

Az általánosított Thurstone módszer és alkalmazhatóságának vizsgálata, valamint az általa kínált előrejelzések pontosságának növelése

doktori (PhD) értekezés

GYARMATI LÁSZLÓ

DOI:10.18136/PE.2026.1004

Konzulensek:

Dr. Mihálykó Csaba,
egyetemi docens

Dr. Mihálykóné dr. Orbán Éva
egyetemi docens

Informatikai Tudományok Doktori Iskola
Pannon Egyetem Veszprém

2026

**Az általánosított Thurstone módszer
és alkalmazhatóságának vizsgálata,
valamint az általa kínált előrejelzések
pontosságának növelése**

Az értekezés doktori (PhD) fokozat elnyerése érdekében készült a Pannon
Egyetem Informatikai Tudományok Doktori Iskolája keretében

informatikai tudományok tudományágban

Írta: Gyarmati László

Témavezetők: Dr. Mihálykó Csaba és Dr. Mihálykóné dr. Orbán Éva

Elfogadásra javaslom: igen / nem.

.....
témavezető

Elfogadásra javaslom: igen / nem.

.....
témavezető

Az értekezés bírálatra bocsátható.

.....
TDHT elnök

A jelölt az értekezés nyilvános vitáján %-ot ért el.

A bíráló Bizottság tagjai:

elnök:

bírálok:

tagok:

Veszprém,

.....
Bíráló Bizottság elnök

A doktori (PhD) oklevél minősítése:

Veszprém,

.....
EDHT elnök

Kivonat

Napjainkban egyre nagyobb jelentőségűvé válik a döntések támogatása, amelynek egyik lehetséges módja a páros összehasonlítási modellek alkalmazása. Ezek a modellek a kiértékelendő objektumok páronkénti összehasonlításából származó verbális, vagy más jellegű döntések eredményeit konvertálják erősségeket jellemző mérőszámokká, sorrendekké. Ezek a modellek akkor is megbízhatóan teljesítenek, amikor nem áll rendelkezésre nagy mennyiségű adat, ami a mesterséges intelligencia alkalmazásának elengedhetetlen feltétele. Dolgozatomban a sztochasztikus háttérű páros összehasonlítási modellekkel foglalkoztam. Ezek a modellek relációs döntéseket kiértékelve számos lehetőséget rejtenek magukban; többek között alkalmasak előrejelzésre. A dolgozat ezen modellek családjával kapcsolatos eredményeimet tartalmazza általános eloszlásfüggvény használata mellett, az adatok kiértékelhetősége, konzisztenciája és a modellek alkalmazhatósága szempontjából mind elméleti, mind szimulációs eszközökkel.

A paraméterek becslésére maximum likelihood becslési módszert alkalmaztam. Bár 3-opciós modelleket gyakran használnak, az adatok kiértékelhetőségének feltétele nem volt tisztázva. Kutatásaim során szimulációs kísérletek alapján megsejtettem az adatok kiértékelhetőségének szükséges és elégséges feltételét, amelyet sikerült matematikai eszközökkel is bizonyítani mind a 3-opciós Thurstone-motivált modellek, mind a Davidson modell esetén. Ezzel egy régóta nyitott kérdésre adtam meg a választ.

Vizsgáltam azt is, hogy milyen opciószámú modelleket érdemes alkalmazni. Megmutattam, hogy elegendő mennyiségű adat rendelkezésre állása esetén érdemesebb a nagyobb opciószámú, ezáltal több információt feldolgozó, összetettebb modelleket használni. Szimulációs vizsgálatok alapján megadtam a szükséges adatszámok küszöbszámait is.

A páros összehasonlítási mátrixokon alapuló modellek esetén központi szerepet játszó konzisztencia fogalom segítségével megadtuk az adatok konzisztenciájának hiányzó definícióját a 2-opciós Bradley-Terry és a Davidson modellben és egy, a konzisztenciával ekvivalens állítást bizonyítottunk az előbbi modellben. Megmutattuk, hogy konzisztens adatok esetén a páros összehasonlítási mátrixokon alapuló modell és a 2-opciós Bradley-Terry modell egyforma kiértékelési eredményeket ad, és inkonzisztens adatok esetén is ugyanazok az optimális összehasonlítási struktúrák. Ezzel kapcsolatot teremtettünk a két, egymástól különböző modellcsalád között.

A modelleket számos esetben alkalmaztam különféle problémák esetén (sporteredmények kiértékelése különböző szempontok szerint, többszintű döntéstámogatás, fényforrások csoportosítása), és az alkalmazások során rámutattam a modellek előnyös tulajdonságaira. Megmutattam, hogy a szokásos három opciónál több opciót alkalmazva, a hazai pálya előnyét és az időbeli súlyozást beépítve a modellek előrejelző képessége javult, versenyképes az irodalomban fellelhető modellekével.

Abstract

Nowadays, decision support is becoming increasingly important, and one possible approach is the application of pairwise comparison models. These models convert the results of verbal or other types of judgments derived from pairwise comparisons of the evaluated objects into numerical measures characterizing strengths and into rankings. They perform reliably even when large amounts of data are not available, which is an essential requirement for the application of artificial intelligence.

In my thesis, I studied pairwise comparison models with a stochastic background. By evaluating relational decisions, these models offer numerous possibilities; among others, they are suitable for prediction. This thesis presents my results related to this family of models using general distribution functions, focusing on the evaluability and consistency of the data, as well as the applicability of the models, using both theoretical and simulation-based methods.

For parameter estimation, I applied the maximum likelihood estimation method. Although three-option models are frequently used, the conditions for data evaluability had not been clarified. During my research, based on simulation experiments, I formulated the necessary and sufficient condition for data evaluability, which was successfully proven using mathematical methods for both three-option Thurstone-motivated models and the Davidson model. In doing so, I provided an answer to a long-standing open question.

I also examined which number of options in models is worth applying. I demonstrated that when a sufficient amount of data is available, it is preferable to use models with a higher number of options, which process more information and are therefore more complex. Based on simulation studies, I also determined the threshold values for the required data size.

In the case of models based on pairwise comparison matrices, the concept of consistency plays a central role. Using this concept, we provided the previously missing definition of data consistency for the two-option Bradley-Terry and Davidson models, and proved a statement equivalent to consistency in the former model. We showed that in the case of consistent data, models based on pairwise comparison matrices and the two-option Bradley-Terry model yield identical evaluation results, and even in the case of inconsistent data, the optimal comparison structures are the same. In this way, we established a connection between two different model families.

I applied these models in various cases to different problems (such as evaluating sports results from different perspectives, multi-level decision support, and grouping light sources), and during these applications, I highlighted the advantageous properties of the models. I demonstrated that by using more than the usual three options, and by incorporating home advantage and time-based weighting, the predictive performance of the models improves and becomes competitive with those found in the literature.

Abstrakt

Heutzutage gewinnt die Entscheidungsunterstützung zunehmend an Bedeutung, wobei Paarvergleichsmodelle eine mögliche Methode darstellen. Diese Modelle transformieren die Ergebnisse verbaler oder sonstiger Entscheidungen aus paarweisen Objektvergleichen in Maßzahlen zur Charakterisierung von Präferenzen sowie in Rangordnungen. Diese Modelle funktionieren auch dann zuverlässig, wenn keine großen Datenmengen zur Verfügung stehen, was eine wesentliche Voraussetzung für den Einsatz künstlicher Intelligenz darstellt.

In meiner Arbeit beschäftigte ich mich mit stochastisch fundierten Paarvergleichsmodellen. Durch die Auswertung relationaler Entscheidungen bieten diese Modelle zahlreiche Möglichkeiten; unter anderem eignen sie sich auch für Prognosen. Die vorliegende Arbeit präsentiert Ergebnisse zu dieser Modellfamilie unter Verwendung allgemeiner Verteilungsfunktionen, wobei die Auswertbarkeit und Konsistenz der Daten sowie die Anwendbarkeit der Modelle sowohl theoretisch als auch simulationsbasiert untersucht werden. Zur Parameterschätzung habe ich die Maximum-Likelihood-Methode angewandt. Obwohl Modelle mit drei Optionen häufig verwendet werden, war die Bedingung für die Auswertbarkeit der Daten bislang nicht geklärt. Im Rahmen meiner Forschung wurden anhand von Simulationsversuchen die notwendigen und hinreichenden Bedingungen für die Auswertbarkeit der Daten ermittelt, deren mathematische Gültigkeit sowohl für die Thurstone motivierte Drei-Optionen-Modelle, und das Davidson-Modell nachgewiesen werden konnte. Damit habe ich eine lange offene Frage beantwortet.

Es wurde außerdem untersucht, welche Anzahl an Optionen in den Modellen sinnvoll bzw. zweckmäßig ist. Dabei habe ich gezeigt, dass es bei ausreichender Datenmenge vorteilhafter ist, Modelle mit einer größeren Anzahl an Optionen zu verwenden, da diese mehr Informationen verarbeiten und somit komplexer sind. Auf Grundlage von Simulationen habe ich auch Schwellenwerte für die erforderliche Datenmenge bestimmt.

Im Falle von Modellen, die auf Paarvergleichsmatrizen basieren, spielt der Begriff der Konsistenz eine zentrale Rolle. Mithilfe dieses Begriffs haben wir die zuvor fehlende Definition der Datenkonsistenz für das Zwei-Optionen-Modell von Bradley-Terry sowie für das Davidson-Modell angegeben und im erstgenannten Modell eine zur Konsistenz äquivalente Aussage bewiesen. Wir haben gezeigt, dass bei konsistenten Daten Modelle auf Basis von Paarvergleichsmatrizen und das Zwei-Optionen-Bradley-Terry-Modell identische Auswertungsergebnisse liefern und dass selbst bei inkonsistenten Daten die optimalen Vergleichsstrukturen übereinstimmen. Damit wurde eine Verbindung zwischen zwei unterschiedlichen Modellfamilien hergestellt.

Ich habe die Modelle in zahlreichen Fällen auf verschiedene Probleme angewandt (wie etwa die Auswertung von Sportergebnissen aus unterschiedlichen Perspektiven, mehrstufige Entscheidungsunterstützung und

die Gruppierung von Lichtquellen) und dabei die vorteilhaften Eigenschaften der Modelle aufgezeigt. Es konnte gezeigt werden, dass durch die Verwendung von mehr als den üblichen drei Optionen sowie durch die Einbeziehung des Heimvorteils und einer zeitlichen Gewichtung die Prognosefähigkeit der Modelle verbessert wird und mit den in der Fachliteratur beschriebenen Modellen konkurrieren kann.

Tartalomjegyzék

1. Téma indoklása és kutatási előzmények	1
1.1. Motiváció	1
1.2. A kutatás háttere	3
2. A vizsgált modellek	7
2.1. Thurstone-motivált modellek	7
2.2. Speciális Thurstone-motivált modellek: a 2- és 3-opciós Bradley-Terry modell	15
2.3. A Davidson modell	17
2.4. Páros összehasonlítási mátrixokon alapuló modellek	19
3. Informatikai megvalósítás	21
4. Az adatok kiértékelhetőségével kapcsolatos kutatások	25
4.1. Korábbi feltételrendszerek	26
4.2. A 2-opciós modellre vonatkozó Ford tétel általánosítása	29
4.3. Egy 3-opciós elégséges feltételrendszer	30
4.4. Az adatok kiértékelhetőségének szükséges és elégséges feltételrendszere 3-opciós modellek esetén	32
4.5. Optimális opciószám választása	39
5. Az adatok konzisztenciája és inkonzisztenciája a modellekben	45
5.1. Konzisztencia a 2-opciós Bradley-Terry modellben	46
5.2. Optimális összehasonlítási struktúrák 2-opciós Bradley-Terry modellek esetén inkonzisztens esetben	48
6. A modellek alkalmazása	52
6.1. Csapatok rangsorolása sportmérkőzések eredményei alapján	53
6.2. Fényforrások csoportosítása erősségük alapján	62
6.3. Többkritériumos döntéshozatal: atléták besorolása szakágakba	64
6.4. Sportesemények előrejelzése	71
6.5. Sportesemények előrejelzésének pontossága	77
7. Összefoglalás	83
A. Appendix	96
A.1. Optimális opciószám választása	96
A.2. Optimális összehasonlítási struktúrák a 2-opciós Bradley-Terry modellek esetén inkonzisztens esetben	99

