



The CLEAN PAPER

WHAT THE PAPER SAYS · WHAT IT DOESN'T · WHY IT MATTERS

Կրիպտոգրաֆիկ ապացույցը կարող է թաքցնել գաղտնիքը՝ simulator-ի բացակայությունը դժվար ապացուցելի դարձնելով

Zero-knowledge proofs-ը թույլ են տալիս որևէ պնդում ապացուցել՝ առանց witness-ը բացահայտելու: Դասական տեսությունն ասում է, որ լիարժեք իմաստով դա հնարավոր չէ մեկ հաղորդագրությամբ, առանց վստահելի setup-ի և perfect soundness-ով: Rahul Ilango-ի paper-ը այդ անհնարինությունը չի վերացնում: Այն փոխում է թիրախը. simulator-ի իրական գոյությունը պահանջելու փոխարեն հարցնում է՝ արդյոք ընտրված proof system-ը կարող է արդյունավետորեն ապացուցել, որ simulator գոյություն չունի: Proof complexity-ի և cryptography-ի խոշոր assumptions-ի ներքո սա բավական է zero-knowledge-ի falsifiable, game-based հետևանքները property-by-property վերականգնելու համար: Սա ինտերնետում անմիջապես կիրառելի primitive չէ, այլ proof-theoretic եղանակ՝ mathematical unprovability-ն cryptographic cover-ի վերածելու համար:

Written by Cinzia Vaglio · Cross-checked by Vera Rubin · Edited by Michele Renda · Figures by Nova Vela · AI-assisted, human-reviewed

Paper: Gödel in Cryptography: Effectively Zero-Knowledge Proofs for NP with No Interaction, No Setup, and Perfect Soundness, Rahul Ilango, FOCS 2025 / IACR ePrint 2025/1296 · DOI: 10.1109/FOCS63196.2025.00058

Հնարքը գաղտնիքի թաքնված լինելը ապացուցելու է

Սկսենք գրոյական գիտելիքի ամենապարզ տարբերակից:

Այլսը ցանկանում է Բորին համոզել, որ Սուդոկուի գլուխկոտրուկը լուծում ունի: Եթե նա ուղարկի լուծումը, Բորը կհամոզվի, բայց գլուխկոտրուկը կկորցնի իմաստը: Այլսին պետք է ավելի տարօրինակ բան՝ ապացույց, որ լուծում գոյություն ունի, առանց լուծումը բացահայտելու:

Սա է **գրոյական գիտելիքի ապացույց**-ի՝ գրոյական գիտելիքի ապացույցի խոստումը: Ապացուցողը՝ Այլսը, համոզում է ստուգողին՝ Բորին, որ պնդումը ճիշտ է, բայց չի բացահայտում պնդման ճշմարտությունից ավել ոչինչ:

Խնդիրն այն է, որ այս խոստումը ինչոր գին ունի: Սովորական մաթեմատիկական ապացույցը երկու հարմար հատկություն ունի: Այն **մեկ հաղորդագրություն** է՝ գրում ես, հանձնում ու հեռանում: Եվ այն կարող է լինել **կատարյալ հուսալի**՝ կեղծ պնդումը որևէ վավեր ապացույց չունի: Դասական անհնարինության արդյունքները ցույց են տալիս, որ գրոյական գիտելիքի ստիպված է հրաժարվել այս հատկություններից երկուսից էլ, և ոչ միայն դրանց համակցությունից. առանձին վերցրած յուրաքանչյուրն էլ անհամատեղելի է լիարժեք դասական գրոյական գիտելիքի հետ:

Առաջինը՝ գրոյական գիտելիքի ապացույցին պետք է փոխգործակցություն: Եթե Այլսը մեկ հաղորդագրություն է ուղարկում՝ առանց նախապես կազմակերպված վստահելի կարգավորման, դասական գրոյական գիտելիքի երաշխիքը փլվում է՝ անկախ նրանից, թե հուսալիությունի ինչքան զիջում եք պատրաստ ընդունել:

Երկրորդը՝ գրոյական գիտելիքի ապացույցին փոքր սխալի հանդուրժողականություն է պետք: Կատարյալ հուսալիություն պահանջելը, պարզվում է, հանգիստ կերպով վերացնում է նաև փոխգործակցությունը. եթե ստուգողը երբեք չի կարող խաբվել՝ անկախ իր պատահական ընտրություններից, ապա կարող է այդ ընտրությունները նախապես ֆիքսել: Իսկ երբ ստուգողը կանխատեսելի է, Այլսը կարող է բոլոր հարցերի պատասխանները մեկանգամից ուղարկել մեկ հաղորդագրությամբ՝ հենց այն դեպքը, որն արդեն անհնար էր դարձել:

Rahul Ilango-ի հոդվածը կրկնակի պատի շուրջ ճանապարհ է փնտրում: Ոչ թե ձևացնելով, որ պատը չկա, և ոչ էլ անհնար պայմաններում դասական գրոյական գիտելիքի կառուցելով: Ծարժումը ավելի նուրբ է. թուլացնել «ոչինչ չի բացահայտում» պահանջը, բայց այնպես, որ պահպանվեն այն անվտանգության հատկությունները, որոնք կրիպտոգրաֆները իրականում կարող են փորձարկել:

Արդյունքը կոչվում է **արդյունավետորեն գրոյական գիտելիքի**՝ արդյունավետորեն գրոյական գիտելիք:

A proof without leaking the witness

The paper keeps the ordinary-proof shape by weakening what privacy must prove.

blocked door 1

interaction

blocked door 2

trusted setup

blocked door 3

imperfect
soundness

the route

the rulebook cannot
efficiently refute the
simulator

Boundary: not classical zero-knowledge; effectively zero-knowledge under a chosen proof system.

Չրոյական գիտելիքի ճանապարհը փակվում է երեք դռների մոտ՝ փոխգործակցություն, վստահելի կարգավորում և ոչ կատարյալ հուսալիություն: Լանգո-ի կառուցումը անցնում է ուրիշ դռնով. ընտրված կանոնագիրքը չի կարող արդյունավետորեն հերքել սիմուլյատորի գոյությունը:

Original diagram — The Clean Paper CC BY 4.0

Classical privacy vs effective privacy

The relaxation is the point: the paper changes what the privacy claim has to establish.

classical zero-knowledge

positive claim

A simulator exists and can fake the
verifier's view without the witness.

effectively zero-knowledge

weaker claim

Your chosen proof system cannot
efficiently prove that no simulator
exists.

Boundary: the construction preserves testable consequences, not the full simulator guarantee.

