták két irodalomban javasolt protokollról (SRP, Ariadne), hogy azok támadhatóak, és javasoltak egy új protokollt, melyről bebizonyították, hogy az biztonságos. A [2]-ben használt módszer alapját a szimulációs paradigma adja, mely egy jól ismert eljárás a kriptográfiai protokollok biztonságának bizonyítására [4,5]. A [2]-ben definiált modell azonban csak egy korlátozott Aktív-1-1 támadót enged meg, mely egyetlen kompromittált csomópontot és egyetlen kompromittált azonosítót használhat. További megkötés, hogy a támadó egyidejűleg csak egy útvonalválasztási folyamat végrehajtását támadhatja.

Jelen cikkben a [2]-ben használt modellt általánosítjuk tetszőleges Aktív-y-x támadó esetére, megengedve párhuzamos protokollfutásokat is. Megmutatjuk továbbá, hogy a [2]-ben javasolt endairA protokoll ebben a bővített modellben is biztonságos. Itt elsősorban csak a fő gondolatokat írjuk le; a munka részletesebb leírását [1] tartalmazza.

2. A modell

2.1. A hálózat modellezése

Az útvonalválasztás vizsgálata során az ad hoc hálózatot statikusnak tekintjük és egy $G(E, V)$ gráffal modellezzük, ahol a csúcsoknak egyértelműen megfelelnek az egyes hálózati csomópontok, és két csúcs között akkor és csak akkor van él, ha az azoknak megfelelő csomópontok hallják egymás adását, vagyis szomszédosak. Feltesszük, hogy a hálózat minden csomópontja rendelkezik egy azonosítóval, amely egyértelműen azonosítja a csomópontot a hálózatban (ez lehet például egy publikus kulcs). Feltesszük, hogy az azonosítók hitelesítettek, de néhány azonosító kompromittálódott, s az ezek hitelesítéséhez szükséges kulcsok a tá-

madó birtokába jutottak. A támadót is tartalmazó hálózatot röviden *konfigurációnak* hívjuk. Formálisan a konfiguráció egy $(G(E, V), V^*, L)$ hármas, ahol $G(E, V)$ a hálózatot reprezentáló gráf, a $V^* \subset V$ a támadó irányítása alatt álló csomópontok halmaza, L pedig egy olyan függvény, amely mindegyik csúcshoz hozzárendeli a hozzátartozó csomópont által birtokolt azonosítók halmazát. Megbízható csomópontok esetén ezek egy elemű halmazok, míg a támadó irányítása alatt álló minden csomópont az összes kompromittált azonosítót tartalmazó halmazt rendeljük.

2.2. A támadó modellezése

A támadóról az alábbiakat feltételezzük:

- fizikailag nem lehet mindenütt jelen, ezért nincs irányítása a teljes hálózat felett;
- x csomópontot irányít és y kompromittált azonosítót birtokol (Aktív- y - x támadó, ahol $x, y \geq 1$);
- a birtokában levő kompromittált azonosítók halmaza és a megbízható csomópontok azonosítóinak halmaza diszjunkt;
- kommunikációs képességeit tekintve egy átlagos csomópont képességeivel rendelkezik, azaz a támadó csomópontok csak szomszédaiknak tudnak üzenetet küldeni és csak azok adását hallják;
- aktív, abban az értelemben, hogy nemcsak üzeneteket hallgat le, hanem képes új üzenetek megalkotására és beszúrására is, valamint képes üzeneteket módosítani, késleltetni stb.;
- nem adaptív abban az értelemben, hogy az egyes protokollfutások ellen végrehajtott támadásai egymástól függetlenek.

Az útvonalválasztást kezdeményező csomóponttól és a célcsomóponttól feltesszük, hogy megbízhatóak.