Ábrák jegyzéke

1.	A döntések és a nekik megfeleltetett intervallumok az s -opciós modellben	8
2.	A programcsomag egyszerűsített folyamatábrája	24
3.	Példa olyan összehasonlítási eredményekre, amikor létezik csupa győzelmekből álló kör	31
4.	Példa olyan összehasonlítási eredményekre, amikor nem létezik csupa győzelmekből álló kör, az adatok mégis kiértékelhetők	31
5.	Az egyes feltételrendszerek egymáshoz való viszonya	33
6.	Az egyes feltételrendszerek hatékonysága a szükséges és elégséges feltételrendszerhez viszonyítva a Bradley-Terry modellben ($n = 10$)	35
7.	Az egyes feltételrendszerek hatékonysága a szükséges és elégséges feltételrendszerhez viszonyítva a Bradley-Terry modellben ($n = 20$)	36
8.	Az egyes feltételrendszerek hatékonysága a szükséges és elégséges feltételrendszerhez viszonyítva a Thurstone modellben ($n = 10$)	37
9.	Az egyes feltételrendszerek hatékonysága a szükséges és elégséges feltételrendszerhez viszonyítva a Thurstone modellben ($n = 20$)	37
10.	Az egyes feltételrendszerek hatékonysága a szükséges és elégséges feltételrendszerhez viszonyítva a Bradley-Terry modellben ($n = 50$)	38
11.	Spearman rangkorreláció az első 22 forduló alapján kialakított rangsor és az összes mérkőzés eredménye alapján kialakított rangsor között. A zöld oszlopok egy-egy új feltételrendszer teljesülését jelzik (5-NS, 6-SC, 11-DAV, 22-MO)	38
12.	Az (40) példának megfeleltetett irányított gráf a 4-opciós modellben	40
13.	A (41) példának megfeleltethető gráf a 4-opciós modellben	41
14.	A (42) példának megfeleltethető gráf az 5-opciós modellben	41
15.	Az angol csapatok becsült és valós Transzfermarket értékei	61
16.	Példa inkonzisztenciára 3 objektum esetén	63
17.	A módszer számítási folyamata	65
18.	A felmérések súlyai a szakágakban	66
19.	Példa a teljes erősség kiszámítására	66
20.	Optimális opciószám kiválasztó algoritmus folyamatábrája	98

21.	Különböző gráfstruktúrák esetén az információ visszakeresés átlagos mérőszámai: az EU_M, EU_W által meghatározott átlagos távolságok, a PE_M, PE_W, Spearman ρ és Kendall τ által meghatározott átlagos korrelációk, $n = 4$ objektum összehasonlítása esetén $l = 0, 15$ -ös perturbáció alkalmazásával	101
22.	Különböző gráfstruktúrák esetén az információ visszakeresés átlagos szórásai: az EU_M, EU_W által meghatározott szórások, a PE_M, PE_W, Spearman ρ és Kendall τ által meghatározott szórások értékei, $n = 4$ objektum összehasonlítása esetén $l = 0, 15$ -ös perturbáció alkalmazásával	102
23.	Az átlagos EU_M és EU_W távolságok a gráfok éleinek számának függvényében ($n = 6$)	102
24.	Átlagos rangkorrelációs (Spearman ρ és Kendall τ) értékek a különböző élszámú struktúrák esetén ($n = 6$)	103
25.	Az adott élszámú optimális struktúrák euklideszi távolsága $n = 6$ objektum esetén	103
26.	Pearson-korrelációs együtthatók az optimális gráfok és a teljes gráf között a gráfok élszámának függvényében	104
27.	Példa arra, hogy a kevesebb összehasonlítás több információt eredményezhet: a 7 élű optimális struktúra, mint egy 8 élű struktúra ($n=6$ objektum és 0,15 perturbáció esetén).	104
28.	EU_M és EU_W távolságok a csillaggráfhoz és a teljes gráfhoz tartozó eredmények között a perturbáció értékének függvényében	105
29.	PE_M és PE_W, Spearman és Kendall rangkorrelációs értékek a csillaggráf és a teljes gráf esetén esetén a különböző perturbáció értékének függvényében (PE_M és PE_W szinte megegyezik)	105
30.	EU_M, EU_W, PE_M, PE_W, ρ és τ értékek a csillaggráf és a teljes gráf esetén különböző perturbáció értékek mellett	106

Táblázatok jegyzéke

1.	Az általam alkalmazott modellek	13
2.	Különféle általam konkrétan alkalmazott modellek opciószámai	13
3.	Modellek előnykezeléseinek fajtái	13
4.	A modellek időbeli súlyozásai	14
5.	A kiértékelhetőségi feltételrendszerekkel kapcsolatos eredmények áttekintése	28
6.	Paraméterválasztás a szimulációkhoz	34

7.	A Premier League csapatainak rangsorai és erősségei különböző módszerekkel, a nemzeti bajnokságban játszott mérkőzéseik alapján	55
8.	Az egyes csapatok Transzfermarket értékei és az általuk elért pontok a Premier League esetén	56
9.	A rangsorok és súlyok a nemzetközi mérkőzések 4-es súllyal történő figyelembevételével	58
10.	A rangsorok korrelációi a teljes rangsor esetén a különféle modellekkel és súlyokkal	59
11.	Az átlagos erősségek országok szerint	60
12.	Átlagos helyezések országokként	60
13.	Az átlagos pénzübeli értékek a Transzfermarket szerint	60
14.	Szakágakba besorolandók eredményei és ajánlások a megfelelő szakágakra	68
15.	A fiú és lány versenyzők eredményei a két rangsor szerint	68
16.	Szakágakban lévő atléták eredményei	69
17.	A régi és az új teszt eredményei a lányoknál	69
18.	A régi és új teszt eredményei a fiúknál	70
19.	Az A csoport eredményei	72
20.	A B csoport eredményei	72
21.	A legjobbak és a legrosszabbak egymás elleni mérkőzéseinek vélt (és valós) eredményeit felhasználó egységes rangsor	73
22.	A legjobb nyolc csapat esetén az előrejelzések (félkövér betűk: helyes előrejelzések)	74
23.	A mérkőzések eredményei	75
24.	A legjobb 8 csapat és erősségeik csak az egymás elleni mérkőzéseik alapján	75
25.	A Final4 várható résztvevői (félkövér betűk: helyes előrejelzések)	75
26.	Az első oszlopban szereplő csapatok győzelmének valószínűsége (félkövér betűk: az eredmény megegyezik a valószínűséggel)	76
27.	A LogLoss értékei az egyes modellek alapján	80
28.	Az RPS értékei az egyes modellek alapján	80
29.	A BS értékei az egyes modellek alapján	81
30.	Az ACC értékei az egyes modellek alapján	82
31.	A 2- és 4-opciós modell összehasonlítása szimuláció segítségével ($m_i : 0 - 2, d_1 = -1, 35$)	96
32.	A 3- és 5-opciós modell összehasonlítása szimuláció segítségével ($m_i : 0 - 2, d_1 = -1, 35, d_2 = -0, 1$)	97
33.	Összehasonlítási struktúrák eredményei 0,05-ös perturbációval $n = 5$ esetén	100

Köszönetnyilvánítás

Szeretnék köszönetet mondani legfőképpen konzulenseimnek, kik időt nem kímélve segítségükkel támogattak, hogy elkészülhessem a disszertáció. Hála a sok tanácsért, amelyekkel vezettek a kutatásban, és segítettek, hogy egyre jobb eredményeket érhessek el.

Szeretnék köszönetet mondani családomnak, barátaimnak is, kik mellett álltak és bátorítottak ezen út megtételében.

Jelölések

LBWA	Level Based Weight Assessment
PCM	Pairwise Comparison Matrix
AHP	Analytic Hierarchy Process
LLSM	logaritmikusan legkisebb négyzetek módszere
MI	mesterséges intelligencia
TH	Thurstone modell
BT	Bradley-Terry modell
ML	maximum likelihood
n	objektumok száma
ξ_i	i-edik objektum mögötti látens valószínűségi változó
m_i	i-edik objektum várható erőssége
C_k	k-adik döntési opció
I_k	k-adik intervallum
d_k	k-adik határoló pont
b_i	i-edik objektum előnykezelésének mértéke
A^0	előnyt nem alkalmazó adatmátrix
A^e	előnyt alkalmazó modellben az előnyös összehasonlításokat tartalmazó adatmátrix
A^h	előnyt alkalmazó modellben a hátrányos összehasonlításokat tartalmazó adatmátrix
q	időbeli súlyozás paramétere
N	normál modell
K	konstans előnyt alkalmazó modell
E	egyedi előnyt alkalmazó modell
V	időbeli súlyozást alkalmazó modell
π_i	i-edik objektum súlya
$\vec{G}^{(2)}$	2-opciós modell adathalmazának irányított gráfja
MO	Mihálykó-Orbán féle feltételrendszer
NS	szükséges és elégséges feltételrendszer
SC	szigorú körös feltételrendszer
DAV	Davidson-féle feltételrendszer
τ	Kendall rangkorreláció
ρ	Spearman rangkorreláció
PE_M	Pearson rangkorreláció erősségekre
PE_W	Pearson rangkorreláció súlyokra
EU_M	Euklideszi távolság erősségekre
EU_W	Euklideszi távolság súlyokra
RN	Rank Net
TM	Transzfermarket
AL	PCM modell LLSM kiértékelővel
P_I_1	egyparaméteres Poisson-modell iteratív változata
LogLoss	Logaritmikus veszteségmutató
RPS	Rangsorolt valószínűségi pontszám
BS	Brier-Score
ACC	Pontosság

1. Téma indoklása és kutatási előzmények

1.1. Motiváció

Rohanó és mindinkább gyorsuló világunkban egyre nagyobb szerepet kap a különféle döntéstámogató modellek/szoftverek használata, alkalmazása. Mesterséges intelligenciával (MI) egyre jobban körül vagyunk véve, mely az esetek nagy részében igen kiváló eredményekkel segít minket munkánk, feladataink folyamán.

Azonban az MI kimondottan akkor teljesít jól, ha rengeteg adat és információ áll rendelkezésünkre. De mi a helyzet akkor, ha csekély mennyiségű adat mellett kell döntéseket támogatni? Ilyenkor jól bevált matematikai modellekre érdemes hagyatkoznunk, amelyek kevés adat esetén is megbízhatóan képesek teljesíteni.

Az információ visszanyerés a döntéstámogatás egyik formája. Sokszor meglévő információt kell kinyernünk, vagy olyanná alakítanunk, amellyel képesek vagyunk utána dolgozni. Az információ visszanyerés esetén közvetlenül nem ismert információkat próbálunk becsülni, melyekről csak valamilyen közvetett információval rendelkezünk. Információ visszanyerés például különféle fagylaltok összehasonlítása: mennyire finom egy fagylalt, melyik a legfinomabb az emberek szerint? (A rejtett információ itt a fagylalt általános finomsága, míg a meglévő információ az, hogy kinek mennyire finom az adott fagylalt.) Tudjuk számszerűsíteni, hogy átlagosan a vanília vagy a csokoládé a jobb és mennyivel? Van-e erre valami jó módszer, hogy információt nyerjünk vissza?

A páros összehasonlítások kiválóan alkalmasak döntéstámogatásra, információ visszanyerésre. A modellek azért hatékonyak az információkinyerésben, mert egyszerű döntési helyzeteken keresztül részletes, számszerű és ellenőrizhető képet adnak a preferenciákról és a kedveltségi/fontossági viszonyokról. Ezen módszerek könnyen alkalmazhatóak, kevés adat esetén is jól teljesítenek. Különféle páros összehasonlítási modelleket ismerünk, például az Analytic Hierarchy Process (AHP) és annak különféle változatai, fuzzy logikát alkalmazó Level Based Weight Assessment (LBWA) és a Thurstone-motivált modellek. Általánosan elmondható, hogy az emberek nehezen adnak meg 1-től 10-es skálán értékeket, és azok megbízhatósága sem elegendő. Két objektumot viszont könnyen össze tudnak hasonlítani, hogy melyik jobb, egyforma, vagy rosszabb. Egyszerű relációs adatokból képesek vagyunk a modell segítségével számszerűsített erősségeket

adni, amelyek már így további feldolgozáson eshetnek át.

A fagylaltos példára visszatérve, két fagylaltot az egyes emberek könnyedén össze tudnak hasonlítani, hogy melyiket érzik finomabbnak, vagy esetleg egyformának minősítik őket. A kinyert adatokból, összehasonlítási eredményekből a Thurstone-motivált modellek segítségével számszerűen megkaphatjuk például, hogy mennyire finom a vanília és a csokoládé ízű fagylalt. Így végül megkapjuk, hogy átlagosan mennyire finom (jó) az egyik illetve másik fagylalt. Természetesen az értékelést több szempont alapján is elvégezhetjük, és a kapott eredményeket aggregálhatjuk.