Դասական գրոյական գիտելիքիը հարցնում է՝ արդյոք սիմուլյատոր գոյություն ունի: «արդյունավետորեն գրոյական գիտելիքի»-ը հարցնում է միայն՝ արդյոք քնտրված կանոնագիրքը կարող է արդյունավետորեն ապացուցել, որ սիմուլյատոր գոյություն չունի: Հենց այս ավելի թույլ հարցն է թույլ տալիս պահպանել մեկ հաղորդագրությունը, կարգավորման բացակայությունն ու կատարյալ հուսալիությունը:

Original diagram — The Clean Paper CC BY 4.0

Հին թեստը. սիմուլյատոր իսկապես գոյություն ունի

Չրոյական գիտելիքիը ձևակերպելու դասական եղանակը օգտագործում է մտացածին օգնական՝ **սիմուլյատոր**:

Գաղափարը սա է. պատկերացրեք Ձեյնին, որը **չգիտի** Ալիսի գաղտնիքը: Եթե Ձեյնը կարող է ինքնուրույն ստեղծել այնպիսի ապացույցներ, որոնք արտաքինից չեն տարբերվում Բորի՝ Ալիսից ստացած ապացույցներից, ապա Ալիսի ապացույցները Բորին նոր բան չեն սովորեցրել: Ձեյնը նույն փորձառությունը կարող էր կեղծել առանց Ալիսի գաղտնիքի:

Ուստի դասական գրոյական գիտելիքիը պահանջում է իրական սիմուլյատոր: Պետք է գոյություն ունենա արդյունավետ ալգորիթմ, որը կարող է գաղտնիքը չիմանալով ստեղծել համոզիչ կեղծ ապացույցներ: Գաղտնիքը տեխնիկական լեզվով **վկա**-ն է՝ վկայությունը: Սուդոկոլի դեպքում վկաը պարզապես լրացված ճիշտ ցանցն է:

Այս սահմանումը շատ ուժեղ է, բայց հենց այստեղ էլ գործում է դասական անհնարինությունը: Ինտուիցիան պարզ է. իսկապես ոչ ինտերակտիվ ապացույցը պարզապես տող է: Երբ Բորը ստանում է այդ տողը, կարող է այն ցույց տալ ուրիշներին. նա արդեն ձեռք է բերել պնդումը ուրիշներին ապացուցելու ունակություն, և դա ինքնին «ոչինչ չսովորելուց» ավել է թվում: Դասական թեորեմներն այս ինտուիցիան դարձնում են վերևում նկարագրված անհնարինություններ:

Երեք հատկություն, որոնցից այս հոդվածը չի հրաժարվում

Հոդվածի վերնագրում նշված են երեք սահմանափակումներ:

Փոխգործակցություն չկա. Ալիսը ուղարկում է մեկ ապացույց string: Չկա հետուդարձով արձանագրություն:

Կարգավորում չկա. Ալիսն ու Բորը չեն հենվում վստահելի տարածված հղումային string-ի կամ նախապես պատրաստված հանրային պատահականության վրա: «ոչ ինտերակտիվ գրոյական գիտելիքի» կոչվող շատ համակարգեր կարգավորում ունեն: Այստեղ խոսքը իրականում zero կարգավորման մասին է:

Կատարյալ հուսալիություն. Կեղծ պնդումը վավեր ապացույց չունի: Ոչ թե

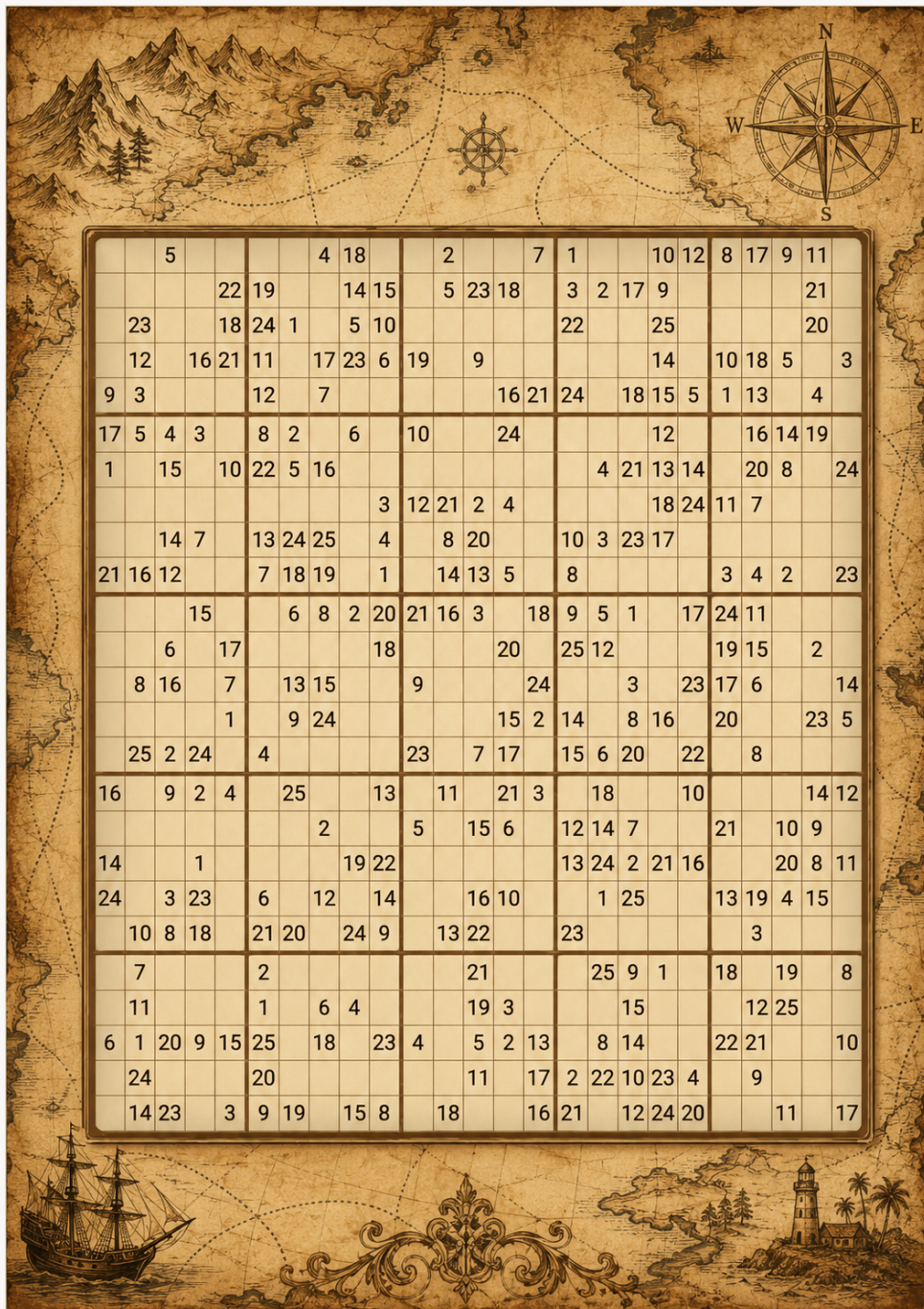
«գրեթե երբեք չի ընդունվի», այլ՝ վավեր ապացույց ընդհանրապես գոյություն չունի:

Սրանք հենց սովորական գրված մաթեմատիկայի երեք հատկություններն են, և դասական գրոյական գիտելիքի, ինչպես վերևում նշվեց, չի կարող դրանք բոլորը պահպանել:

MegaSudoku-ով զգալ տարբերությունը

Ահա միտումնավոր պարզեցված եղանակ՝ տարբերությունն զգալու համար:

Անալոգիայի լուրջ մասի համար սովորական 9×9 Սուդոկու մի օգտագործեք: Այն չափազանց փոքր ու վերջավոր է. համակարգիչը պարզապես կարող է լուծել այն կամ ապացուցել, որ լուծում չունի: Փոխարենը պատկերացրեք **MegaSudoku(n)** գլոխկոտրուկների ընտանիք: Ընդլայնենք կանոնը. ընտրենք բլոկի չափ n , դնենք $N = n^2$ և կառուցենք $N \times N$ ցանց, որը բաժանված է $n \times n$ բլոկների և ունի N նշան: Սովորական Սուդոկուն պարզապես փոքր $n = 3$, $N = 9$ դեպքն է՝ 9×9 ցանց, 3×3 բլոկներ և ինը նշան: ապացույց բարդությունի պատմությունը սկսվում է միայն այն ժամանակ, երբ n -ը կարող է աճել, և ցանցում կարելի է ավելացնել գաջերներ, որոնք այն դարձնում են Սուդոկուի տեսքով SAT բանաձև: SAT բանաձևը պարզապես այո/ոչ սահմանափակումների ցանկ է. հնարավոր է փոփոխականներին ճիշտ/սխալ արժեքներ տալ այնպես, որ բոլոր սահմանափակումները բավարարվեն:



25×25 Սուդոկու. կանոնները կարելի է ստուգել առանց ավարտված ցանցը բացահայտելու՝ որպես թաքնված լուծումը՝ վկալ, հաստատող ապացույցի տեսողական անալոգիա:

AI-generated editorial thumbnail — The Clean Paper CC BY 4.0

Sudoku-ն և SAT-ը. նույն խնդիրը երկու զգեստով

Այն պնդումը, որ Սուդոկուն կարող է «SAT բանաձևի պես վարվել», պարզապես փոխաբերությունն է: Թարգմանությունը գործում է երկու ուղղությամբ, և հեշտ ուղղությունը կարելի է ամբողջությամբ գրել:

Սուդոկուից դեպի SAT: SAT-ը խոսում է միայն ճիշտ/սխալ լեզվով, ուստի

յուրաքանչյուր (տող, սյուն, արժեք) եռյակի համար ներմուծենք մեկ բուլյան փոփոխական. $x(r,c,v)$ նշանակում է « r տողի, c սյան բջիջը պարունակում է v արժեքը»: 4×4 Սուդոկուին՝ 2×2 բլոկներով և 1-4 արժեքներով, պետք է $4 \cdot 4 \cdot 4 = 64$ փոփոխական, իսկ դասական 9×9 Սուդոկուին՝ 729: Այնուհետև Սուդոկուի յուրաքանչյուր կանոն վերածվում է դիզյունկոնյունկտիվ խմբի: դիզյունկոնյունկտիվ փոփոխականների կամ դրանց ժխտումների OR-ն է, իսկ ամբողջ բանաձևը բոլոր դիզյունկոնյունկտիվ AND-ն է:

Յուրաքանչյուր բջիջ ունի առնվազն մեկ արժեք՝ մեկ դիզյունկոնյունկտիվ յուրաքանչյուր բջջի համար.

$$x(1,1,1) \vee x(1,1,2) \vee x(1,1,3) \vee x(1,1,4)$$

Յուրաքանչյուր բջիջ ունի առավելագույնը մեկ արժեք՝ արժեքների յուրաքանչյուր զույգի համար «ոչ երկուսը միաժամանակ» դիզյունկոնյունկտիվ.

$$\neg x(1,1,1) \vee \neg x(1,1,2) \quad \neg x(1,1,1) \vee \neg x(1,1,3) \quad \dots \text{ և այդպես բոլոր վեց զույգերի համար:}$$

Յուրաքանչյուր տող պարունակում է յուրաքանչյուր արժեքը: Առաջին տողի և 3 արժեքի համար՝ առնվազն մեկ անգամ.

$$x(1,1,3) \vee x(1,2,3) \vee x(1,3,3) \vee x(1,4,3)$$

և առավելագույնը մեկ անգամ՝ $\neg x(1,1,3) \vee \neg x(1,2,3)$, և նույնը տողի բջիջների յուրաքանչյուր զույգի համար:

Մյուսներն ու բլոկները՝ նույնատիպ խմբեր, փոխվում է միայն բջիջների խումբը: Վերին ձախ բլոկի և 2 արժեքի համար.

$$x(1,1,2) \vee x(1,2,2) \vee x(2,1,2) \vee x(2,2,2)$$

գումարած զույգ առ զույգ «ոչ երկուսը միաժամանակ» դիզյունկոնյունկտիվ:

Տպված հուշումները ամենապարզ մասն են. յուրաքանչյուր հուշում դառնում է մեկ փոփոխական պարունակող դիզյունկոնյունկտիվ: Վերին ձախ անկյունում տպված 3-ը դառնում է.

$$x(1,1,3)$$

Այս ամենի AND-ը բավարարելի է ճիշտ այն դեպքում, երբ Սուդոկուն լուծում ունի, իսկ satisfying արժեքագրումը հենց լուծումն է. կարդում ենք, թե որ $x(r,c,v)$ -երն են ճիշտ, և լրացնում ցանցը: 9×9 Սուդոկուի համար ստացվում է 729 փոփոխական և մի քանի հազար դիզյունկոնյունկտիվ, ինչը ժամանակակից SAT solver-ը լուծում է միլիվայրկյաններում: Ուշադրություն դարձրեք հուշում դիզյունկոնյունկտիվ $x(1,1,3)$ -ին. այն ասում է «այս բջիջը ճիշտ 3 է», ոչ թե «այս բջիջները բոլորը տարբեր են»: Հենց այս անհամաչափությունն է ստորև ներկայացված արձանագրության note-ում ստիպելու լրացուցիչ հնարք օգտագործել հուշում բջիջների համար:

SAT-ից դեպի Սուդոկո: Հոդվածին պետք է հակառակ, ավելի դժվար ուղղությունը. կամայական SAT բանաձևից կառուցել mega-Սուդոկո, որը լուծում ունի ճիշտ այն դեպքում, երբ բանաձևը բավարարելի է: Սուդոկոյի բնիկ կանոնները կարող են միայն ասել «այս բջիջները բոլորը տարբեր են», ուստի կամայական տրամաբանական սահմանափակումները պետք է *կառուցել*: Հենց սա են անում գաղտնիքները: Գաղտնիք փոքր, նախապես կառուցված բջիջների կլաստեր է՝ մեկ հատ բանաձևի յուրաքանչյուր դիգյունկտի համար, որտեղ ընտրված բջիջները խաղում են փոփոխականների դերը՝ դրանց նշանը կոդավորում է ճիշտ կամ սխալ, իսկ ներքին սահմանափակումներն այնպես են նախագծված, որ օրինական լրացումները համապատասխանեն միայն տվյալ դիգյունկտը բավարարող արժեքաբաշխումներին: Սա NP-completeness ապացույցների դասական տեխնիկա է: Generalized Սուդոկոյի համար այն 2003-ին կառուցել են Yato-ն և Seta-ն:

Երկու ուղղությունները միասին ասում են, որ $N \times N$ Սուդոկոն և SAT-ը նույն խնդրի երկու տարբեր ներկայացումներ են: Հենց սա է թույլ տալիս և՛ այս հոդվածին, և՛ հոդվածին ամբողջ NP-ի մասին պատմել ցանցերի ու նշանների լեզվով:

վկար դեռ հեշտ է պատկերացնել: Այլսը գիտի mega-Սուդոկոյի ամբողջական, վավեր լրացումը: Բոլոր ցանկանում է համոզվել, որ այդպիսի լրացում գոյություն ունի, բայց Այլսը չի ցանկանում այն բացահայտել: Եթե նա ուղարկի ամբողջ ցանցը, Բոլոր կհամոզվի, բայց գաղտնիքը կկորչի:

Դասական գրոյական գիտելիքի տարբերակում Այլսն ու Բոլոր փոխգործակցում են: Հին ոճի մտավոր մոդելներից մեկը օգտագործում է փակված խաղաքարեր: Այլսը ծածկում է լուծված ցանցը, յուրաքանչյուր փուլից առաջ գաղտնի վերանվանում նշանները, իսկ Բոլին թույլ է տալիս ստուգել պատահական ընտրված մեկ տեղային սահմանափակում՝ տող, սյուն, բլոկ կամ գաղտնիք: Եթե բացված բջիջները ցույց են տալիս բոլորը տարբեր նշաններ, Բոլի վստահությունը մեծանում է: Այնուհետև ամեն ինչ կրկին ծածկվում է, և նշանները նորից են պատահական վերանվանվում: Մի դժվարություն կա. տպված հուշումները լրացուցիչ հնարք են պահանջում, որովհետև նշանների վերանվանումը դրանք էլ է թաքցնում: Ստորև note-ը բացատրում է դասական արձանագրությունների լուծումը:

Ինչպես են դասական արձանագրություններն իրականում աշխատում clue բջիջների հետ

Վերանվանման հնարքն ունի կույր կետ: Տողի, սյան և բլոկի կանոնները ասում են «այս բջիջները բոլորը տարբեր են», և *բոլորը տարբեր* հատկությունը պահպանվում է նշանների ցանկացած վերանվանման դեպքում: Բայց հուշումն ասում է «այս բջիջում ճիշտ 5 է»: Վերանվանումից հետո Բոլը տեսնում է միայն $\sigma(5)$ ՝ դիմակավորված ինչոր նշան, բայց σ վերանվանումը չգիտի: Նա ոչինչ չի կարող ստուգել: Եթե սա չտկվի, Այլսը կարող է ապացուցել, որ *որևէ* վավեր

ցանց գոյություն ունի՝ ամբողջությամբ անտեսելով տպված հուշումները, ինչը տվյալ կոնկրետ գլուխկոտրուկի մասին ոչինչ չի ասացուցում: Դասական գրականությունն ունի երկու ստանդարտ լուծում:

Palette-ը: Թաքնված ցանցին ավելացվում է N բջիջներից մեկ լրացուցիչ տող՝ palette, որը Ալիսը լրացնում է $1...N$ նշաններով ֆիքսված հանրային հերթականությամբ և վերանվանում է մնացած ամեն ինչի հետ, այնպես որ այնտեղ հայտնվում են $\sigma(1)...\sigma(N)$: Բոբի պատահական փորձարարական վարակումը հիմա ունի մեկ լրացուցիչ տարբերակ: Տող, սյուն, բլոկ կամ գաջեր՝ բացելու փոխարեն նա կարող է ընտրել **palette-ը և մեկ հուշում բջիջ**: Ալիսը բացում է երկուսն էլ: Palette-ը բացահայտում է տվյալ փուլի վերանվանումը, և Բոբը ստուգում է, որ հուշում բջիջը ցույց է տալիս տպված հուշումի հենց վերանվանված տարբերակը: Սա շարունակում է լինել գրոյական գիտելիքի, որովհետև Բոբը սովորում է միայն σ -ն, որը յուրաքանչյուր փուլում նոր պատահական permutation է և ինքնուրույն ոչինչ չի ասում, ինչպես նաև այն բջջի արժեքը, որը գլուխկոտրուկից արդեն գիտեր: Գաղտնի բջիջներից ոչինչ չի արտահոսում, իսկ սիմուլյատորը կարող է պատկերը կեղծել պարզապես պատահական σ ընտրելով: հուսալիությունը պահպանվում է, որովհետև խաբող Ալիսը յուրաքանչյուր փուլում ֆիքսված հավանականությամբ բռնվում է, և փուլերը կրկնվում են այնքան, մինչև կասկածը աննշան դառնա:

հուշումների վերացում կառուցմամբ: Ավելի կառուցվածքային տարբերակում հատուկ փորձարարական վարակումը հանվում է: հուշումի արժեքը *ստուգելու* փոխարեն այն պարտադրվում է տարբերության սահմանափակումներով. հուշում բջիջը կապվում է palette-ի բոլոր բջիջների հետ, բացի այն բջջից, որը կրում է հենց իր արժեքը՝ «տարբեր $\sigma(1)$ -ից, տարբեր $\sigma(2)$ -ից, ..., տարբեր ամեն ինչից բացի $\sigma(5)$ -ից»: Միակ նշանը, որը բջիջը կարող է օրինականորեն ունենալ, հուշումի նշանն է: Բոլոր սահմանափակումները նորից «այս երկուսը տարբեր են» տեսակի են՝ invariant վերանվանման նկատմամբ և ստուգելի ճիշտ ինչպես տողը: Նույն հնարքը դասական graph-coloring արձանագրությունում օգտագործվում է նախապես գունավորված գազաթների համար: Սա նաև վերևում նշված *գաջեր* գաղափարի ոգին է. MegaSudoku-as-SAT պատկերում հուշումները կոմպիլացվում են inequality գաջերների մեջ, ինչպես ցանկացած այլ սահմանափակում:

Ֆիզիկական արձանագրությունը: Սուդոկուի իրական քարտային գրոյական գիտելիքի արձանագրությունը (Gradwohl, Naor, Pinkas և Rothblum, 2007) ընդհանրապես վերանվանում չի օգտագործում և հուշումները հաստատում է մինչև թաքցնելը: Յուրաքանչյուր բջջի համար Ալիսը դնում է նույն արժեքով երեք նույնական քարտ՝ գաղտնի բջիջների համար երեսով ներքև, բայց **հուշում բջիջների համար երեսով վերև**, որպեսզի Բոբն իր աչքով տեսնի, որ տպված հուշումները հարգված են, մինչև քարտերը շրջվեն: Այնուհետև յուրաքանչյուր բջջից մեկ քարտ մտնում է իր տողի փաթեթ, մեկը՝ սյան, մեկը՝ բլոկի: Յուրաքանչյուր փաթեթ խառնվում և բացվում է, և Բոբը ստուգում է, որ այն պարունակում է բոլոր N նշանները: Խառնումը ոչնչացնում է դիրքային տեղեկատվությունը՝ սա է գրոյական գիտելիքի մասը,

իսկ հուշումները արդեն հաստատված են բաժանման պահին:

Երկու դեպքում էլ դասը նույնն է. գրոյական գիտելիքի արձանագրությունը շատ խիստ հաշվառում է, թե *որ փաստերն* են դիմանում թաքցմանը: Վերանվանումը պահպանում է «բոլորը տարբեր են»-ը և ջնջում «հավասար է 5»-ը, ուստի «հավասար է 5»-ը պետք է այլ ճանապարհով վերադարձնել:

Սա հոդվածի արձանագրությունը չէ: Սա դասական գրոյական գիտելիքի մտավոր մոդելն է.

- Այնպես որ Բորը հետադարձով փոխգործակցում են:
- Բորը պատահական checks է ընտրում:
- Այնպես բացահայտում է միայն տեղային համահունչությունը, ոչ ամբողջ լուծումը:
- գաղտնիությունի ապացույցը ցույց է տալիս, որ Բորի տեսած պատկերը կարելի էր ստեղծել առանց Այնի գաղտնի լուծման:

Այսպիսով դասական գրոյական գիտելիքը կառուցված է դրական փաստի շուրջ.