2.3. A plauzibilis útvonal definíciója

A biztonság fogalmának definiálása nem egyértelmű feladat. Egy kézenfekvő megoldás lehet az optimális (egyes esetekben a legrövidebb) útvonalak visszaadásának követelménye, ez viszont az eltérő csomóponti késleltetések és a szinte mindig alkalmazott optimalizációk miatt irreális [2]. Továbbá az is látható, hogy nem tudjuk megakadályozni (az útvonalválasztás szintjén), hogy a támadó az általa birtokolt azonosítók közül tetszőleges számút hozzáadjon egy lehallgatott útvonalkéréshez, és hasonlóan nem tudjuk megakadályozni a szomszédos támadó csomópontok kooperációját sem [1]. Látható tehát, hogy bizonyos támadások ellen egyáltalán nem tudunk védekezni, vagy ha tudunk, akkor a gyakorlati esetek többségében a védelem kialakítása túl költséges lenne. Ezért körültekintően kell megadni a biztonság fogalmát: ha túl szigorú definíciót adunk, akkor a fent említett elkerülhetetlen támadások miatt egyetlen protokoll sem fogja kielégíteni definícióinkat, míg ellenkező esetben olyan protokollt is biztonságosnak nyilváníthatunk, mely nemcsak az elkerülhetetlen támadások ellen védtelen.

E problémát úgy oldjuk meg, hogy a „helyes” útvonal fogalmába „beépítjük” az elkerülhetetlen támadá-

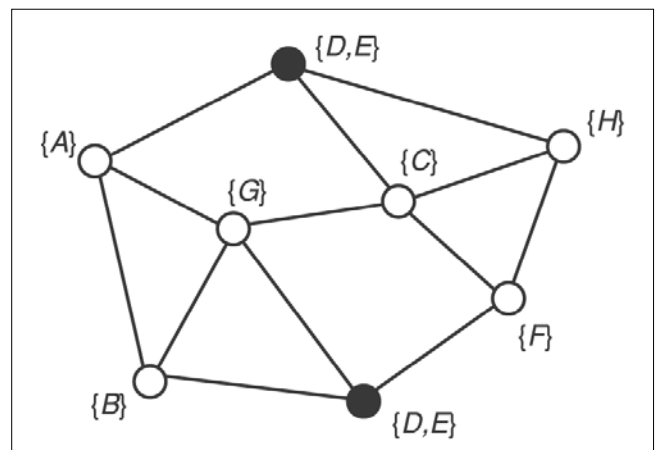
sok hatásait. A „helyes” útvonalakat *plauzibilis* útvonalaknak nevezzük és az alábbi módon definiáljuk.

Minden konfiguráció egyértelműen redukálható egy olyan másik konfigurációra, amelynek megfeleltetett gráf nem tartalmaz szomszédos támadó csomópontokat. Más szavakkal a redukált konfigurációban összevonjuk a szomszédos támadó csomópontoknak megfeleltetett csúcsokat. A redukált konfigurációhoz tartozó redukált gráfot $\underline{G}$ -vel jelöljük.

Definíció:

Egy útvonalat alkotó csomópontazonosítók sorozata akkor *plauzibilis*, ha

- nem tartalmaz ismétlődő azonosítókat és
- *particionálható* olyan részsorozatokra, amely részsorozatok egyértelműen hozzárendelhetők olyan csúcsokhoz a redukált gráfban, hogy ezen csúcsoknak megfeleltetett csomópontok rendelkeznek a hozzájuk tartozó partíció összes azonosítójával, és ezen csúcsok egy útvonalat alkotnak a redukált gráfban.



1. ábra

A csúcsösszevonás valamint a plauzibilis útvonal definíciójának értelmét a fentebb tárgyalt elkerülhetetlen támadások adják. Az 1. ábrán egy redukált konfiguráció látható. A telített csúcsok reprezentálják az összevont támadó csomópontokat, melyek a D és E kompromittált azonosítóval rendelkeznek. Látható, hogy ebben a konfigurációban az A, D, E, C, F útvonal plauzibilis, melynek egy helyes particionálása $A|D, E|C|F$, viszont A, B, D, E, H nem plauzibilis útvonal, hiszen az E és H azonosítónak megfeleltetett csomópontok nem szomszédosak.