A páros összehasonlítási modellek közül kiemelkednek a Thurstone-motivált modellek, ami a valószínűségszámítási háttérnek köszönhető. A modellek nem csak rangsorokban, erősségek meghatározásában képesek minket segíteni, de jövőbeli összehasonlítások kimenetelére is képesek valószínűségeket megadni.

A modellek kiválóan alkalmasak általánosításra, különféle döntéstámogatási, információ visszanyerési problémák esetén való alkalmazásra. Amellett, hogy több döntési opció is megadható az általánosított modellekben, még az összehasonlítások esetén felmerülő esetleges előnyt is képesek kezelni. További bővítésekkel, általánosításokkal a modellek kiválóan alkalmasak többkritériumos döntéstámogatásra. Azonban vannak nyitott kérdések is velük kapcsolatban. A 3-opciós modellek esetén nem volt ismert a kiértékelhetőség szükséges és elégséges feltétele. A konzisztencia fogalmával sem foglalkoztak a kutatók, illetve a különféle modellek sem voltak összehasonlítva részletesen információ kinyerés szempontjából.

Dolgozatomban a Thurstone-motivált modelleket kutatom; több kérdésre keresem a választ az általánosított Thurstone modellekkel kapcsolatban. Mi az a minimális adathalmaz amely esetén ki tudjuk értékelni a modelleket? Kis adathalmazok esetén hogyan viselkednek? Más modellekkel milyen kapcsolatban állnak? Javítja-e az információ visszanyerést, előrejelzést az előny kezelése? Hogyan lehet tovább általánosítani a modelleket? De talán a legfontosabb kérdés mind közül, mire is alkalmazhatóak a modellek és milyen hatékonysággal?

Munkámban szeretném bemutatni, hogy a vizsgált modellek kiváló alternatívákat jelentenek döntéstámogatásra kis adathalmazok esetén.

1.2. A kutatás háttere

A páros összehasonlítási modellek két fő családba sorolhatóak; ezek a PCM (Pairwise Comparison Matrix) alapuló modellek és a sztochasztikus háttérű modellek.

A páros összehasonlítási mátrixokon alapuló módszerek családján belül, az Analytic Hierarchy Process (AHP) az egyik legelterjedtebb.

Saaty 1977-ben [1]-ben lefektette a módszer alapjait, és 1990-ben [2] összefoglalta az AHP filozófiáját. A PCM-et sajátvektor módszerrel értékelte ki. Módszerével csak teljes összehasonlítási mátrixok esetén lehetett a kiértékelést elvégezni, ami jelentősen szűkítette az alkalmazhatóságot. Saaty 1977-es cikkében bevezette a konzisztencia mutatókat is, amelyek a döntések következetességét hivatottak mérni.

A teljes összehasonlításra vonatkozó kíváncsi elhagyható, ha a PCM-ek kiértékelése a logaritmikusan legkisebb négyzetek módszerével (LLSM) történik. Ezt elemezték Bozókiék 2010-ben írt cikkükben [3]. Megadtak továbbá egy olyan jobb/rosszabb arányt használó PCM kitöltést, amelynek segítségével már sokkal több értéket fel tudtak venni a mátrix elemei az alapmodellhez képest.

Brunelli 2018-ban konzisztencia- és inkonzisztencia-indexekről írt áttekintő cikket, amelyben bemutatott sok gyakran használt konzisztencia mérőszámot, és felsorolta azon fontos tulajdonságokat, amelyeket az ilyen módszereknek teljesíteni kell [4]. Ezen tulajdonságok más páros összehasonlítási modelleknél is kívánatosak.

A Thurstone modell nagy múltra tekint vissza. Az eredeti elképzelést Thurstone 1927-ben publikálta [5]. Thurstone pszichológusként ezt páciensei állapotának mérésére fejlesztette ki. Ez a páros összehasonlítási modell normális eloszlású valószínűségi változókat feltételez az egyes összehasonlítani kívánt objektumok mögé. Mindössze két kimenetelt engedett az összehasonlítások eredményeként, a jobbat és a rosszabbat. A legkisebb négyzetek módszerét alkalmazta az egyes erősségek becslésére a modellben. Thurstone munkája egyedülálló és az alapja az általunk alkalmazott modelleknek, ezért is nevezzük az általunk alkalmazott modelleket Thurstone-motivált modelleknek. Thurstone nem adott meg feltételrendszer 2-opciós modellje alkalmazhatóságára.

Mosteller (1951) cikkében [6] az általános modellt az egyenlő szórások és a nulla korreláció feltételezésével szűkítette, ami a modell gyakorlati

alkalmazását egyszerűbbé tette.

Bradley és Terry 1952-ben megjelent cikkükben [7] elhagyták a normális eloszlást és logisztikus eloszlást alkalmaztak helyette. Ez azért is volt fontos, mert az erősségeket ekkor már maximum likelihood becslés segítségével próbálták becsülni, amely ugyan jobb eredményeket ad, de számítási igénye jelentősen nagyobb, főleg a nehezebben számolható normális eloszlás esetén. Ekkor még nem volt megadva, hogy melyek azok a szükséges és elégséges feltételek, amik mellett kiértékelhető a modell.

Ford 1957-es cikkében [8] adta meg és bizonyította a 2-opciós modell esetén, hogy mik ezek a feltételek logisztikus eloszlás alkalmazása mellett. Innentől kezdve ismert volt a 2-opciós modell kiértékelhetőségének szükséges és elégséges feltétele, mely kulcsfontosságú a kiértékelhetőség eldöntésekor maximum likelihood becslés esetén.

Glenn és David 1960-ban cikkükben [9] a Thurstone–Mosteller modell 3-opciós általánosítását alkalmazták, amelynek segítségével már az "egyforma" döntés lehetőségét is beépítették a modellbe. Ennek köszönhetően már sokkal általánosabb az alkalmazás lehetősége: ha csak a sportmérkőzésekre gondolunk, mivel sok sportág megengedi a döntetlen lehetőségét, és a 3-opciós modell ezen adatok kiértékelésének is utat nyitott.

Rao és Kupper [10] cikkükben 1967-ben 3-opciósra bővítették a Bradley-Terry modellt, meghagyva a maximum likelihood paraméterbecslés eljárását. Mivel a Bradley-Terry modell számításigénye jelentősen kisebb, mint a Thurstone modellé így több alkalmazást tett lehetővé a modell.

Davidson 1970-ben [11] bemutatta a sztochasztikus páros összehasonlítási modellek családjába tartozó 3-opciós modelljét, mely szakít a korábbi Thurstone-motivált modellekkel, és amely a Luce féle választási axiómát teljesíti. Ezen modell általánosítása a 2-opciós Bradley-Terry modellnek. Davidson emellett egy elégséges feltételt adott 3-opciós modelljében az adatok kiértékelhetőségére. Ennek köszönhetően, most már lett egy feltételrendszer, mellyel el lehetett dönteni, hogy az adathalmaz kiértékelhető-e.

Agresti 1992-ben [14] cikkében bemutatta, hogy a jobb, rosszabb, egyforma kategóriákat még tovább lehet bontani, így 3-nál többopciós modelleket lehet alkalmazni. A paraméterbecslés módszere a többopciós modellek esetére újra a legkisebb négyzetek módszere lett, a relatív gyakoriságok és az eloszlásfüggvény bizonyos pontokban vett értékei alapján. A legegyszerűbb esetektől eltekintve a becslések kifinomult numerikus módszerek segítségével

hajthatóak végre. Ezek a numerikus módszerek kisebb erőforrás igényűek legkisebb négyzetek módszerén alapuló eljárások esetén, mint maximum likelihood becslés esetén. A likelihood becsléshez sem a feltételrendszer nem volt meg többopciós esetben, sem pedig a számítási teljesítmény nem állt még rendelkezésre ez időtájbán.

Az irodalomban foglalkoztak azzal az esettel is, amikor az adatok a maximum likelihood becsléssel nem kiértékelhetők. Conner és Grant 2000-ben [15] 2-opciós Bradley-Terry modellt alkalmaztak cikkükben, Zermelo 1929-es [16] cikkét (így indirekt a Bradley-Terry modellt) általánosították. Olyan esetekre adtak megoldást kiegészítések beszúrásával, amikor az összehasonlítási gráf nem erősen összefüggő, így korábbi módszerrel nem kaphattunk volna teljes rangsort. Yan 2016-os cikkében [17] a 3-opciós Bradley-Terry modellre adott olyan feltételeket, melynek segítségével a nem erősen összefüggő gráf kiértékelhető. Pontos egymáshoz viszonyított erősséget csak az erősen összefüggő komponenseken belül lehet meghatározni, de az egyes komponensek között lehet relációkat megadni, így részleges rangsor készíthető. Újabb eredményeket a nem kiértékelhető adathalmazokkal kapcsolatban a 2-opciós modellben Kovács és társai 2025-ben megjelent cikkükben [18] publikáltak.

Hunter 2004-ben numerikus algoritmusokat adott meg cikkében a Bradley-Terry modellhez [19]. Ezeket alkalmazva az adott modell szerinti kiértékelés számításigénye jelentősen csökkent.

Orbán-Mihálykó és társai 2019-ben megjelent cikkeikben [20, 21] az s -opciós modelleket már maximum likelihood becslés segítségével értékelték ki, bemutatva a maximum likelihood becslésekben rejlő lehetőségeket konfidenciaintervallumok szerkesztésére és hipotézisvizsgálatok végzésére. Általánosították az alkalmazható eloszlásfüggvények családját szigorúan log-konkáv sűrűségfüggvénnyel rendelkező eloszlásfüggvényekre [21]. Elégséges feltételt adtak meg az egyes modellek kiértékelhetőségére. Ennek következményeképpen a többopciós modellek esetén megszületett egy elégséges feltétel és ez idő tájban már a maximum likelihood becslés számítás igénye sem volt túlzott a kor technikájához mérve. A szerzők 2019-ben a 3-, 4-, és 5-opciós modellekre elégséges feltételrendszereket fogalmaztak meg egy 4-opciós modell alkalmazása mellett [22].

Ezután csatlakoztam a kutatásokhoz, hogy a modelleket továbbfejlesszük és alkalmazásokon keresztül is bemutassuk ezen modellek előnyeit.

A PCM alapú és a sztochasztikus háttérű modellelcsaládok jelentős eltéréseket mutatnak; vannak területek, ahol az egyik és vannak területek, ahol a másik szerepel jobban. A sztochasztikus modellek valószínűségi számítási háttérrel rendelkeznek, ami nem mondható el a PCM-ekről, így lehetőséget adnak előrejelzésekre. A PCM alapú modellek esetén a kiértékelhetőség sokkal egyszerűbben ellenőrizhető [3], illetve magának a kiértékelésnek is kevesebb a számításigénye. A konzisztencia mérés kezelésében a PCM alapú modellek jobban teljesítenek, míg a bizonytalanságot a sztochasztikus modellek jobban kezelik. Az összehasonlítási skálától jelentősen függenek a PCM alapú modellek, míg a sztochasztikus családban ezen függőség elhanyagolható, a relációs összehasonlításoknak köszönhetően.

A Thurstone-motivált modellek esetén a 3-opciós modellben még a közelmúltban sem volt ismert a szükséges és elégséges feltételrendszer, vagyis az, hogy mikor lehet kiértékelni az adathalmazokat. Továbbá a modell esetén nem vezettek be még konzisztencia fogalmat, mely szükséges az adatok tisztításához, és azok megbízhatóságáról való meggyőződéshez. A modelleket korábban nem vizsgálták információ visszanyerés szempontjából, sem optimális opciószám, sem optimális összehasonlítási struktúrák tekintetében. Ezen kutatási területeken le volt maradva a Thurstone-motivált modellek családja a PCM családtól. Kutatásomnál célul tűztem ki a modellek vizsgálatát az adatok kiértékelhetősége, különböző opciószámú modellek egymáshoz viszonyított viszonya és az alkalmazásuk tekintetében. Fontosnak éreztem továbbá a PCM alapú és Thurstone-motivált modellek közötti esetleges kapcsolatok felkutatását, ami a különböző filozófiájú módszerek közös gyökerének a megtalálását is szolgálja.

A páros összehasonlítási modelleknek rengeteg alkalmazása volt már korábban, legyen szó oktatásról [23, 24], sportról [25], azon belül labdarúgásról [26, 27, 28], információ visszanyerésről [29], energetikáról [30], pénzügyi szektorról [31], menedzsmentről [32], élelmiszeriparról [33], képek összehasonlításáról [34]. Szerettem volna bővíteni az alkalmazások körét, az alkalmazások segítségével szerettem volna bemutatni a Thurstone-motivált modellek előnyös tulajdonságait.