սիմուլյատոր իսկապես գոյություն ունի:

Հիմա հեռացնենք հարմար մասերը: Այնպես ուղարկում է մեկ ապացույց string և հեռանում: Չկա վստահելի կարգավորում, նախապես պատրաստված ընդհանուր պատահական string, իսկ Բորը երբեք չպետք է ընդունի կեղծ գլուխկոտրուկ: Սա հենց այն միջավայրն է, որտեղ դասական գրոյական գիտելիքը չի գոյատևում:

Հնարքի համար ևս մեկ կերպար է պետք: Ֆիքսենք **կանոնագիրք**՝ logician-ի իմաստով ֆորմալ ապացույց համակարգ՝ արքիտմների ֆիքսված հավաքածու և գրավոր մաթեմատիկական ապացույցները մեխանիկորեն ստուգելու կանոններ: ZFC-ը՝ մաթեմատիկայի ստանդարտ արքիտմները, canonical օրինակն է: Այս պահից ամեն պնդում հարաբերական է նախապես ընտրված կանոնագիրքին, բայց ընտրությունը ճկուն է. կառուցումը աշխատում է ցանկացած ֆիքսված կանոնագիրքի համար, ներառյալ ZFC-ը:

(հոդվածի բառապաշարից կարևոր տարբերակում. այստեղ **ապացույց համակարգ** միշտ նշանակում է հենց այս կանոնագիրքը՝ ֆորմալ համակարգը, որը ստուգում է մաթեմատիկական ապացույցները, ոչ երբեք Այնի ուղարկված հաղորդագրությունները: Այնի և Բորի մեքենաներն անվանվում են ապացուցող և ստուգող:)

Գյոդելյան տարբերակը պահպանում է MegaSudoku-ի պատմությունը, բայց փոխում ապացույցը:

Ընտրենք նույն ներկայացվող չափի երկրորդ սահմանափակում համակարգ, որը կոչենք **D**: Պատմության մեջ S-ն ու D-ն նույն ձևաչափի երկու MegaSudoku(n) գլուխկոտրուկ են: Կուլիսներում D-ն կարող էր սկսել այլ չափի դժվար տրամաբանական բանաձևից. անհրաժեշտության դեպքում այն կարելի է harmless dummy սահմանափակումներ երկարացնել, որպեսզի նույն ցանցի չափին համապատասխանի: D-ն կառուցվում

Է մի բանաձևից, որը իրականում **անբավարարելի** է. չկա արժեքաբաշխում, որը բավարարի բոլոր սահմանափակումները, ինչպես կոտրված գլուխկոտրուկը չունի օրինական լրացում: Խաղալիք օրինակ կլինեն բանաձև, որը միաժամանակ պահանջում է «X-ը ճիշտ է» և «X-ը սխալ է»: Այսինքն՝ D-ն լուծում չունի:

Բայց D-ն չպետք է լինի այնպիսի կոտրված գլուխկոտրուկ, որի կոտրված լինելը *հեշտ է ցույց տալ*: «X և ոչ-X» խաղալիք օրինակը ձախողվում է, որովհետև ցանկացած կանոնագիրք այն մեկ տողով կհերքի: D-ն պետք է կեղծ լինի այնպիսի ձևով, որ ընտրված կանոնագիրքը չկարողանա կարճ ապացույցով հաստատել դրա անհնար լինելը: Եթե կանոնագիրքը կարճ ապացույցով կարողանար հերքել D-ն, ստորև ներկայացված կառուցումը կփլվեր. այն այլընտրանքային ճանապարհը, որը կարող էր ապացույց ստեղծել առանց Ալիսի գաղտնիքի, ֆորմալ կերպով կբացառվեր, և դրա հետ գաղտնիություն երաշխիքն էլ կկորչեր: Ուստի D-ն ընտրվում է այնպիսի ընտանիքից, որի unsatisfiability-ն ֆիքսված կանոնագիրքը չի կարող արդյունավետորեն հերքել. կանոնագիրքի ներսում չկա կարճ ապացույց, որ D-ն լուծում չունի:

Ալիսի մեկ հաղորդագրությամբ ապացույցը հետո վերաբերում է «կամ/կամ» պնդման.

կա՛մ իրական MegaSudoku S-ը լուծում ունի, կա՛մ decoy D-ն լուծում ունի:

Սա է տրամաբանական կապը: D-ն **չի** ստեղծվում կախարդական եղանակով S-ը ճիշտ դարձնելու համար: ապացույցը չի ասում՝ «D-ն լուծում չունի, հետևաբար S-ը լուծում ունի»: Այն ապացուցում է disjunction՝ **S կամ D**: Կատարյալ հուսալիությունը նշանակում է, որ կեղծ disjunction-ը վավեր ապացույց չունի: Քանի որ D-ն իրականում կեղծ է՝ լուծում չունի, disjunction-ը կարող է ճիշտ լինել միայն այն դեպքում, երբ S-ը ճիշտ է: Ուստի եթե ապացույցը ընդունվում է, S-ը պետք է լուծում ունենա: Decoy-ը չի կարող կեղծ S-ը ճիշտ դարձնել:

Բայց գրոյական գիտելիքին նման մասի համար հարցերը՝ ինչ կլինեն, եթե D-ն *իսկապես* լուծում ունենար: Այդ decoy լուծումը կդառնար այլընտրանքային վկա և թույլ կտար ապացույցներ ստեղծել՝ առանց Ալիսի իրական MegaSudoku լուծումը իմանալու: Այսինքն՝ այն սիմուլյատորի դեր կխաղար: Իրականում D-ն լուծում չունի, ուստի սիմուլյատորի այս ճանապարհը փակ է: Բայց առանցքային կետն այն է, որ կանոնագիրքը չի կարող արդյունավետորեն ապացուցել, որ այն փակ է:

Այսպիսով D-ն երկու աշխատանք ունի: **հուսալիությունի** համար D-ն կեղծ է, և «S կամ D» վավեր ապացույցը ստիպում է S-ին ճիշտ լինել: **արդյունավետ գրոյական գիտելիքի** համար D-ն դժվար է հերքել, ուստի կանոնագիրքը չի կարող արագ բացառել decoy ճանապարհը, որը սիմուլյատորը հնարավոր կդարձներ:

Ուստի անվտանգության հարցը այլևս սա չէ.

Կարո՞ղ ենք ապացուցել, որ սիմուլյատոր իսկապես գոյություն ունի:

Այն դառնում է.

Կարո՞ղ է ձեր կանոնագիրքը արդյունավետորեն ապացուցել, որ սիմուլյատորը անհնար է:

Եթե պատասխանը ոչ է, հետևում է զարմանալիորեն ուժեղ բան: Չրոյական գիտելիքից բխող յուրաքանչյուր անվտանգության երաշխիք, որը (ա) կարելի է դիտարկել որևէ թեստ գործարկելով և (բ) կանոնագիրքի ներսում ապացուցելիորեն հետևում է սիմուլյատորի գոյությունից, իրականում պահպանվում է: Դրանցից որևէ մեկի հաջող հարձակումը ինքնին կտար այն կարճ refutation-ը, որը բացակայում է: Հենց սա է *արդյունավետ* մասը արդյունավետորեն զրոյական գիտելիքիում:

Դասարանային տարբերակումը հետևյալն է.

Դասական զրոյական գիտելիքի. ապացույցները անվտանգ են, որովհետև սիմուլյատոր գոյություն ունի:

Գյուղեյան արդյունավետ զրոյական գիտելիքի. observable անվտանգության թեստերի համար ապացույցները դիտվում են որպես անվտանգ, որովհետև կանոնագիրքը չի կարող արդյունավետորեն ապացուցել, որ սիմուլյատորն անհնար է:

Երկրորդ պնդումն ավելի թույլ է: Բայց հենց դրա շնորհիվ հոդվածը կարողանում է պահել դասական տարբերակը կոտրած երեք հատկությունները՝ մեկ հաղորդագրություն, կարգավորման բացակայություն և կատարյալ հուսալիություն:

Նոր թեստը. դուք չեք կարող ապացուցել, որ սիմուլյատորը բացակայում է

Ilango-ի relaxation-ը փոխում է հարցը:

Դասական զրոյական գիտելիքի հարցնում է.

սիմուլյատոր գոյություն ունի՞:

արդյունավետորեն զրոյական գիտելիքի հարցնում է ավելի թույլ բան.