2.4. A szimulációs paradigma

A modell egy sarkalatos pontja a biztonság fogalmának precíz definiálása. Erre a széleskörben használt szimulációs paradigmát [4,5] szeretnénk alkalmazni. A szimulációs paradigma alapja, hogy egy valós és egy ideális modellt definiálunk, ahol a valós modell leírja a protokoll valóságos működését, míg az ideális modell azt reprezentálja, hogy mit várunk el egy biztonságos (ideális) protokolltól. A valós modellre gondolhatunk úgy, mint a protokoll implementációjára, míg az ideális mo-

dellet tekinthetjük a feladat specifikációjának. Mindkét modellnek része egy támadó is. Az ideális modellbeli támadó reprezentálja az elkerülhetetlen támadásokat. Ezzel szemben, a valós modellbeli támadó erejét nem korlátozzuk, csak annyiban, hogy a valós támadó által végrehajtható lépések száma az alkalmazott biztonsági paraméter polinomiális függvénye kell legyen.

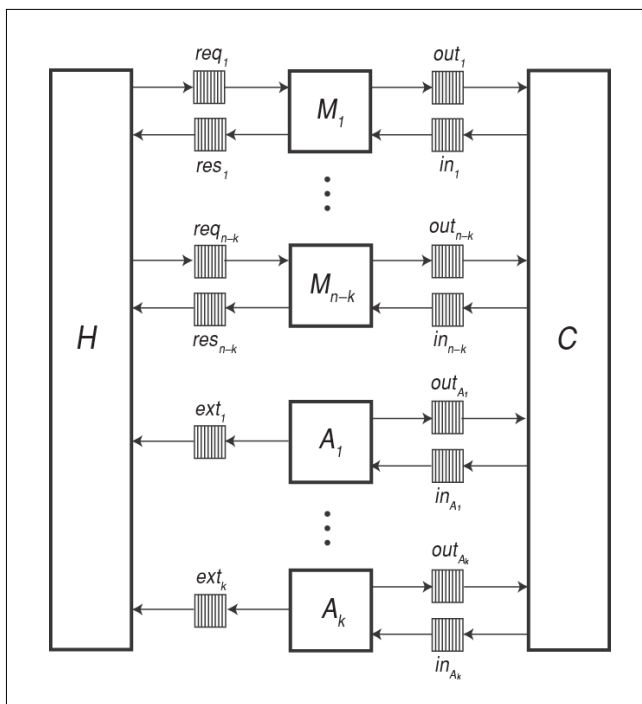
Egy protokollt ezek után akkor tekintünk biztonságosnak, ha a valós és az ideális modellek megkülönböztethetetlenek a becsületes résztvevők számára. Precízebben, egy protokoll akkor biztonságos, ha bármely valós modellbeli A támadóhoz létezik olyan ideális modellbeli A' támadó, hogy minden amit A elér a valós modellben, azt A' is elér az ideális modellben ugyanolyan mennyiségű erőforrással. Más szavakkal bármely valós támadó tevékenységének hatását szimulálni lehet az ideális modellben. Mivel azonban az ideális modellben definíció szerint csak az elkerülhetetlen támadások lehetségesek, ezért következik, hogy ennél többet egy valós támadó sem tud elérni a valós modellben.

A szimulációs paradigmának megfelelően a következőkben definiáljuk az ad hoc útvonalválasztás valós és ideális modelljét, majd pontosan megadjuk, hogy milyen értelemben követeljük meg a kettő megkülönböztethetlenségét.

2.5. A valós modell

A valós modellt a 2. ábra szemlélteti. A modell interaktív és véletlenszerű Turing gépekből épül fel, melyek közös szalagok segítségével kommunikálnak. Az egyes gépek az egyes protokoll résztvevők, valamint a támadó működését és kommunikációját modellezzik. $M_1, \dots, M_{n-k}$ reprezentálja a megbízható működésű csomópontokat, melyek pontos működése az adott útvonalválasztó protokolltól függ.