A dolgozat felépítése a következő. A 2. fejezetben bemutatom az általam kutatott modelleket és a PCM alapú modelleket. A 3. fejezetben az informatikai megoldásokra térek ki. A 4. fejezetben részletezem a kiértékelhetőséggel kapcsolatban elért eredményeimet. Az 5. fejezetben a konzisztencia, inkonzisztencia vizsgálat területén elért eredményeimet mutatom be. A 6. fejezetben az általam készített alkalmazásokból mutatok be néhányat. A 7. fejezetben összefoglalom az általam elért eredményeket.

2. A vizsgált modellek

2.1. Thurstone-motivált modellek

Ebben a fejezetben bemutatom a munkám során alkalmazott Thurstone-motivált modelleket és mindegyiknek saját jelölést adok, így a későbbi fejezetekben csak hivatkozni fogok rájuk. Ezen fejezet végén különféle ábrákon, táblázatokban bemutatom a különbségeket és hasonlóságokat.

A sztochasztikus háttérű páros összehasonlítási modelleknél az egyes rangsorolni, értékelni kívánt objektumok esetén valószínűségi változókat alkalmazunk, amelyek reprezentálják az adott objektumokat. Az egyes objektumokat jelöljük sorra $1, 2, 3, \dots, n$ -nel, a mögötte meghúzódo valószínűségi változókat pedig $\xi_1, \xi_2, \xi_3, \dots, \xi_n$ -nel. Várható értéküket, ami az átlagos erősségüket reprezentálja, jelöljük $m_1, m_2, m_3, \dots, m_n$ -nel. Az $\underline{m} = (m_1, m_2, m_3, \dots, m_n)$ vektort szeretnénk becsülni a kiértékelések során, ami megadja számunkra az objektumok egymáshoz viszonyított erősségeit, amiket keresünk az információ visszakeresés folyamán.

Amikor összehasonlítunk két objektumot, az i -ediket és a j -ediket, akkor a két objektum viszonyát vizsgáljuk, és a szóbeli értékelés mögött a valószínűségi változók különbsége húzódik meg. Az adott döntést a látens valószínűségi változók aktuális értékének különbsége határozza meg. Ez a különbség ebben az alakban írható fel:

$$\xi_i - \xi_j = m_i - m_j + \eta_{i,j}. \quad (1)$$

Az $\eta_{i,j}$ valószínűségi változókról feltételezzük, hogy azonos eloszlású valószínűségi változók közös F eloszlásfüggvénnyel.

Az alkalmazásokban általam használt modellekben két eloszlást alkalmaztam:

- Normális eloszlást, a Thurstone modellek esetén (TH);
- Logisztikus eloszlást, a Bradley-Terry modellek esetén (BT).

Az általános modellben az F eloszlásfüggvényre igaz, hogy $0 < F(x) < 1$, háromszor folytonosan differenciálható. Az $f = F'(x)$ sűrűségfüggvény 0-ra szimmetrikus, azaz $f(-x) = f(x)$, $x \in \mathbb{R}$ esetén és szigorúan log-konkáv. Az ilyen F eloszlásfüggvények halmazát jelöljük $\mathbb{F}$ -fel.

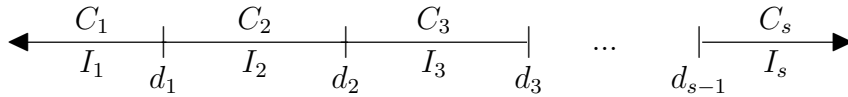
Az s -opciós modell esetén s lehetséges kimenetele lehet a döntésnek, pl. 4-opciós modellben C_1 = "sokkal rosszabb", C_2 = "kicsit rosszabb",

$C_3 = \text{"kicsit jobb"}$, $C_4 = \text{"sokkal jobb"}$. Legyenek az egyes kimenetek, vagyis döntések sorban $C_1, C_2, \dots, C_s$ és ezen döntésekhez tartozó intervallumok sorra $I_1, I_2, \dots, I_s$, ahol $I_l \cap I_k = \emptyset$ ha $l \neq k$ és $\mathbb{R} = \bigcup_{k=1}^s I_k$. Sorban $d_1, d_2, \dots, d_{s-1}$ -el jelöljük az egyes intervallumok határolópontjait, azaz az I_k -kat elválasztó pontokat. Ezekre teljesül, hogy $-\infty < d_1 < d_2 < \dots < d_{s-1} < \infty$. Azt feltételezzük, hogy C_k döntés akkor születik, ha a valószínűségi változók aktuális különbsége az I_k intervallumba esik.

Az i és j objektumok közti döntések szimmetriája miatt

- ha s páratlan, akkor $D = d_{(s-1)/2} < 0$, valamint $d_{s-l} = -d_l$, $l = 1, \dots, (s-1)/2$. Ekkor $C_{(s+1)/2}$ az "egyforma" opció.
- ha s páros, akkor $D = d_{s/2} = 0$, $d_{s-l} = -d_l$, $l = 1, \dots, s/2$.

A döntéseket és a nekik megfeleltetett intervallumokat az 1. ábra mutatja.



1. ábra. A döntések és a nekik megfeleltetett intervallumok az s -opciós modellben

Az alapmodell nem kezeli azt az esetet amikor valamelyik objektum előnyben van a másikkhoz képest, ilyen előny például a sakkban a kezdés joga, vagy más sportokban a hazai pálya. A modellbe előnykezelést is be lehet építeni. Ekkor a hazai csapat előnyben lehet, míg az idegen helyen játszó hátrányban. Nem kell sportra gondolnunk, hogy előnykezelés alkalmazása javasolt legyen. Több kutatás megállapítja, hogy akár az összehasonlítás sorrendjének is jelentős szerepe van az összehasonlításban, például az ízlésnél [35]. Az előny kezelését a $\underline{b}$ előnyt reprezentáló vektorváltozó segítségével vagyunk képesek beépíteni a modellbe.

Ha az i -edik csapat előnyben van (ξ_i^e), akkor növeljük erősségét a rá jellemző előnykezelés mértékével, b_i -vel, vagyis az objektum átlagos erőssége $m_i + b_i$ lesz.

Ha pedig a j -edik csapat hátrányban van (ξ_j^h), akkor csökkentjük az átlagos erősségét a rá jellemző mértékkel, h_j -vel, vagyis az objektum erőssége $m_j - h_j$ lesz hátrányban. Az általánosság korlátozása nélkül feltételezhetjük, hogy az előny mértéke megegyezik a hátrány mértékével. Vagyis ha lenne $\underline{h}$ hátrányvektor, ami a csapat hátránya esetén csökkentené a csapat erősségét,

akkor megmutatható, hogy $h_i = b_i$ feltételezhető. Emellett szinte mindig igaz, hogy ha az egyik objektum előnyben van, akkor a másik hátrányban van, például sportmérkőzések (hazai pálya előnye/idegen pálya hátránya). A sakkban ez erőteljesen megjelenik, ugyanis a fehérrel játszó előnye jelentős a fekete színnel játszó játékoskal szemben, így a végeredményt nagyban befolyásolja, hogy milyen színnel játszott az adott játékos. Tehát az összehasonlítások esetén a következőt kapjuk:

$$\xi_i^e - \xi_j^h = (m_i + b_i) - (m_j - b_j) + \eta_{i,j}. \quad (2)$$

Munkám során az előnykezelés 3 formáját különböztettem meg.

- Normál modell: nincs előnykezelés, vagyis minden $b_i = 0$, ez az alapmodell (N).
- Konstans modell: konstans előnykezelés van, vagyis minden objektumra azonos mértékű előnyt feltételezünk, $b_i = b, i = 1, 2, \dots, n$ esetén (K).
- Egyedi modell: egyedi előnykezelést alkalmaz, vagyis minden objektum esetén más-más mértékű előnyt jelent az előnyös helyzet. Azaz egyedi előnykezelést alkalmazunk, b_i -t (E).

Jól látható, hogy a K modell általánosítása az N modellnek, míg az E modell a K modellnek. A modellekre az alábbiéppen fogunk hivatkozni:

{modell típusa}_ {opciószám}_ {előnykezelés}.

Például a 3-opciós konstans előnyt alkalmazó Thurstone modell jelölése: TH_3_K.

Az összehasonlítások eredményeit reprezentálja az $\mathbf{A}^0$ $n \times n \times s$ dimenziós mátrix, ahol a 0 felső index mutatja azt, hogy nem különböztetjük meg, hogy előnyben vagy hátrányban történt-e az összehasonlítás. Az $\mathbf{A}^0$ mátrix $A_{i,j,k}^0$ eleme azt adja meg, hogy az i -edik objektumot hasonlítva a j -edik objektumhoz hányszor születik C_k döntés. A döntések szimmetriája miatt fennáll, hogy $k = 1, \dots, s$ esetén $A_{i,j,k}^0 = A_{j,i,s-k+1}^0$, ha $i \neq j$, és $A_{i,j,k}^0 = 0$, ha $i = j$.

Ha az i -edik objektum előnyben volt az összehasonlítás során, akkor az $\mathbf{A}^e$ mátrixot használjuk az adatok tárolására, ha pedig hátrányban, akkor az $\mathbf{A}^h$ mátrixot, ahol az e illetve a h felső index utal az i objektum helyzetére. Most is $k = 1, \dots, s$ esetén $A_{i,j,k}^e = A_{j,i,s-k+1}^h$, a főátlóbeli elemek pedig most is nullával egyenlők.

Mivel mindig a különbségekről döntünk, ezért rögzíthetjük az első objektum erősségét, legyen $m_1 = 0$.

A likelihood függvénnyel vagyunk képesek megadni, hogy mekkora az A^0, A^e, A^h összehasonlítási mátrix kialakulásának a valószínűsége az $\underline{m}, \underline{b}, \underline{d}$ bemeneti paraméterek mellett. A maximum likelihood becslést használjuk az objektumok erősségének meghatározásához [36].

Az N modellre felírva a likelihood függvényt, független döntéseket feltételezve:

$$L(A^0 | m_1, m_2, \dots, m_n, d_1, d_2, \dots, d_{s-1}) = \prod_{k=1}^s \prod_{i=1}^{n-1} \prod_{j=i+1}^n (P(\xi_i - \xi_j \in I_k))^{A_{i,j,k}^0}. \quad (3)$$

Az i és j objektumokhoz tartozó valószínűségi változók közötti különbségnek az I_k intervallumba esési valószínűsége $2 < s$ esetén a következőképpen fejezhető ki:

$$p_{i,j,1} = P(\xi_i - \xi_j \in I_1) = F(d_1 - (m_i - m_j)) \quad (4)$$

$$p_{i,j,k} = P(\xi_i - \xi_j \in I_k) = F(d_k - (m_i - m_j)) - F(d_{k-1} - (m_i - m_j)), k = 2, \dots, s-1 \quad (5)$$

és

$$p_{i,j,s} = P(\xi_i - \xi_j \in I_s) = 1 - F(d_{s-1} - (m_i - m_j)). \quad (6)$$

$s = 2$ esetén, $d_1 = 0$. Itt nem releváns (5). Ebben az esetben a (4) és (6) valószínűségek a következőképpen fejezhetők ki:

$$p_{i,j,1} = P(\xi_i - \xi_j \in I_1) = F(0 - (m_i - m_j)) \quad (7)$$

és

$$p_{i,j,2} = P(\xi_i - \xi_j \in I_2) = 1 - F(0 - (m_i - m_j)). \quad (8)$$

A likelihood becslés alapján kapott becslt paraméterérték előnyt nem alkalmazó modell esetén az a paramétervektor, amelynél a (3) likelihood függvény maximális, azaz

$$(\hat{\underline{m}}, \hat{\underline{d}}) = \arg \max_{\underline{m} \in \mathbb{R}^n, m_1=0, d_1 < d_2 < d_3 < \dots < d_{\lfloor s/2 \rfloor} = D} L(A^0 | \underline{m}, \underline{d}). \quad (9)$$

Ha egyedi előnyt alkalmazunk a modellben, akkor a likelihood függvény:

$$L(A^h, A^e | m_1, m_2, \dots, m_n, d_1, d_2, \dots, d_{s-1}, b_1, b_2, \dots, b_n) = \prod_{k=1}^s \prod_{i=1}^{n-1} \prod_{j=i+1}^n (P(\xi_i^e - \xi_j^h \in I_k))^{A_{i,j,k}^e} \cdot (P(\xi_i^h - \xi_j^e \in I_k))^{A_{i,j,k}^h} \quad (10)$$

Ebben az esetben az i és j objektumokhoz tartozó valószínűségi változók közötti különbségnek az I_k intervallumba esési valószínűségét megadó képlet kicsit változik; a $2 < s$ esetén a következőképpen fejezhető ki i előnye esetén:

$$p_{i,j,1}^e = P(\xi_i^e - \xi_j^h \in I_1) = F(d_1 - ((m_i + b_i) - (m_j - b_j))) \quad (11)$$

$$p_{i,j,k}^e = P(\xi_i^e - \xi_j^h \in I_k) = F(d_k - ((m_i + b_i) - (m_j - b_j))) - F(d_{k-1} - ((m_i + b_i) - (m_j - b_j))), k = 2, \dots, s-1 \quad (12)$$

és

$$p_{i,j,s}^e = P(\xi_i^e - \xi_j^h \in I_s) = 1 - F(d_{s-1} - ((m_i + b_i) - (m_j - b_j))). \quad (13)$$

Ha $s = 2$, akkor $d_1 = 0$. Ebben az esetben a (11) és (13) valószínűségek a következőképpen fejezhetők ki:

$$p_{i,j,1}^e = P(\xi_i^e - \xi_j^h \in I_1) = F(0 - ((m_i + b_i) - (m_j - b_j))) \quad (14)$$

és

$$p_{i,j,2}^e = P(\xi_i^e - \xi_j^h \in I_2) = 1 - F(0 - ((m_i + b_i) - (m_j - b_j))). \quad (15)$$

Hasonlóan felírhatóak abban az esetben is a valószínűségek, amikor az i hátrányban van.