Կարո՞ղ է ձեր ընտրված կանոնագիրքը արդյունավետորեն ապացուցել, որ սիմուլյատոր գոյություն չունի:

Սա կարող է տեխնիկական շրջանցում թվալ, բայց հենց սա է հիմնական գաղափարը: Կառուցումը գտնվում է տարօրինակ վիճակում. սիմուլյատոր իրականում **չկա**, և հոդվածը սա բացահայտ ասում է, բայց ձեր ֆիքսած կանոնագիրքը չի կարող արդյունավետորեն ապացուցել դրա բացակայությունը: Եթե ձեզ հետաքրքրող բոլոր վատ հետևանքների համար նման refutation անհրաժեշտ լիներ, համակարգը այդ հետևանքների նկատմամբ շարունակում է զրոյական գիտելիքի պես վարվել:

Այստեղ է հայտնվում Գյոդելը: Ոչ որպես զարդարանք և ոչ «Գյոդելը crypto-ն անվտանգ է դարձնում» կարգախոսով: Կապը ապացույց տեսությունի հետ է: Կանոնագիրքը կոչվում է **օպտիմալ**, եթե ճշգրիտ իմաստով լավագույնն է. երբ որևէ կանոնագիրք համապատասխան տեսակի բանաձևը կարճ ապացույցով կարող է հերքել, օպտիմալ կանոնագիրքն էլ կարող է՝ առավելագույնը *polynomially* ավելի երկար ապացույցով: Krajíček-ը և Pudlák-ը 1989-ին վարկած են ձևակերպել, որ **օպտիմալ ապացույց համակարգ գոյություն չունի**: Այսինքն՝ որ կանոնագիրքն էլ ֆիքսեք, մեկ ուրիշ կանոնագիրք որոշ ճշմարիտ պնդումների ընտանիք շատ ավելի կարճ ապացույցներով է ապացուցում: Սա ապացույց բարդությունի կենտրոնական բաց վարկածներից է և Գյոդելի անավարտություն թեորեմի վերջավոր, բարդություն-theoretic ազգականը. որոշ ճշմարիտ պնդումներ ձեր ֆիքսած կանոնագիրքում կարճ ապացույց չունեն՝ ոչ այն պատճառով, որ սկզբունքորեն անապացուցելի են, այլ որովհետև յուրաքանչյուր ֆիքսված կանոնագիրք որոշ կարճ ճշմարտությունների համար կարճ ապացույց չի տալիս:

հոդվածը ընդունում է այս վարկածը՝ մի փոքր ավելի ուժեղ **infinitely often** ձևով, որը գաղտնագրությունում ենթադրություններ օգտագործելիս ստանդարտ է: Krajíček-Pudlák թեորեմը դրա դիմաց տալիս է կոնկրետ արդյունք. յուրաքանչյուր կանոնագիրքի համար գոյություն ունի իսկապես անբավարարելի բանաձևերի հաջորդականություն, որը կանոնագիրքը կարճ ապացույցներով չի կարող հերքել, և, վճռականորեն, այդ բանաձևերը կարելի է **արդյունավետ ալգորիթմով ստեղծել**: Վերջին հատկությունը՝ միատեսակությունն, գաղափարը պարզապես գոյության պնդումից վերածում է Ալիսի իրականում գործարկելի ալգորիթմի. decoy D-ները գալիս են «արտադրական գծից», ոչ դատարկությունից:

Կրիպտոգրաֆիկ քայլն այդ ապացույց վիճակագրական հզորությունի պակասն օգտագործելն է որպես ռեսուրս:

Ինչ է իրականում անում կառուցումը

Ահա հոդվածի կառուցումը՝ մերկ կառուցվածքով:

Ֆիքսենք կանոնագիրք, ասենք ZFC: ապացույցբարդություն ենթադրությունի ներքո գոյություն ունի արդյունավետորեն գեներացվող բանաձևերի հաջորդականություն, որոնք իրականում անբավարարելի են, բայց կանոնագիրքը չունի կարճ ապացույց, որ դրանք անբավարարելի են:

Հիմա կառուցենք մեկ հաղորդագրությամբ այս ձևի ապացույց.

կա՛մ իրական պնդումը բավարարելի է, կա՛մ այս հատուկ դժվար բանաձևը բավարարելի է:

Հատուկ դժվար բանաձևը իրականում բավարարելի չէ: Հետևաբար, եթե հիմքում ընկած ապացույց machinery-ն կատարյալ sound է, հաղորդագրության ընդունումը

դեռ նշանակում է, որ իրական պնդումը ճիշտ է: Սա տալիս է կատարյալ հուսալիություն:

Բայց գրոյական գիտելիքին նման անվտանգության համար պատկերացրեք, որ հատուկ դժվար բանաձևը *բավարարելի լինել*: Այդ դեպքում դրա վկար կարելի էր օգտագործել ապացույցներ simulate անելու համար՝ առանց իրական վկար իմանալու: Իրականում այն բավարարելի չէ, բայց կանոնագիրքը չի կարող դա արդյունավետորեն ապացուցել: Ուստի չի կարող արդյունավետորեն ապացուցել, որ սիմուլյատորն անհնար է:

Սա է առանցքը: Համակարգը գաղտնիքը չի թաքցնում դասական սիմուլյատոր ստեղծելով: Observable անվտանգության թեստերի մեծ դասի համար այն գաղտնիքը պահում է կանոնագիրքի այն անկարողության հետևում, որ նա չի կարող հաստատել սիմուլյատորի բացակայությունը:

Ինչ է պնդում հոդվածը

Հիմնական թեորեմն ունի մի քանի շերտ: Կորիզը սա է:

Ստանդարտ գաղտնագրային ենթադրությունի՝ **ոչ ինտերակտիվ վկա indistinguishable ապացույցներ**-ի գոյության ներքո, որոնք լավ ուսումնասիրված օբյեկտներ են և հետևում են մի քանի հաստատված ենթադրություն package-ներից, և ապացույցբարդություն վարկածի ներքո, որ **(infinitely often) օպտիմալ ապացույց համակարգ գոյություն չունի**, հոդվածը յուրաքանչյուր ընտրված կանոնագիրքի համար կառուցում է NP/SAT-ի մեկ հաղորդագրությամբ ապացուցող և ստուգող՝ **կատարյալ հուսալիություն**-ով, առանց կարգավորման, որոնք տվյալ կանոնագիրքի նկատմամբ արդյունավետորեն գրոյական գիտելիքի են: NP/SAT-ը գլոխկոտրուկի նման խնդիրների ստանդարտ «ամենադժվար ընդհանուր հայտարարն» է, իսկ mega-Մուդոկուն դրա մի ներկայացումն է:

հերքելի անվտանգության հատկությունները ավելի լայնորեն պահպանելու թեորեմի համար հոդվածը ավելացնում է ևս մեկ ստանդարտ ենթադրություն՝ derandomization-ի **P = BPP** համոզմունքը, մոտավորապես՝ պատահականությունը ալգորիթմներին էական լրացուցիչ ուժ չի տալիս:

Թեորեմի լեզվից դուրս սա նշանակում է.