2. ábra A valós modell

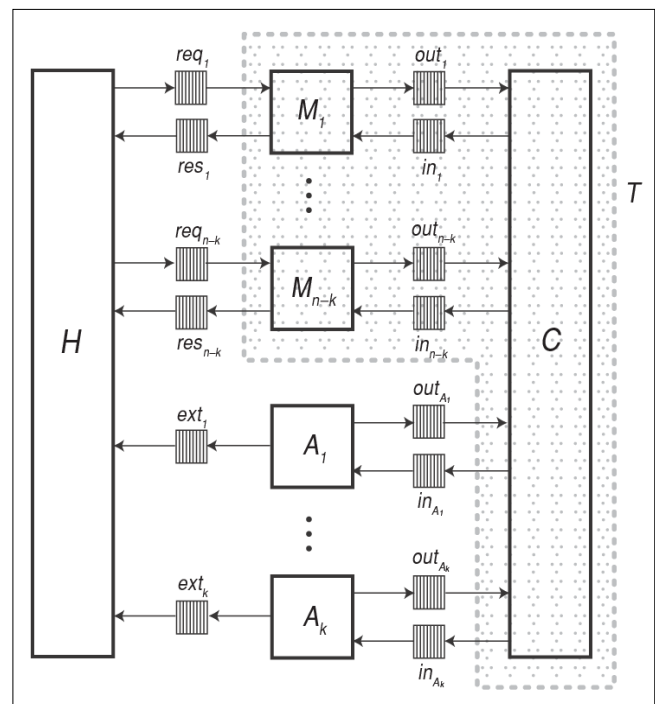


$A_1, \dots, A_k$ a támadó csomópontokat reprezentálja, melyek tetszőleges polinom idejű gépek lehet. C gép modellezi a kommunikációs kapcsolatokat (azaz $\mathcal{G}$ éleit). C feladata, hogy az egyes gépek kimeneti szalagjain (out_i és out_{A_i}) megjelenő protokollüzeneteket átmozgassa a gépek ($\mathcal{G}$ szerinti) szomszédainak bemeneti szalagjára (in_i és in_{A_i}). Végül H reprezentálja az útvonalválasztás feletti protokollrétegeket (és a becsületes felhasználókat). A modellben H kezdeményezi az útvonalválasztó protokoll futtatását a req_i szalagokra elhelyezett kérésekkel, melyekre a választ a res_i szalagokon kapja meg. Az ext_i szalagokra írt utasítások segítségével a támadó befolyásolni tudja, hogy a becsületes felhasználók mely csomópontokról indítsanak útvonalválasztást és mely csomópontok felé. Mivel azonban a támadó nem adaptív, ezeket a szalagokat csak a végrehajtás elején használhatja, s így az itt elhelyezett befolyásoló utasítások nem függnek a támadó által később megfigyelt üzenetektől.

Minden gép kezdetben inicializálódik valamilyen bemeneti adattal (például kriptográfiai kulcsokkal) és valamilyen véletlen bemeneti adattal (ami a véletlenszerű működéshez szükséges). A gépek reaktív módon működnek, azaz aktiválni kell őket ahhoz, hogy elvégezzék a feladatukat. Aktiválás során, egy gép először beolvassa a bemeneti adatokat a bemeneti szalagjairól, majd állapottranzíciók után valamilyen kimenetet generálhat a kimeneti szalagjaira. Ezután visszatér inaktív állapotba, s várja a következő aktiválást. A gépeket egy feltételezett, hipotetikus ütemező aktiválja menetekben, előre meghatározott sorrendben. A számításnak akkor van vége, amikor H eléri az elfogadó állapotát.

A valós modell kimenete a H -nak visszaadott útvonalak halmaza, amit adott $conf$ konfiguráció és A támadó esetén $real_out_{conf,A}(r)$ -rel jelölünk.

3. ábra Az ideális modell



Itt $r = (r_i, r_{M1}, \dots, r_{Mn-k}, r_{A1}, \dots, r_{Ak}, r_C)$ az egyes gépek véletlen bemeneteit tartalmazó vektor és r_i a kriptográfiai kulcsok generálásánál használt véletlen bemenet. $real_out_{conf,A}$ jelöli a kimenetet leíró valószínűségi változót, amikor r -et egyenletes eloszlás szerint választjuk.