A likelihood becslés eredménye egyedi előnyt alkalmazó modell esetén:

$$(\hat{\underline{m}}, \hat{\underline{d}}, \hat{\underline{b}}) = \arg \max_{\underline{m} \in \mathbb{R}^n, m_1=0, d_1 < d_2 < d_3 < \dots < d_{\lfloor s/2 \rfloor} = D, \underline{b} \in \mathbb{R}^n} L(A^h, A^e | \underline{m}, \underline{d}, \underline{b}). \quad (16)$$

A (3), (10) függvényeket szeretnénk maximalizálni, megtudni, hogy mikor a legnagyobb a függvény értéke, mikor alakulhat ki a legnagyobb valószínűséggel az A^0, A^e, A^h mátrix, az $\underline{m}, \underline{d}, \underline{b}$ paraméterek mellett. A függvény maximuma azonban nem mindig létezik és nem biztos, hogy egyetlen helyen veszi fel azt. Ehhez bizonyos feltételeknek eleget kell tenni az A^0, A^e, A^h összehasonlítási mátrixoknak. A feltételrendszereket az N modell esetében a következő fejezetben részletesen ismertetem.

A függvény számítása hatványozás miatt elég költséges, így érdemesebb a logaritmusát venni. Ez esetben a hatványkitevőkből szorzótényezők lesznek.

Mivel a logaritmus szigorúan monoton növekvő függvény, így csak a maximum értéke változik, a maximum helye nem.

Ha $b_i = 0$ $i = 1, 2, \dots, n$, azaz $\underline{b}$ nem szerepel a (10)-ben, akkor nem különböztetendő meg az előnyben/hátrányban megvalósuló összehasonlítások eredménye, így A^e és A^h helyett az A^0 használható. Így a (3) és a (10) likelihood függvények azonosak, így azonos eredményt ad a két becslés. Így látható, hogy az E általánosítása a K és az N modellnek.

A modellekbe még beépíthető az adatok időbeli súlyozása is, mert vannak olyan esetek, amikor az időben közelebbi összehasonlításokat nagyobb súllyal érdemes figyelembe venni, míg a távolabbiakat kisebbel. Ennek különösen előrejelzéseknél van jelentősége. Ehhez egy $0 < q \leq 1$ paramétert alkalmazok. Az egyes összehasonlítások súlyát q^l szorzóval módosítjuk, ahol l jelenti, hogy hány időegységnyi távolságra van az összehasonlítás időpontja az aktuális időponttól. Ha $q = 1$ akkor nem beszélünk időbeli súlyozásról. Ha $q < 1$, akkor időbeli súlyozásról beszélünk, melynek mértéke annál nagyobb, minél kisebb a q értéke.

Ha például 30 időegység távolságra van az összehasonlítás a jelenlegi pillanattól, akkor a korábbi $A_{i,j,k}^0$ helyére $A_{i,j,k}^0 \cdot q^{30}$ kerül, azáltal az időben közelebbi adatok nagyobb súllyal szerepelnek a kiértékelésben. Az időbeli súlyozást ugyanúgy alkalmaztam, ahogyan azt Hvattum [37] és Held [38] is tették a saját cikkeikben.

Az új $(v)A^0$ mátrix $(v)A_{i,j,k}^0$ értékeit így számítom:

$$(v)A_{i,j,k}^0(r) = A_{i,j,k}^0(r)(t) \cdot q^t, i, j = 1, \dots, n, k = 1, \dots, s, t = 0, \dots, T, \quad (17)$$

ahol t jelenti az eltelt időegységet az összehasonlítás pillanatától a jelen pillanatig, T a legnagyobb eltelt időegységet az összes összehasonlítás között, v jelzése a súlyozásnak, míg r jelenti azt, hogy az egyes megfigyelésekre külön alkalmazom. Az aggregált adatmátrix ezen $(v)A_{i,j,k}^0(r)$ adatmátrixok összege.

Súlyozás esetében kiértékelhető marad az adathalmaz, ha nem súlyozottan kiértékelhető volt, és ha nem volt kiértékelhető eredetileg, akkor nem válik azzá a súlyozás következtében sem, ugyanis a súlyozás nem változtat a kiértékelhetőségi feltételeken.

A $(v)A^0$ -hoz hasonlóan képezhetjük előnyben és hátrányban a $(v)A^e$, $(v)A^h$ mátrixokat. Az időbeli súlyozást, változást a későbbiekben V -vel jelölöm. Ha alkalmazok időbeli súlyozást, akkor a modell jelölése $_V$ -vel bővül, míg ha nem, akkor nem változik a jelölés.

Munkám során az általános s -opciós modellnek különféle változatait alkalmaztam kiértékelésekre, elemzésekre, szimulációkra. A következő

táblázatokban foglalom össze a hasonlóságait és különbségeiket.

Az általam konkrétan alkalmazott modellek típusai az 1. táblázatban találhatóak.

1. táblázat. Az általam alkalmazott modellek

	Modellek	
neve	Thurstone	Bradley-Terry
jelölés	TH	BT
eloszlás	normális	logisztikus

Az általam konkrétan alkalmazott modellek opciószámai a 2. táblázatban találhatóak.

2. táblázat. Különféle általam konkrétan alkalmazott modellek opciószámai

	Modellek				
opciószám	2	3	4	5	7
jelölés	2	3	4	5	7
döntések	jobb	jobb	sokkal jobb	sokkal jobb	sokkal jobb
			-	-	jobb
			kicsit jobb	kicsit jobb	kicsit jobb
	-	egyforma	-	egyforma	egyforma
			kicsit rosszabb	kicsit rosszabb	kicsit rosszabb
	rosszabb	rosszabb	-	-	rosszabb
			sokkal rosszabb	sokkal rosszabb	sokkal rosszabb

Az általam konkrétan alkalmazott modellek előnykezelésének fajtái a 3. táblázatban találhatóak.

3. táblázat. Modellek előnykezeléseinek fajtái

	Modellek		
jelölés	N	K	E
előnykezelés	nincs	konstans	egyedi
mértéke	$b_i = 0$	$b_i = b$	b_i

Az általam konkrétan alkalmazott modellek súlyozásának fajtái a 4.

táblázatban találhatóak.

4. táblázat. A modellek időbeli súlyozásai

	Modellek	
neve	súly nélküli	súlyozott
jelölés		V
q értéke	$q = 1$	$0 < q < 1$

A PCM alapú modelleknél a kiértékelés eredménye egy egyre normált csupa pozitív koordinátákból álló n hosszúságú vektor. Ahhoz, hogy a Thurstone-motivált modellek eredményeit össze tudjam vetni más modellekkel, a becsült várható értékekből a Thurstone-motivált modellek esetén is súlyok képzése ajánlott. A súlyvektort jelölje a $\underline{w}$ vektor, mely az egyes objektumok súlyait tartalmazza sorban. A súlyvektor számítása a következőképpen történik:

$$\underline{\hat{w}} = \left(\frac{\exp(\widehat{m}_1)}{\sum_{i=1}^n \exp(\widehat{m}_i)}, \dots, \frac{\exp(\widehat{m}_n)}{\sum_{i=1}^n \exp(\widehat{m}_i)} \right). \quad (18)$$

A $\underline{\hat{w}}$ vektorra igaz, hogy $\hat{w}_i > 0$ $i = 1, \dots, n$ esetén és $\sum_{i=1}^n \hat{w}_i = 1$. Az így kapott eredmények számszerűleg is összevethetők a PCM alapú módszerek kiértékelési eredményeivel. Mivel a súlyokat előállító transzformáció szigorúan monoton növekedő, ezért a $\hat{w}_i$ -k ugyanazt a sorrendet adják, mint a várható értékek vektora. Az m_i -k csak egymáshoz viszonyított erősségbeli különbségeket adnak meg számunkra. A súlyvektorképzés azt is biztosítja, hogy a súlyok értéke nem függ attól, hogy melyik m_i -t rögzítettük. Azért is érdemes ezen súlyvektorképzést alkalmazni, mert ez a Bradley-Terry modellben egyszerű számítási módra ad lehetőséget, ami miatt a Bradley-Terry modellt előszeretettel használják a Thurstone modellel szemben.

2.2. Speciális Thurstone-motivált modellek: a 2- és 3-opciós Bradley-Terry modell

A Thurstone-motivált modellek speciális esetei a Bradley-Terry modellek. Ezen modellek esetén az egyes valószínűségek egyszerűbb formában megadhatóak, mint más Thurstone-motivált modellek esetén.

A Bradley-Terry modellben az alábbi eloszlásfüggvényt használjuk:

$$F(x) = \frac{1}{1 + e^{-x}}, x \in \mathbb{R}. \quad (19)$$

Ha bevezetjük a

$$\pi_i = \frac{e^{m_i}}{\sum_{k=1}^n e^{m_k}} \quad (20)$$

jelölést, akkor ezen súlyokkal egyszerűen írhatóak fel a valószínűségek, így a likelihood függvény is könnyebben maximalizálható numerikusan. A 2-opciós modell valószínűségei felírhatóak az alábbi alakban:

$$p_{i,j,1} = P(i \text{ 'rosszabb' mint } j) = \frac{\pi_j}{\pi_i + \pi_j}, \quad (21)$$

$$p_{i,j,2} = P(i \text{ 'jobb' mint } j) = \frac{\pi_i}{\pi_i + \pi_j}. \quad (22)$$

Ebben az esetben teljesül a Luce-féle választási axióma [12], és annak általánosítása [13], vagyis az, hogy a 2-opciós modellben a két objektum összehasonlítási eredményének valószínűsége csak az egymáshoz viszonyított erősségein múlik. Ha vesszük a két objektum közül az egyik illetve a másik választásának valószínűségét és ezeket egymással elosztjuk, akkor a kapott érték a két objektumhoz rendelt súlyok aránya lesz, más paraméterek nem befolyásolják. Tehát

$$\frac{p_{i,j,2}}{p_{i,j,1}} = \frac{\frac{\pi_i}{\pi_i + \pi_j}}{\frac{\pi_j}{\pi_i + \pi_j}} = \frac{\pi_i}{\pi_j}, \quad (23)$$

ami a Luce-féle választási axióma teljesülését jelenti. A 3-opciós modellben ahhoz, hogy kezelni tudjuk a három opciót, szükségünk van a $\theta = e^d, \theta > 1$ változóra. Ezzel a változóval felírva a valószínűségeket a 3-opciós modellben ezt kapjuk:

$$p_{i,j,1} = P(i \text{ 'rosszabb' mint } j) = \frac{\pi_j}{\pi_i \theta + \pi_j}, \quad (24)$$

$$p_{i,j,2} = P(i \text{ 'egyforma' } j\text{-vel}) = \frac{\pi_i \pi_j (\theta^2 - 1)}{(\pi_i \theta + \pi_j)(\pi_j \theta + \pi_i)}, \quad (25)$$

$$p_{i,j,3} = P(i \text{ 'jobb' mint } j) = \frac{\pi_i}{\pi_j\theta + \pi_i}. \quad (26)$$

3-opciós modellben már nem csak az objektumok erősségétől függenek a valószínűségek, hanem a θ paramétertől is, ezt pedig befolyásolják a többi objektum összehasonlításából származó eredmények. Tehát

$$\frac{p_{i,j,3}}{p_{i,j,1}} = \frac{\frac{\pi_i}{\pi_j\theta + \pi_i}}{\frac{\pi_j}{\pi_i\theta + \pi_j}} \neq \frac{\pi_i}{\pi_j}, \quad (27)$$

így már csak speciális esetben teljesül a Luce-féle választási axióma, pl. ha minden objektum azonos erősségű. Davidson pontosan ezért alkotta meg a maga modelljét, hogy a Luce-féle választási axióma teljesülését biztosítsa.

