



Budapesti Műszaki és Gazdaságtudományi Egyetem
Méréstechnika és Információs Rendszerek Tanszék

Tranzíciós rendszerek analízise absztrakciós módszerekkel

Farkas Rebeka I. évf, (MSc) mérnökinformatikus szakos hallgató

Konzulensek: Vörös András tudományos segédmunkatárs, MIT

Tóth Tamás doktorandusz, MIT

Kritikus rendszerek specializáció

Önálló laboratórium 1. összefoglaló

2014/15. II. félév

A biztonságkritikus rendszerek esetében elvárás a matematikailag bizonyíthatóan hibamentes működés. Modell alapú fejlesztés segítségével még az implementáció előtt eldönthető egy rendszerről, hogy helyesen van-e megtervezve.

A BSc képzés során az önálló laboratórium, a TDK dolgozat, valamint a szakdolgozatom keretein belül is modellezéshez használt formalizmusok – így Petri-hálók, valamint a valamivel nagyobb leíró, illetve kifejezőerővel rendelkező Lineáris Tranzíciós Rendszerek (LTR) – elérhetőségét vizsgáltam absztrakt interpretáció segítségével. A Petri-háló és az LTR csak egy speciális esete az összefoglaló néven tranzíciós rendszereknek nevezett formalizmusnak, ahogy az absztrakt interpretáció is csak egy az elérhetőségi analízishez felhasználható absztrakciós módszerek közül.

Az MSc Képzés során az elsődleges cél egy olyan komplex algoritmus kifejlesztése, amely tranzíciós rendszerek egy kiterjedt csoportjának elérhetőségi analízisére alkalmas.

A félév során először irodalomkutatót végeztem, hogy minél több absztrakt domaint (absztrakt interpretáció által állapotter-reprezentációra használt tartomány) és absztrakciós módszert megismerjek. Így megismertem a három értékű logikán alapuló *ternary simulation*-t, kétértékű változók absztrakciójához; *zóna*, és *régió* absztrakt domaint óraváltozók analíziséhez, Gauge domaint egymásba ágyazott ciklusokban módosított változók értékének becslésére; valamint megvizsgáltam, hogyan lehet változók tömbjeit hatékonyan reprezentálni.

Emellett az állapotter csökkentéséhez megismertem, kiterjesztettem általános tranzíciós rendszerekre, és implementáltam a COI algoritmust. A *Cone of Influence reduction* algoritmus alapja az a felismerés, hogy az elérhetőségi analízis illetve az állapotterbecslés célja általában nem annak a kiderítése, hogy egy konkrét állapot (egy konkrét változó értékadás) elérhető-e, hanem állapotok egy halmazának elérhetőségére kíváncsi a felhasználó. Nem érdemes tehát az egész rendszert alávetni ezen vizsgálatoknak, ki kell választani azokat a részeket, amelyek valamilyen módon (akár közvetetten is!) tudják befolyásolni a kérdéses változók értékeit. Erre alkalmazható a COI.

Megismertem az extrapolációs operátorral, előnyeivel, hátrányaival, egy lehetséges felhasználási módjával, valamint implementáltam egy alternatív megoldást. Az extrapolációs operátor az absztrakt interpretációs algoritmus során alkalmazott widening operátor miatt megjelenő túlzott felső becslést hivatott elkerülni, azonban ennek az ára, hogy az algoritmus terminálódása már nem lesz garantálva. Ennek ellenére érdemes foglalkozni vele, mivel segítségével megvalósítható ellenpélda-alapú absztrakció finomítás, ami egy nagyon elegáns és hatékony módja az elérhetőségi analízisnek.