2.6. Az ideális modell

Az ideális modellt a 3. ábra szemlélteti. Mint látható, az ideális modell felépítése hasonlít a valós modell felépítéséhez, ezért itt csak a különbségeket említjük:

- Mielőtt C' egy M_i gép in_i szalagjára helyezne egy útvonalválaszt ($rrep$), ellenőrzi, hogy az tartalmaz-e nem-plauzibilis útvonalat. Ha igen (és csak ekkor) megjelöli korruptként az üzenetet. Egyébként C' úgy működik mint C .
- Amikor M_i egy olyan útvonalválaszt kap, amelyhez tartozó útvonalkérést ő kezdeményezte, akkor az üzeneten először elvégzi a protokoll által megkövetelt ellenőrzéseket. Ha ezek sikeresek, akkor ellenőrzi, hogy az üzenet meg van-e jelölve korrupciós jelzővel. Ha igen, akkor M_i eldobja az üzenetet. Egyébként M_i úgy viselkedik mint M_i .

Az ideális modell kimenete a H -nak visszaadott útvonalak halmaza. A kimenetet $ideal_out_{conf,A}(r)$ -vel jelöljük, ahol r értelmezése a valós modell esetéhez hasonló. $ideal_out_{conf,A}$ jelöli a kimenetet leíró valószínűségi változót, mikor r -t egyenletes eloszlás szerint választjuk.

Mint az a leírásból látható, az ideális modellben H soha nem kap vissza olyan útvonalválaszt, mely nem-plauzibilis útvonalat tartalmaz. Az ideális modellt pontosan ebben az értelemben értjük ideálisnak.

2.7. A biztonságos útvonalválasztás formális definíciója

Az elkerülhetetlen támadásokkal kapcsolatos megjegyzéseket figyelembe véve, egy biztonságos útvonalválasztó protokollt azt várjuk el, hogy az csak elenyésző valószínűséggel adjon vissza nem-plauzibilis útvonalakat. Formálisan, a fent bevezetett két modell segítségével, és a szimulációs paradigma szellemében, ezt a következőképpen írhatjuk le:

Definíció:

Az útvonalválasztó protokoll akkor biztonságos, ha bármely $conf$ konfigurációhoz és bármely A valós támadóhoz létezik egy olyan A' ideális támadó, hogy $ideal_out_{conf,A}$ és $real_out_{conf,A'}$ eloszlások statisztikailag megkülönböztethetetlenek.

Vegyük észre, hogy nem az eloszlások pontos egyezését követeljük meg, hiszen azt a gyakorlatban egyetlen protokoll sem tudná garantálni, mivel a támadó, bár elenyésző valószínűséggel, de mindig képes sikeres támadást végrehajtani az alkalmazott kriptográfiai primitívek ellen (például megtippel egy digitális aláírást). Megjegyezzük még, hogy a fenti definíció gyengíthető, ha statisztikai megkülönböztethetlenség helyett számításelméleti megkülönböztethetlenséget követelünk meg. Jelen cikkben azonban erre nem lesz szükségünk.

3. Az endairA biztonsága

Ebben a fejezetben, az eddigieket egy rövid példán keresztül szeretnénk szemléltetni. Először röviden ismertetjük az endairA protokollt [2,1], majd bebizonyítjuk, hogy az biztonságos a fenti értelemben. A protokoll működését szemléltetik a következő üzenetek, ahol sig_x jelöli az x azonosítójú célcsofópont digitális aláírását az üzeneten, id pedig egy nem predikálható véletlen tranzakció-azonosító:

Útvonalkérés:

$S \rightarrow^* : [rreq, S, D, id, ()]$

$B \rightarrow^* : [rreq, S, D, id, (B)]$

$C \rightarrow^* : [rreq, S, D, id, (B, C)]$

Útvonalválasz:

$D \rightarrow C : [rrep, S, D, id, (B, C), (sig_D)]$

$C \rightarrow B : [rrep, S, D, id, (B, C), (sig_D, sig_C)]$

$B \rightarrow S : [rrep, S, D, id, (B, C), (sig_D, sig_C, sig_B)]$

Kezdetben az útvonalkeresést kezdeményező S csomópont küld egy $rreq$ útvonalkérést, ami tartalmazza saját és a D célcsofópont azonosítóját, egy frissen generált véletlen tranzakció-azonosítót, és egy üres azonosítólístát. Minden egyes közbelső csomópont, mely megkapja a kérést, hozzáfűzi saját azonosítóját a kérésben található azonosítólístához, majd a kérésűzenetet többszűrésessel (broadcast) továbbítja. Amikor a kérés eléri a célt, az egy $rrep$ útvonalválaszt generál. A válasz tartalmazza a kezdeményező és a cél azonosítóját, valamint a kérésben talált tranzakció-azonosítót és azonosítólístát (azaz a felfedezett útvonalat). A választ a célcsofópont digitális aláírásával látja el. A válaszüzenetet minden érintett csomópont az üzenetben található útvonal fordítottján továbbítja a megfelelő szomszédos csomópontnak. Továbbítás előtt azonban minden közbelső csomópont ellenőrzi a következőket:

- szerepel-e a saját azonosítója az útvonalban;
- a saját azonosítóját követő és megelőző azonosítók valóban szomszédos csomópontokhoz tartoznak-e;
- a D célcsofópont aláírása helyes-e.

Ha minden ellenőrzés sikeres, akkor a közbelső csomópont aláírja a válaszüzenetet, majd továbbítja azt. Ellenkező esetben eldobja a válaszüzenetet továbbítás nélkül. Amikor a kezdeményező megkapja a válaszüzenetet, ellenőrzi, hogy

- az útvonalban szereplő első azonosító valóban egy szomszédos csomópont azonosítója;
- az útvonalban szereplő minden azonosítóhoz található aláírás az üzenetben;
- minden aláírás helyes.

Ha minden ellenőrzés sikeres, akkor a kezdeményező elfogadja a válaszüzenetben kapott útvonalat.

Az alábbi tétel az endairA protokoll biztonságára vonatkozik:

Tétel:

Az endairA útvonalválasztó protokoll statisztikailag biztonságos, ha a felhasznált digitális aláíró séma biztonságos választott nyílt szövegű támadás esetén.

Bizonyításvázlat:

Egy útvonalválasztó protokoll statisztikailag biztonságos, ha nem-plauzibilis útvonalat csak elhanyagolható valószínűséggel ad vissza bármely *conf* konfigurációra és bármely *A* támadó esetén. Pontosabban azt kell belátni, egy *rrep* üzenet az ideális rendszerben elhanyagolható valószínűséggel kerül eldobásra a korrupciós jelző miatt.

Tegyük fel, hogy a következő üzenet eldobásra kerül a korrupciós jelzője miatt az ideális rendszerben, míg a valós modellben nem:

[*rrep*, *S*, *D*, *id*, (*N*₁, *N*₂, ..., *N*_p), (*sig*_{*D*}, *sig*_{*N*_p}, ..., *sig*_{*N*₁})]

Ekkor a következőket állapíthatjuk meg:

- nincs ismétlődő azonosító a $\pi = (S, N_1, N_2, \dots, N_p, D)$ útvonalban;
- *N*₁ csomópont *S* szomszédja;
- minden aláírás érvényes;
- *S* és *D* megbízható csomópontok;
- minden köztes csomópont (nagy valószínűséggel) azt az útvonalat látja amit *D* küldött (azaz π -t), mivel *D* azt aláírta, és ezt mindegyik csomópont ellenőrzi;
- ennek ellenére π egy nem plauzibilis útvonal $\underline{G}$ gráfban, ahol $\underline{G}$ a redukált konfiguráció gráfja.