A likelihood függvény ezen formája esetén fixpont-iteráció is felírható, így fixpont-iterációs módszerrel is számolhatjuk a súlyokat, és később belőlük az erősségeket. Mivel ilyen módon sokkal gyorsabb a számítása az erősségeknek, így én is előszeretettel alkalmaztam a különféle Bradley-Terry modelleket nagyszámú szimulációk esetében.

2.3. A Davidson modell

Davidson a 3-opciós modell esetére szeretne volna, hogy a Luce féle választási axióma teljesüljön, így létrehozta saját modelljét [11].

1. Tétel. *A Davidson modell nem a Thurstone-motivált modellek családjába tartozik.*

A bizonyítás a [39] publikációban található. Ennek ellenére a modell sok hasonlóságot mutat a Thurstone-motivált modellekkel.

A modell a következő:

Legyen $\underline{\pi} = (\pi_1, \dots, \pi_n)$ súlyvektor, amire igaz, hogy $0 < \pi_i$ és $\sum_{i=1}^n \pi_i = 1$. Az objektumok sorra $1, 2, \dots, n$, az i -edik objektum súlyát π_i reprezentálja. A összehasonlítás során a "rosszabb", "egyforma", és "jobb" döntések valószínűsége legyen a következő:

$$p_{i,j,1} = P(i \text{ 'rosszabb' mint } j) = \frac{\pi_j}{\pi_i + \pi_j + \nu \sqrt{\pi_i \pi_j}}, \quad (28)$$

$$p_{i,j,2} = P(i \text{ 'egyforma' } j) = \frac{\nu \sqrt{\pi_i \pi_j}}{\pi_i + \pi_j + \nu \sqrt{\pi_i \pi_j}}, \quad (29)$$

$$p_{i,j,3} = P(i \text{ 'jobb' mint } j) = \frac{\pi_i}{\pi_i + \pi_j + \nu \sqrt{\pi_i \pi_j}}, \quad (30)$$

ahol ν egy pozitív paraméter, mely az "egyforma" döntések valószínűségét hivatott kezelni. Ha ν kicsi akkor arányaiban kevés "egyforma" döntés született, míg ha ν nagy, akkor arányaiban sok "egyforma" döntés született az összehasonlítások során. Ha $\nu = 0$, akkor a 2-opciós Bradley-Terry modellt kapjuk vissza. Ez a modell tehát, hasonlóan a BT3 modellhez, a BT2 modell általánosítása.

Ha π_i értékeket megszorozzuk egy pozitív számmal akkor a valószínűségek nem változnak.

A (28), (29) és (30) valószínűségek miatt látható, hogy

$$\frac{p_{i,j,2}}{\sqrt{p_{i,j,1} \cdot p_{i,j,3}}} = \nu. \quad (31)$$

Mivel

$$\frac{p_{i,j,3}}{p_{i,j,1}} = \frac{\pi_i}{\pi_j}, \quad (32)$$

a modellben a Luce féle választási axióma teljesül [12, 13]. Davidson pontosan ezzel a tulajdonsággal indokolta a modell bevezetését [11].

Az összehasonlítások eredményeit a korábbiakhoz hasonlóan az A^0 mátrixban tároljuk.

Az adatok hasonlóképpen értékelhetőek ki a maximum likelihood becsléssel, mint a Thurstone-motivált modelleknél. A likelihood függvény a Davidson modellben:

$$L(A|(\pi_1, \dots, \pi_n, \nu)) = \prod_{k=1}^3 \prod_{i=1}^{n-1} \prod_{j=i+1}^n (p_{i,j,k})^{A_{i,j,k}^0}. \quad (33)$$

A maximum likelihood becslés a Davidson modell esetén:

$$(\hat{\pi}, \hat{\nu}) = \arg \max_{\pi > 0, \nu > 0, \sum_{i=1}^n \pi_i = 1} L(A^0|(\pi_1, \dots, \pi_n, \nu)). \quad (34)$$

Davidson [11] cikkében egy elégséges feltételrendszert adott a maximum likelihood becslés létezésére és egyértelműségére, amelyet a 4.1 fejezetben mutatok be. A közelmúltig bezárólag ezt a feltételrendszert használták az alkalmazások során annak ellenőrzésére, hogy a maximum likelihood becslés végrehajtható-e.

2.4. Páros összehasonlítási mátrixokon alapuló modellek

Napjainkban nagyon elterjedtek a páros összehasonlítási mátrixokon alapuló módszerek. Saaty 1990-ben megjelent cikkére [2] több tízezer hivatkozás született a Google Scholar szerint. Bár a munkámban ezeket csak összehasonlítás érdekében használtam, röviden mégis bemutatom őket.

Az AHP [1, 2] modell az egyik legelterjedtebb páros összehasonlítási modell. A modell esetén, amikor összehasonlítunk két objektumot, i -t és j -t, arról születik döntés, hogy i hányszor jobb, mint j . Általában az 1, 3, 5, 7, 9-es számokat engedik meg, illetve ezek reciprokát. Ezen számokból egy $n \times n$ -es mátrixot képeznek, amit páros összehasonlítási mátrixnak (PCM) neveznek. A mátrix reciprok-szimmetrikus tulajdonságú, ami azt jelenti, hogy $a_{i,j} = 1/a_{j,i}$, minden i, j pár esetén. Ezt a páros összehasonlítási mátrixot általában a sajátvektor módszer segítségével értékelik ki [40]. A módszer egyszerűsége miatt nagy népszerűségnek örvend, az irodalomban igen sokféle alkalmazásra találunk példát. A módszer hátránya, hogy csak teljes összehasonlítás esetére alkalmazható, vagyis akkor, amikor minden objektum minden más objektummal össze lett hasonlítva, ami azonban nem mindig áll fenn, különösen nagyszámú összehasonlítandó objektum esetén. A módszer egyik központi kérdése az összehasonlítási eredmények konzisztenciája; erre a későbbiek során visszatérünk.

A páros összehasonlítási mátrixok más módon is megalkothatók, a követelmény velük szemben, hogy az elemei pozitívak legyenek, és teljesítsék a reciprok-szimmetrikus tulajdonságot. Ha vannak nem összehasonlított objektumok, akkor a hiányzó elemek helyét üresen hagyják.

Bozóki és társai az LLSM-et javasolták a PCM-ek kiértékelésére [3], ami szintén egy optimalizációs módszer. Ennek segítségével nemteljes összehasonlítások esetére is kiértékelhetővé vált a modell. Emellett azt a páros összehasonlítási mátrix képzést javasolták, amiben az $a_{i,j}$ arányt az i "jobb", mint j és az i "rosszabb", mint j döntések számának hányadosaként értelmezik, például ha két objektum között hét esetben volt "jobb" az egyik és öt esetben a másik, akkor $\frac{7}{5}$ legyen a PCM mátrix adott helyén az érték. Ezzel kibővítették a PCM mátrixba beírható értékek tartományát. Abban az esetben, ha egyik 0-szor volt jobb mint a másik, ezen fajta hányadosképzés nem működik, és a PCM elemet vagy mesterségesen kell létrehozni, vagy ki kell hagyni. A kihagyás persze komoly információvesztést okoz, a mesterséges értékek beírása pedig torzulást okozhat. Bozóki és társai [3]-ban azt is ellenőrizték, hogy mikor értékelhető ki az LLSM segítségével egy PCM. Bebizonyították, hogy akkor

alkalmazható az LLSM kiértékelési módszer, hogy ha az összehasonlításokból képzett gráf összefüggő. Ez a gráf a következőképpen hozható létre:

A gráf csúcsai az összehasonlított objektumok, él van behúzva két csúc között, ha a megfelelő mátrixelem ki van töltve.

Felhívjuk a figyelmet arra, hogy a PCM alapú módszerek a PCM kialakításakor a döntéseket direkt módon számmá konvertálják. Ez nem áll fenn a Thurstone-motivált modellek esetén, ott a döntések relációknak felelnek meg.

Felmerül a kérdés, hogy a kialakított páros összehasonlítási mátrix milyen mértékben tartalmaz/ nem tartalmaz ellentmondásokat. Ezt a konzisztencia fogalmán keresztül vizsgálják, és különböző konzisztencia mérőszámokon keresztül mérik [4].

A konzisztencia fogalmát először Saaty vetette fel az AHP esetén [1]. Az adatok ellentmondásmentessége teljes összehasonlítás esetén azt jelenti, hogy a páronkénti összehasonlításokból nem keletkezik logikai ellentmondás. Ez egy példán keresztül bemutatva azt jelenti, hogy ha i 3-szor olyan jó, mint k és k 5-szor olyan jó, mint j , akkor i $3 \cdot 5 = 15$ -ször olyan jó, mint j . Ez már az AHP esetén gyakran használt $\{1, 3, \dots, 9, \frac{1}{3}, \dots, \frac{1}{9}\}$ értékek esetén sem mindig áll fenn.

Teljes összehasonlítás esetén konzisztenciáról akkor beszélünk, ha

$$a_{i,j} = a_{i,k} \cdot a_{k,j}, \text{ minden } i, j, k\text{-ra.} \quad (35)$$

Nemteljes összehasonlítások esetén, mivel előfordulhat eset, amikor i össze van hasonlítva j -vel és k -val, de a k nincs összehasonlítva j -vel a fenti definíció nem alkalmazható. Helyette Bozóki és társai [41]-ben bevezették azt a definíciót, amely az összehasonlítási gráfban szereplő körök mentén kívánják meg a mátrixelemek szorzatáról, hogy eggyel legyenek egyenlők. Azaz nemteljes összehasonlítás esetén a PCM mátrix elemei akkor konzisztensek, ha minden az összehasonlítási gráfban levő $(i_1, i_2, \dots, i_l, i_1)$ kör mentén

$$a_{i_1, i_2} \cdot a_{i_2, i_3} \cdot \dots \cdot a_{i_l, i_1} = 1. \quad (36)$$

teljesül.

3. Informatikai megvalósítás

A különféle páros összehasonlítási modellek esetén a kiértékelés az esetek többségében igen számításigényes, nehéz feladat. Ezért szükség van a számítások számítógépes elvégzésére. Ezen felül a legtöbb esetben ezek a modellek szimulációk segítségével érthetőek meg jobban, vagy hasonlíthatóak össze. Szimulációk segítségével sejthetünk meg feltételeket a modellek kapcsolatban. Ha egy kiértékelés jelentős erőforrás-igénnyel rendelkezik, akkor nagyszámú szimuláció esetén szükséges hatékony alkalmazás elkészítése, mely a futtatásokat kezeli, és azokat optimálisan minél gyorsabban végrehajtja.

Elkészítettem egy olyan alkalmazást/függvényekből álló nagy csomagot, amely egymással kompatibilis részeit könnyen össze lehet fűzni a különféle szimulációk, futtatások, tesztek elvégzéséhez.

Programozási nyelvnek a C#-ot választottam gyorsasága és elegáns megoldásai miatt.

Normális eloszláshoz segítségül a Math.NET Numerics függvény könyvtárat használtam, míg az adatok json formátumban történő írásáért és olvasásáért a Json.NET - Newtonsoft felelt. A GPU programozáshoz az ILGPU megoldását alkalmaztam.

Az alkalmazás-csomag hat fő részből áll.

Az első az adatok beolvasásáért felelős modul. Ez adatokat, például sportolók eredményeit, különféle adatformátumokból képes beolvasni, és egy olyan adathalmazba rendezi az adatokat, amely alkalmas a további feldolgozásra. Összehasonlítási mátrixban van tárolva az egymás ellen elért eredmény, míg az egyes objektumok egyéb tulajdonságai, illetve a modell tulajdonságai az adathoz kapcsoltnak vannak tárolva. Jelenleg az általam leprogramozott modell 7-opciót enged meg előnykezelés mellett maximálisan, így ehhez van igazítva az adattárolás, mely egységes a különféle modellek esetén.