- ապացույցը մեկ հաղորդագրություն է:
- Trusted կարգավորում չկա:
- Կեղծ պնդումները հնարավոր չէ ապացուցել:
- ապացուցողը դասական գրոյական գիտելիքի չէ. այն սիմուլյատոր չունի:
- Բայց դասական գրոյական գիտելիքի յուրաքանչյուր հերքելի, խաղային փորձով ստուգվող անվտանգության հետևանք կարելի է ստանալ այս միջավայրում:

հերքելի բառը կարևոր է: Այն նշանակում է, որ անվտանգության ձախողումը կարելի է փորձարկել՝ հակառակորդն խաղի մեջ գործարկելով: Շատ գաղտնագրային

անվտանգությունն սահմանումներ հենց այս ձևն ունեն. կարող է հակառակորդն տարբերել երկու ciphertext, շրջել ֆունկցիան, գտնել վկա կամ հաղթել որևէ հստակ սահմանված փորձարկում: Թեորեմը յուրաքանչյուր նման հերքելի հատկության համար տալիս է ապացուցող՝ մեկ առ մեկ: Մի ապացուցող, որը միաժամանակ ունենա *քոլոր* հերքելի հատկությունները, հավանաբար անհնար է: Հին վերօգտագործելիություն հարձակումը՝ «Բոքը կարող է ապացույցը ցույց տալ ուրիշներին», ինքն էլ հերքելի հատկություն է և այստեղ իրականում ձախողվում է: հողվածի առաջարկն այն է, որ մեկ ապացուցողը plausibly կարող է ծածկել բոլոր *բնական* հերքելի հատկությունները՝ նրանք, որոնք իրական գաղտնագրային practice-ում հանդիպում են, բայց այս մասը պայմանական թեորեմ է՝ հիմնված «բնական»-ի ոչ ֆորմալ գաղափարի և առանձին վարկածի վրա: Երաշխիքը ուղղված է observable failures-ին, ոչ գաղտնիության ամեն փիլիսոփայական կամ սիմուլյացիա-based իմաստին:

Մեկ կոնկրետ corollary արժե անվանել. կառուցումը տալիս է առաջին ոչ ինտերակտիվ **վկա-hiding** ապացույցները միատեսակ ապացուցողով՝ «գլոխկոտրուկի ապացույցը չի օգնում գտնել դրա լուծումը», առանց փոխազդեցությունի և առանց կարգավորման: Համեստ հնչող օբյեկտ է, որը տասնամյակներ շարունակ չէր հաջողվում կառուցել:

Ինչ սա չի ասում

Այս բաժինն է հողվածը պահում ազնիվ:

Սա **չի** ասում, որ հին impossibility թեորեմները սխալ էին: Կառուցումը դրանցից խուսափում է սահմանումը փոխելով:

Սա **չի** տալիս սովորական, դասական գրոյական գիտելիքի՝ առանց փոխազդեցությունի, առանց կարգավորման և կատարյալ հուսալիությունով: հողվածը բացահայտ ասում է, որ կառուցված ապացուցողը սիմուլյատոր չունի:

Սա **չի** նշանակում, որ ապացույցը հնարավոր չէ կրկին օգտագործել: Մեկ հաղորդագրությամբ ապացույցը դեռ կարելի է ցույց տալ ուրիշներին: հողվածը չի պահպանում հերքելիությունի նման հատկություններ: Վստահելի կարգավորումով ոչ ինտերակտիվ գրոյական գիտելիքին էլ նույն սահմանափակումն ունի:

Սա **չի** նշանակում, որ ունենք կիրառումին պատրաստ գործնական արձանագրություն: Սա բարդություն տեսություն և գաղտնագրային foundations է: Արդյունքը կախված է ապացույց բարդությունի և գաղտնագրությունի խոշոր ենթադրություններից, իսկ կառուցումը վերաբերում է սկզբունքային հնարավորությանը:

Եվ սա **չի** դարձնում «Գյոդել»-ը կախարդական անվտանգության պրիմիտիվ: Գյոդելի կապը ապացույց համակարգերի, օպտիմալ ապացույց համակարգերի և անավարտությունի վերջավոր անալոգների միջոցով է: Օգտակար ինտուիցիան «անավարտությունը պաշտպանում է ձեր password-ը» չէ: Այն սա է. եթե կանոնագիրքը չի կարող արդյունավետորեն ապացուցել, որ սիմուլյատորն անհնար է, ապա այդ ապացույցը պահանջող հարձակումները կարելի է արգելափակել անվտանգություն

սահմանումների մակարդակում:

Ինչու է սա այնուամենայնիվ հետաքրքիր

գաղտնագրությունն հաճախ դժվարությունը վերածում է անվտանգության: Factoring-ը դժվար է, ուստի RSA-ի նման ենթադրությունները օգտակար են: Lattice խնդիրները դժվար են, ուստի lattice գաղտնագրությունն օգտակար է: Այստեղ դժվարությունը ավելի տարօրինակ է. ոչ «գաղտնիքը հաշվելն է դժվար», այլ «դժվար է ապացուցել, որ որոշակի ապացույց օբյեկտ չի կարող գոյություն ունենալ»:

Հենց դա է հոդվածը անսովոր դարձնում: Այն արսիումներն ու կանոնագիրքները գրեթե գաղտնագրային ռեսուրսի պես է վերաբերվում: Սովորական impossibility-ն ասում է՝ հուսալիությունն և սիմուլյացիան միջև լարվածություն կա: Plango-ի քայլը այդ լարվածությունը տեղափոխում է ապացույց-theoretic վարագույրի հետևել. սիմուլյատոր չկա, բայց ֆորմալ համակարգը չի կարող արդյունավետորեն բացահայտել այդ բացակայությունը:

Ընթերցողի համար զարմանալին այն չէ, որ սա փոխարինելու է այսօրվա զրոյական գիտելիքի համակարգերը: Հավանաբար անմիջապես չի փոխարինի: Չարմանալին այն է, որ մաթեմատիկական logic-ի սահմանափակումը կարելի է կառուցողականորեն օգտագործել՝ ոչ միայն որպես պատ, այլ նաև որպես ծածկույթ:

Որքա՞ն ուժեղ են ապացույցները

Սա թեորեմ հոդված է, ուստի «ապացույց» բառը այստեղ այլ բան է նշանակում, քան կենսաբանության կամ աստղագիտության հոդվածում: Հարցը այն չէ, թե փորձը վերարտադրվե՞լ է: Հարցն այն է՝ սահմանումները, ենթադրությունները և ապացույց շղթաը արդյոք աջակցո՞ւմ են պնդմանը:

ապացույցը ֆորմալ է, և հոդվածը բացահայտ է իր ենթադրությունների մասին: ենթադրությունները պատահական չեն: ոչ ինտերակտիվ վկա indistinguishable ապացույցները գաղտնագրությունում ստանդարտ օբյեկտներ են և հետևում են մի քանի հայտնի ենթադրություն package-ներից: No-օպտիմալապացույցհամակարգ վարկածը ապացույց բարդությունի կենտրոնական վարկած է: $P = BPP$ -ն ստանդարտ derandomization belief է, որն օգտագործվում է միայն հերքելի հատկություն ավելի լայն թեորեմի համար:

հոդվածը նաև փաստարկում է, որ ենթադրությունները պատահական հենարան չեն, այլ մոտավորապես ճիշտ «գինն» են: Այն ապացուցում է converse, ըստ որի դրանք էապես անհրաժեշտ են. եթե նման constructions ընդհանրապես գոյություն ունեն, ապա ոչ ինտերակտիվ վկա indistinguishable ապացույցները պետք է գոյություն ունենան, և ստանդարտ միակողմանի ֆունկցիաները ընդունելու դեպքում օպտիմալ

ապացույց համակարգ չպետք է գոյություն ունենա: Ենթադրությունները նաև «win-win» բնույթ ունեն. դրանցից որևէ մեկի հերքումն ինքնին նշանակալի հայտնագործություն կլիներ ապացույց բարդությունում, գաղտնագրությունում կամ բարդություն տեսությունում:

Բայց քանի որ արդյունքը պայմանական է, վստահությունն էլ պայմանական է: Եթե ենթադրությունները սխալ լինեն, թեորեմի մեկնաբանությունը փոխվում է: Եվ նույնիսկ եթե ճիշտ են, երաշխիքը լիարժեք դասական գրոյական գիտելիքի չէ. այն հոդվածի relaxed, ապացույց-theoretic տարբերակն է:

Ուստի ճիշտ վստահության աստիճանը սա է. բարձր՝ որ հոդվածը հաստատում է հետևողական պայմանական possibility արդյունք, միջին՝ որ ենթադրությունները նկարագրում են այն գաղտնագրային աշխարհը, որտեղ իրականում ապրում ենք, և ցածր՝ անմիջական գործնական consequence-ի համար:

Ինչու է սա կարևոր

հոդվածը ճանապարհ է բացում այնտեղ, որտեղ ըստ դասական տեսության դուռը փակ էր:

Դասական տեսությունն ասում է՝ լիարժեք գրոյական գիտելիքի չի կարող լինել մեկ հաղորդագրություն առանց կարգավորման և չի կարող ունենալ կատարյալ հուսալիություն: Flango-ի հոդվածը ասում է. եթե մենք պահանջենք գրոյական գիտելիքի այն հետևանքները, որոնք կարելի է փորձարկել անվտանգություն games-ով, և եթե անվտանգություն սահմանումը թույլ տանք կախված լինել նրանից, թե կանոնագիրքը ինչ կարող է կամ չի կարող արդյունավետորեն հերքել, ապա օգտակար վարքագծի մեծ մասը կարելի է վերականգնել՝ մեկ հաղորդագրությամբ, առանց կարգավորման և կատարյալ հուսալիությունով:

Սա փոքր բառային փոփոխություն չէ: Սա գաղտնագրային guarantees-ի մասին մտածելու այլ ձև է: Միայն «ինչ գոյություն ունի» հարցնելու փոխարեն հարցրեք՝ «ձեր կանոնագիրքը ի՞նչ կարող է բացառել»: Unprovability-ին փիլիսոփայական անհարմարություն համարելու փոխարեն օգտագործեք այն որպես կառուցվածք:

Գործնական աշխարհը գուցե վաղը չփոխվի: Բայց գաղափարական քարտեզը փոխվում է: Այժմ ֆորմալ իմաստով հնարավոր է, որ «ոչ ոք չի կարող արդյունավետորեն ապացուցել, որ գաղտնիքը արտահոսել է» պնդումը բավական ուժեղ լինի վերականգնելու համար այն խաղային փորձով ստուգվող պաշտպանություններից շատերը, որոնք ուզում էինք «գաղտնիքը չի արտահոսել» ավելի ուժեղ պնդումից:

Հենց դրա համար Գյոդելը տեղ ունի վերնագրում:

Կարճ ամփոփում

Չրոյական գիտելիքի ապացույցները թույլ են տալիս ապացուցողին համոզել ստուգողին, որ պնդումը ճիշտ է՝ առանց վկար բացահայտելու: Դասական impossibility արդյունքները ասում են, որ զրոյական գիտելիքի չի կարող սեղմվել մեկ հաղորդագրության մեջ առանց կարգավորման և չի կարող ունենալ կատարյալ հուսալիություն: Rahul Ilango-ի հոդվածը չի հերքում այդ անհնարիությունները: Այն սահմանում է ավելի թույլ գաղափար՝ արդյունավետորեն զրոյական գիտելիքի: սիմուլյատորի իրական գոյությունը պահանջելու փոխարեն այն պահանջում է, որ ընտրված ապացույց համակարգը՝ ZFC-ի նման ֆորմալ կանոնագիրքը, չկարողանա արդյունավետորեն ապացուցել, որ սիմուլյատոր գոյություն չունի: գաղտնագրությունի կարևոր ենթադրությունների՝ ոչ ինտերակտիվ վկա indistinguishable ապացույցների, և ապացույց բարդությունի «օպտիմալ ապացույց համակարգ գոյություն չունի» վարկածի ներքո հոդվածը կառուցում է NP/SAT-ի համար մեկ հաղորդագրությամբ provers՝ առանց կարգավորման և կատարյալ հուսալիությունով, որոնք հատկություն-by-հատկություն ստանում են զրոյական գիտելիքի հերքելի, խաղային փորձով ստուգվող հետևանքները: Մեկ ապացուցող, որը ծածկի բոլոր «բնական» նման հատկությունները, հետագա, մասամբ conjectural extension է, իսկ բառացիորեն բոլոր հերքելի հատկությունները միաժամանակ պահելն, ամենայն հավանականությամբ, անհնար է, որովհետև ապացույցները շարունակում են վերագրագործելի լինել: Արդյունքը տեսական և պայմանական է, ոչ կիրառումին պատրաստ պրիմիտիվ, բայց այն ցույց է տալիս ապացույց-theoretic unprovability-ն որպես գաղտնագրային ռեսուրս օգտագործելու նոր ձև:

Առանց չափազանցության

Ինչ է ցույց տալիս հոդվածը. Նշված ենթադրությունների ներքո հնարավոր է NP/SAT-ի համար կառուցել մեկ հաղորդագրությամբ, կարգավորում չունեցող, կատարյալ հուսալի provers, որոնք ցանկացած ընտրված ապացույց համակարգի նկատմամբ արդյունավետորեն զրոյական գիտելիքի են և ստանում են դասական զրոյական գիտելիքի յուրաքանչյուր հերքելի խաղային փորձով ստուգվող հետևանքը:

Ինչն է հավանական, բայց անվերապահ ապացուցված չէ. Որ անհրաժեշտ ապացույցբարդություն և գաղտնագրային ենթադրությունները ճիշտ են: Դրանք լուրջ, լավ ուսումնասիրված ենթադրություններ են, և հոդվածը ցույց է տալիս, որ դրանք ըստ էության ոչ միայն բավարար, այլ նաև անհրաժեշտ են, բայց nonetheless ենթադրություններ են:

Ինչ չի ցույց տալիս. Դասական զրոյական գիտելիքի՝ առանց փոխազդեցությունի, կարգավորման և կատարյալ հուսալիությունով, կիրառումին պատրաստ գործնական համակարգ, ապացույցների հերքելիություն կամ non-վերագրագործելիություն, կամ այն, որ Գյոդելի անավարտություն թեորեմը ինքնուրույն crypto է ապահովում:

Հիմնական սահմանափակումները. Երաշխիքը գրոյական գիտելիքի relaxation է, ամենալայն տարբերակը կախված է մի քանի ենթադրություններից, մեկ-universal-ապացուցող պնդումները մասամբ conjectural են, իսկ արդյունքը հիմնականում foundational է:

Որքա՞ն վստահություն պետք է ունենա ընդհանուր ընթերցողը. Բարձր՝ որ սա կարևոր պայմանական տեսություն արդյունք է, եթե սահմանումները ընդունվում են: Միջին՝ որ ենթադրությունները ճիշտ են իրական գաղտնագրային աշխարհում: Ցածր՝ անմիջական կիրառումի համար: Անվտանգ եզրակացությունը սա է. հոդվածը չի կոտրում գրոյական գիտելիքի impossibility-ները, այլ գտնում է նոր ապացույց-theoretic ճանապարհ՝ շրջանցելու դրանց այն մասերը, որոնք շատ անվտանգություն games-ի համար կարևոր են:

Աղբյուրներ

Ilango, IACR ePrint 2025/1296

<https://eprint.iacr.org/2025/1296>

FOCS 2025, DOI 10.1109/FOCS63196.2025.00058

<https://doi.org/10.1109/FOCS63196.2025.00058>