Bebizonyítjuk, hogy ez csak úgy lehetséges ha az *A* támadó hamisította legalább egy megbízható csomópont aláírását. Tudjuk, hogy a redukált konfiguráció gráfjában nincsenek szomszédos támadó csomópontoknak megfelelő szomszédos csúcsok. Így a π útvonal egyértelműen particionálható részsorozatokra a plauzibilis útvonal definíciója szerint mégpedig úgy, hogy minden megbízható csomópontoz tartozó azonosító külön partíciót alkot, míg a szomszédos kompromittált azonosítók szintén külön partíciót alkotnak. Ha a π útvonal nem plauzibilis, akkor a következő állítások közül legalább egy igaz:

1. $P_j = \{N_j\}$ és $P_{j+1} = \{N_{j+1}\}$ megbízható csomópontokhoz tartozó partíciók, mely csomópontok nem szomszédosak $\underline{G}$ gráfban.
2. $P_j = \{N_j\}$ egy megbízható *v* csomópontoz tartozó partíció, míg P_{j+1} egy *v'* támadó csomópontoz tartozó partíció, és az *N_j* azonosítóhoz tartozó *v* csomópontnak
 - a) nincs szomszédja $\underline{V}^*$ csúcshalmazban
 - b) van szomszédja $\underline{V}^*$ csúcshalmazban.
3. P_j egy *v'* támadó csomópontoz tartozó partíció, míg $P_{j+1} = \{N_j\}$ egy megbízható *v* csomópontoz tartozó partíció, és az *N_j* azonosítóhoz tartozó *v* csomópontnak
 - a) nincs szomszédja $\underline{V}^*$ csúcshalmazban
 - b) van szomszédja $\underline{V}^*$ csúcshalmazban.

Az 1., 2.a, és 3.a esetekben *v* becsületes csomópont detektálja, hogy *v'* egy nem szomszédos csomópont, így nem írja alá az üzenetet. Ezért *S* csak úgy kaphatta ezt az üzenetet, hogy a támadó hamisította az *N_j*-hez tartozó aláírást. 2.b és 3.b esetben *v* szom-

szédja legyen *q*, ahol *q* támadó csomópont, melynek azonosítója legyen *v'* azonosítója. *q* és *v'* nem lehetnek szomszédosak, ezért *q* csak úgy járhat sikerrel, ha egy új üzenetet konstruál és így meghamisítja *D* aláírását.

Ezért arra jutottunk, hogyha az ideális modellben nem-elenyésző valószínűséggel fordul elő egy ilyen válaszüzenet, akkor a támadó nem-elenyésző valószínűséggel hamisítani képes valamelyik becsületes résztvevő aláírását, ez pedig ellentmond azon feltevésünknek, hogy az alkalmazott aláírás séma biztonságos.

4. Összegzés

A cikkben egy olyan modellt mutattunk be, melyben precízen definiálható, hogy mit értünk biztonságos útvonalválasztás alatt, és így lehetőség nyílik annak igazolására (vagy cáfolására), hogy egy adott, igény szerinti, forrás alapú, ad hoc útvonalválasztó protokoll biztonságos-e. A módszer gyakorlatban történő használatát egy példán keresztül szemléltettük, melyben igazoltuk, hogy az endairA protokoll biztonságos.

A jövőben hasonló modellt szeretnénk kidolgozni a táblázat alapú ad hoc útvonalválasztó algoritmusokra is (például SEAD, ARAN, S-AODV). Továbbá szeretnénk a bizonyításokat automatizálni valamilyen alkalmas formális nyelv (például processz algebra) és kapcsolódó ellenőrző eszköz segítségével.

Köszönetnyilvánítás

A bemutatott munka támogatói a következők voltak: OM (BÖ2003/70), OTKA (T046664), és IKMA.

Irodalom

- [1] Ács G.,
Ad hoc útvonalválasztó protokollok bizonyított biztonsága, TDK dolgozat, 2004. nov. 9.
- [2] L. Buttyán, I. Vajda,
Towards provable security for ad hoc routing protocols, In Proceedings of the 2nd ACM Workshop on Security of Ad Hoc and Sensor Networks, Washington DC, USA, October, 2004.
- [3] Y.-C. Hu, A. Perrig,
A survey of secure wireless ad hoc routing, IEEE Security and Privacy Magazine, June 2004.
- [4] O. Goldreich,
The Foundations Of Modern Cryptography, Cambridge University Press, 2001.
- [5] M. Bellare, R. Canetti, H. Krawczyk,
A modular approach to the design and analysis of authentication and key exchange protocols, In Proceedings of the ACM Symposium on the Theory of Computing, 1998.