A második fő rész az adatok generálása modul. Bizonyos esetekben szükséges nagyszámú szimuláció elvégzése. Ehhez külön adatgeneráló modult készítettem, mely sokféleképpen képes randomizált adatokat generálni az elvárt információk figyelembevételével. Ebben a részben kapott helyet az adatok súlyozását vezérlő almodul is.

A harmadik rész a feltételrendszerek ellenőrzésével kapcsolatos modul. Ez azért is szükséges, mert ezen modul segítségével ellenőrizhető, hogy egy adathalmaz kiértékelhető-e, van-e egyértelmű maximumhelye a likelihood függvénynek. Az előnyt nem alkalmazó modellek esetén a 2-, 3-opciós modellekhez van a szükséges és elégséges feltételeket ellenőrző kiértékelő, míg a 4-, 5-, 6-, 7-opciós modellek esetén csak elégséges feltételrendszer segítségével ellenőrzöm az adathalmazt. Előnyt, hátrányt alkalmazó esetekben elégséges

feltételrendszert alkalmazok az adathalmaz kiértékelhetőségének vizsgálatára. Emellett különféle maximumhely keresőket is alkalmazok annak eldöntésére, hogy egy maximumhely létezik és az egyedi-e a nem elégséges esetekben. Ezen modul kimondottan bonyolult módosított szélességi kereső gráf algoritmusokat tartalmaz, melyeket az adott igényekhez módosítottam, hogy azok a lehető leggyorsabbak legyenek.

A következő rész a kiértékelő modul. Az egyik leginkább erőforrás-igényes számítás a maximum likelihood becslések elvégzése. Maga a likelihood függvény számítása is elég költséges, ezért érdemes a logaritmusát venni, ahol a szorzás összeadássá és a hatványozás szorzássá szelődül. Több tíz változóban (esetenként 100-as nagyságrendben) kell keresni az optimum helyét m , b , d változóiban, továbbá esetenként bizonyos feltételeket is teljesíteni kell a bemeneti adatnak, ezért szükséges numerikus optimumhely keresőt alkalmazni a becslés elvégzésekor. Mivel a sűrűségfüggvény szigorúan log-konkáv, így ha bizonyos feltételeknek eleget tesz a bemeneti összehasonlítás mátrix, akkor egyértelmű maximumhelye lesz a likelihood/log-likelihood függvénynek. Bár a Matlab, az R és a Python programok kínálnak sok beépített optimalizálót, de nagy szimuláció szám esetében, ahol az idő kritikus tényező, szükség van gyors, megbízható megoldásra. Olyan numerikus szélsőérték kereső módszert kerestem ehhez a feladathoz, amely gyors, módosítható és log-konkáv esetben jól konvergál. A Nelder-Mead módszerre esett a választásom [42] általános esetben, míg a Bradley-Terry modell esetén fixpont-iterációs megoldást alkalmazok.

A program rendkívüli sebességet képes elérni a gyorsaságra optimalizált Nelder-Mead kódnak és a nagyfokú párhuzamosításnak köszönhetően. Időről-időre GPU programozást is bevettem a sebesség növelése érdekében.

A Nelder-Mead optimalizálót nem csak a Thurstone-motivált modelleknél, de a Davidson modelleknél és az LLSM, AHP-LLSM módszereknél is alkalmazni tudtam. Mivel bizonyos esetekben 10^{10} darab szimulációt is le kellett futtatnom, így szükségem volt egy még gyorsabb megoldásra is. A Bradley-Terry modell esetén lehetőség van arra, hogy parciális deriváltak felhasználásával fixpont-iterációs módszert használjak a 2-, 3-, 4- és 5-opciós modellekre, melyek további jelentős gyorsulást eredményeztek. A fixpont-iterációs megoldás 3-opciót alkalmazó modellek esetén átlagosan 50-szeres gyorsulást eredményezett, amivel így hasonló futtatási idő mellett majdnem két nagyságrenddel nagyobb szimulációs szám is elérhetővé vált.

Az egyszerű implementálás és nagyfokú kompatibilitás miatt az ILGPU megoldását használtam. Ezen függvénykönyvtár segítségével átlagosan 50-szeres sebesség növekedést értem el nagyszámú szimuláció esetén. Ezen megoldást azon esetben alkalmaztam előszeretettel, amikor 10^6 -t meghaladó számú maximum likelihood kiértékelést kellett végezni. A GPU-t csak a

kiértékelésekhez alkalmaztam, a többi modul továbbra is CPU-n futott.

A következő programrész az adatok kiértékeléséhez, elemzéséhez kapcsolódó modul. Sok esetben összehasonlításokat kellett végezni különböző kiértékelési módszerek segítségével kapott eredmények közt. Euklideszi távolságot, Spearman, Pearson, Kendall rangkorrelációt számoltam. Ezen mérőszámok számolását vezérlő kódok ebbe a modulba kerültek, ahonnan könnyedén meghívhatóak a különféle eredményekre, rangsorokra alkalmazva. Emellett ide kerültek a jövőbeli események előrejelzésével kapcsolatos programrészletek is.

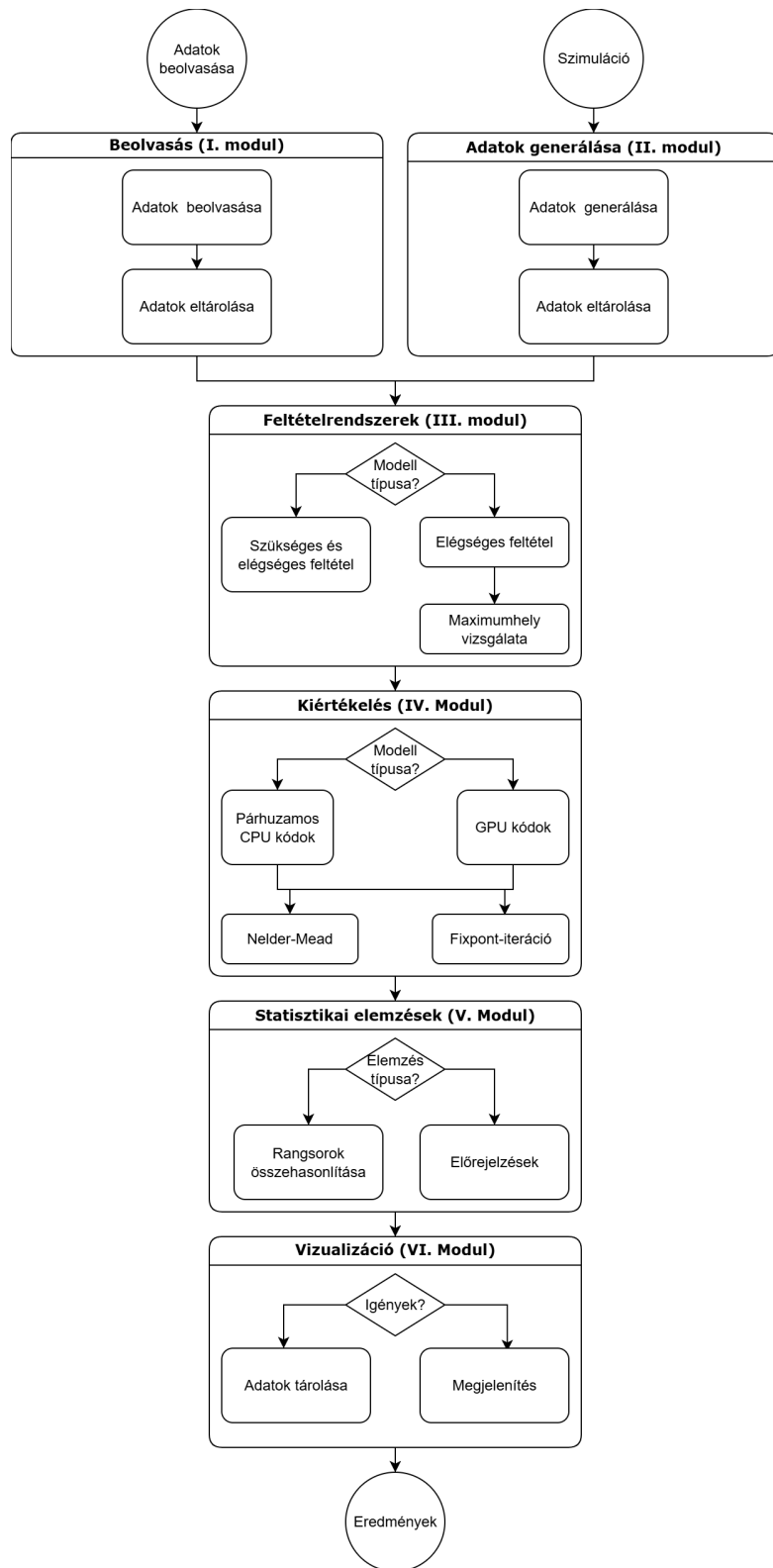
A hatodik modul pedig az adatok megjelenítését hivatott szolgálni. Itt a kapott eredményeket az igényeknek megfelelően fájlba mentem. A fájl formátuma és kinézete nagyon különböző az egyes esetekben. Vannak olyan megoldások, ahol html kódot generáltam, van olyan megoldás, ahol word dokumentumba mentettem a kész táblázatokat, és vannak egyszerűbb csv fájlba mentett kimeneti adatok és json fájlok is. Utóbbiak főleg az újra futtathatósággal kapcsolatos adatokat tárolják.

A teljes programcsomag sokszor grafikus felületet is kapott, de az esetek többségében, mikor szimulációk elvégzése volt a feladat, csak parancssoros felületet kapott az adott program, ahol a szimuláció előrehaladása nyomon követhető volt. Mivel sokszor a szimulációk akár napokig is futottak, így szükséges volt a részeredmények időről-időre történő automatikus mentése, esetleges áramszünet, vagy más jellegű meghibásodásra való tekintettel.

Az elkészült programcsomaggal képes voltam különféle modellek esetén nagyszámú szimulációt elvégezni, relatív rövid idő alatt, és azok eredményeit könnyen értelmezhető módon megjeleníteni a későbbi elemzések lehetőségének biztosítására.

A 2. ábrán látható a programcsomag folyamatábrája.

A dolgozat további fejezeteiben bemutatott eredmények nagy része szervesen, vagy teljes mértékben kötődik a szimulációkhoz és a különféle adathalmazok kiértékeléséhez, melyeket ezen programcsomag segítségével végeztem el.



2. ábra. A programcsomag egyszerűsített folyamatábrája

4. Az adatok kiértékelhetőségével kapcsolatos kutatások

Mit jelent a kiértékelhetőség? Miért is fontos tudni azt, hogy mikor kiértékelhető egy adott összehasonlítási adathalmaz? Az előző fejezetben bemutatott modelleknél a paraméterbecslést a maximum likelihood módszerrel végeztük. A maximum likelihood becslés egy gyakran alkalmazott, általános esetben is rendkívül jó tulajdonságokkal rendelkező, optimalizáláson alapuló becslési módszer [36]. Egy bonyolult függvényt kell esetünkben optimalizálni. Felmerül ezért az a kérdés, hogy az optimalizálandó függvény felveszi-e a maximumát és a maximum helye egyértelmű-e. Ez attól függ, hogy milyen adathalmaz áll rendelkezésre. A módszer páros összehasonlítások esetén nem minden adathalmaz esetén ad megoldást. Ha például a fagyaltos példában a vaníliát mindenki "jobb"-nak mondja a csokoládé ízűnél, akkor tudjuk a sorrendet, de az erősségekről, vagyis, hogy *mennyivel* jobb, nincs információnk. Ez esetben az adataink nem kiértékelhetőek, vagyis a relációkból, mint rendelkezésre álló összehasonlítási eredményekből nem kapunk számszerűsített erősségeket.

Akkor nevezünk egy adathalmazt kiértékelhetőnek, ha van egyértelmű maximumhelye a likelihood függvénynek, vagyis pontosan egy olyan $(\underline{m}, \underline{d}, \underline{b})$ paraméter vektor van, amelyben a likelihood függvény felveszi a legnagyobb értékét. Ha kiértékelhetőek az összehasonlítások, akkor egyértelműen meg tudjuk adni az egyes erősségeket, a határoló pontokat, míg ha nincs maximumhely vagy az nem egyértelmű, akkor nem tehető meg. Amikor sok összehasonlítási adat áll rendelkezésre, akkor általában kiértékelhetőek az adatok. Amikor kevés információ áll rendelkezésre (kevés az összehasonlítás), akkor ez nem biztos, hogy teljesül.

A kiértékelhetőség témájában az elmúlt 5 évben rengeteg kutatást végeztem. Kutatásom egyik fő célja a kiértékelhetőség szükséges és elégséges feltételrendszerének megtalálása volt. Folyamatosan sikerült általánosítani azt a feltételrendszert, amely teljesülése esetén kiértékelhető egy adott összehasonlítási adathalmaz. Négy cikkben jelent meg kiértékelhetőséggel kapcsolatos eredményünk [43, 44, 45, 39]. A kérdés, amire a választ kerestük, az, hogy mikor elegendő a rendelkezésre álló, relációkat tartalmazó információmennyiség arra, hogy számszerűsítsuk az objektumok egymáshoz viszonyított erősségét. Ebben a fejezetben bemutatom az eredményeket. Először azokat a feltételrendszereket mutatom be, amelyeket kutatásba való bekapcsolódásom előtt dolgoztak ki és publikáltak, majd rátérek a saját

eredményeimre.

4.1. Korábbi feltételrendszerek

A kiértékelhetőség feltételei egyrészt az alkalmazott opciószámától, másrészt az eloszlásfüggvény típusától függően találhatók az irodalomban. Ebben a fejezetben mindig az alapmodellt vizsgálom, így az adatmátrix felső 0 indexét nem használom, helyette felül zárójelben az opciószámot tüntetem fel.

Korábbi feltételrendszerekből ebben az alfejezetben hármat mutatok be. Az első az 1957-ben Ford [8] által ismertetett szükséges és elégséges feltétel a 2-opciós Bradley-Terry modellre. Erre Ford-féle feltétel néven hivatkozok. A második 3-opciós modell esetén 1970-ben Davidson által publikált elégséges feltételrendszer [11], melyre Davidson-féle feltételrendszer néven hivatkozok majd. A harmadik a 2019-ben konzulenseim által publikált elégséges feltételrendszer általános s -opciós modell esetén [21]. Mihálykó-Orbán néven fogok erre a feltételrendszerre hivatkozni.

- Ford-féle feltétel:

Ford 1957-ben mutatta be 2-opciós Bradley-Terry modellre vonatkozó a szükséges és elégséges feltételét. A bizonyításnál erősen támaszkodott a likelihood függvény speciális alakjára, így ez a bizonyítás nem általánosítható. A feltételek a következőképpen fogalmazhatók meg:

1. Definíció. *Definiáljuk a $\vec{G}^{(2)}$ gráfot, ami egy összehasonlítási eredményeket reprezentáló irányított gráf. A gráf csúcsai legyenek az objektumok. Legyen a gráfban i -ből j -be mutató irányított él, ha $A_{i,j,2}^{(2)} > 0$.*

A Ford-féle feltétel azt jelenti, hogy a $\vec{G}^{(2)}$ gráf erősen összefüggő. Másképpen megfogalmazva, bárhogyan bontjuk két diszjunkt nem üres részhalmazzra az objektumok (csúcsok) halmazát, ha van az egyikből a másikba és a másikkól az egyikbe mutató irányított él, akkor teljesül a feltétel.

Ford az állította, hogy az A adathalmaz pontosan akkor kiértékelhető, ha a $\vec{G}^{(2)}$ gráf erősen összefüggő [8].

- Davidson-féle feltételrendszer:

Davidson 1970-ben publikálta feltételrendszerét, mely a Ford-féle feltételrendszer 3-opciós modellre történő adaptációja.

2. Definíció. *Definiáljuk a $\vec{G}^{(3)}$ irányított gráfot. Az objektumoknak egyes csúcsok vannak megfeleltetve. Két csúcs, i és j között i -ből j -be mutató irányított él van, ha $0 < A_{i,j,3}^{(3)}$.*

Megjegyezzük, hogy a gráfban a "jobb" objektumtól a "rosszabb" felé mutató él van. Előfordulhat persze, hogy mindkét irányú él létezik két objektum között. Az adatok kiértékelhetőségének Davidson-féle feltételrendszere a következő:

- a) van "egyforma" összehasonlítási eredmény valamely objektumpár között
- b) a $\vec{G}^{(3)}$ irányított gráf erősen összefüggő

Davidson tétele azt mondja ki, hogy ha ezen a), b) feltételek teljesülnek, akkor az $A^{(3)}$ adathalmaz kiértékelhető.

- Mihálykó-Orbán- féle feltételrendszer:

A Mihálykó-Orbán feltételrendszert konzulenseim 2019-ben publikálták. Ez a Thurstone-motivált modellek esetén ad elégséges feltételt az általános szigorúan log-konkáv sűrűségfüggvény esetén. Most a 3-opciós modellre vonatkozó feltételrendszert ismertetem.

A feltételrendszer a következő:

3. Definíció. *Definiáljuk a következő $G^{(3)}$ gráfot. Az objektumoknak feleljenek meg az egyes csúcsok. Az i és j csúcs között legyen él behúzáva, ha $A_{i,j,2}^{(3)} > 0$, vagy ha $A_{i,j,1}^{(3)} \cdot A_{i,j,3}^{(3)} > 0$.*

Szóban megfogalmazva: akkor van éllel összekötve két csúcs, ha van "egyforma", vagy van mindkét irányban "jobb" döntés közöttük. A [21] cikkben bemutatott feltételrendszer a következő:

- a) van legalább egy (i, j) pár, amire $A_{i,j,2}^{(3)} > 0$
- b) van legalább egy (i, j) pár amire $A_{i,j,1}^{(3)} \cdot A_{i,j,3}^{(3)} > 0$

– c) a $G^{(3)}$ gráf összefüggő.

A [21]-ben bizonyított állítás pedig az, hogy ha teljesül a fenti 3 feltétel, akkor az adatok kiértékelhetők a Thurstone-motivált modellek esetén.

A Davidson és Mihálykó-Orbán feltételrendszerek nagy segítséget nyújtottak a kiértékelések során annak eldöntésében, hogy elvégezhető-e a kiértékelés, vagyis ténylegesen van-e maximumhelye a likelihood becslésnek. Azonban sok esetben nem lehetett eldönteni ezek alapján, hogy kiértékelhető vagy sem az adathalmaz, mert bár numerikusan kaptunk eredményt, de kérdéses volt, hogy ez egy tényleges és egyértelmű maximumhely-e. Ennek okán szükség volt a feltételrendszerek fejlesztésére, hogy kevés információ esetén is biztosan el tudjuk dönteni, hogy kiértékelhető-e vagy sem az adathalmaz.

A munkám során alkalmazott feltételrendszereket a következő 5. táblázat mutatja be. A későbbiekben az itt bemutatott rövidítéseket használok táblázatokban és ábrákon.

5. táblázat. A kiértékelhetőségi feltételrendszerekkel kapcsolatos eredmények áttekintése

Model	Opciószám	Feltétel nevének rövidítése és / Feltétel típusa	Szerző(év)
BT	2	Ford szükséges és elégséges	Ford (1957)
Davidson	3	DAV elégséges	Davidson (1970)
Thurstone-motivált	3	MO elégséges	Mihálykó-Orbán et al. (2019)
Thurstone-motivált	2	Ford szükséges és elégséges	Gyarmati et al. (2023)
Thurstone-motivált	3	SC elégséges	Gyarmati et al. (2023)
Thurstone-motivált Davidson	3	NS szükséges és elégséges	Gyarmati et al. (2026)

4.2. A 2-opciós modellre vonatkozó Ford tétel általánosítása

Ford 1957-ben [8]-ban ismertette a 2-opciós Bradley-Terry modell esetére az adatok kiértékelhetőségét biztosító szükséges és elégséges feltételt. A bizonyítás során kihasználta a modell speciális tulajdonságait. Azonban az állítás általános $F \in \mathbb{F}$ esetén is igaz. Az általánosítás bizonyítással együtt megtalálható a [45] publikációnkban.

A tétel a következő:

2. Tétel. *Legyen $s=2$ és $F \in \mathbb{F}$. A maximum likelihood becslés (MLE) létezésének és egyértelműségének szükséges és elégséges feltétele az $m_1 = 0$ feltétel mellett a következő: az objektumok S és $\bar{S}$ tetszőleges, nem üres felosztása esetén legalább egy olyan S -beli elem létezik, amely "jobb", mint $\bar{S}$ valamely eleme, és legalább egy olyan $\bar{S}$ -beli elem is létezik, ami "jobb" valamely S -beli elemnél. Másként fogalmazva, az adatok kiértékelhetőségének szükséges és elégséges feltétele egy paraméter rögzítése után az 1. Definícióban lévő $\vec{G}^{(2)}$ gráf erős összefüggősége.*

A részletes bizonyítást a [45] publikáció tartalmazza. A bizonyításban az erős összefüggőség alapján megmutatható, hogy az $\underline{m}$ paraméter koordinátái korlátos zárt halmazba foglalhatók, vagyis azon kívül biztosan kisebb értéket vesz fel a likelihood és a log-likelihood függvény, mint a halmazon belül. Mivel a likelihood és a log-likelihood függvény folytonos, így a korlátos zárt halmazban felveszi a maximumát. A maximumhely egyértelműséget a log-konkáv mértékek elmélete alapján a sűrűségfüggvény szigorú log-konkávítása biztosítja [46].

Ford 1957-es bizonyításában felhasználta a Bradley-Terry modell (21), (22) alakjait, így az általa adott bizonyítás csak a logisztikus eloszlás esetén volt érvényes. Ezen általánosítás segítségével már nem csak logisztikus, hanem normális eloszlásra (Thurstone modell), gamma eloszlásra (Stern) [47] is alkalmazható a szükséges és elégséges feltételrendszer.

A tétel alapján látható, hogy az a lényeges, hogy bizonyos esetekben a megfelelő döntések számának pozitivitása (nem 0 tulajdonsága) teljesüljön, maga az egyes $A_{i,j,k}$ -k konkrét értéke nem számít a kiértékelhetőség ellenőrzésekor.

4.3. Egy 3-opciós elégséges feltételrendszer

Ebben a fejezetben bemutatok egy újabb feltételrendszert, amely általánosítása a Davidson és Mihálykó-Orbán feltételrendszereknek. A feltételrendszert a [45] cikkben publikáltuk. Ez a feltételrendszer továbbra is csak elégséges feltételt ad a kiértékelhetőségre.

4. Definíció. *Definiáljuk a $\overleftarrow{G}^{(3)}$ gráfot a következőképpen. A gráf csúcsai legyenek az objektumok. $A_{i,j,2}^{(3)} > 0$ esetén oda-vissza mutató irányított "egyforma" él legyen behúzva az i és a j csúcs között, míg $A_{i,j,3}^{(3)} > 0$ esetén i csúcsból a j csúcsba mutató egyirányú irányított "jobb" él.*

Ekkor igaz a következő:

3. Tétel. *Legyen egy 3-opciós Thurstone-motivált modellünk $F \in \mathbb{F}$ eloszlásfüggvénnnyel. Tegyük fel, hogy az $A^{(3)}$ mátrixra teljesül az alábbi 3 feltétel:*

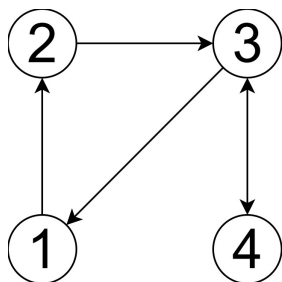
1. *Van legalább egy olyan (i, j) indexpár amelyre $0 < A_{i,j,2}^0$.*
2. *A $\overleftarrow{G}^{(3)}$ gráf erősen összefüggő.*
3. *Létezik olyan irányított kör a $\overleftarrow{G}^{(3)}$ gráfban, ahol "jobb" élek mentén körbeérünk.*
Ekkor az $m_1 = 0$ és $0 < d$ feltételek mellett létezik és egyértelmű a likelihood függvény maximuma.

A [45] cikkben megtalálható a tétel bizonyítása.

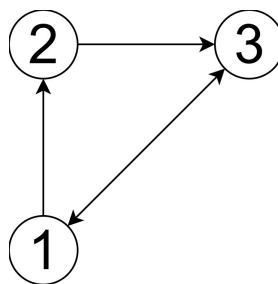
Belátható, hogy a 3. Tételben szereplő feltételrendszer általánosítása a Davidson és a Mihálykó-Orbán feltételrendszereknek. Az első feltétel mindhárom feltételrendszerben megvan. A Davidson-féle feltételrendszerben erősen összefüggőnek kell lennie a gráfnak, de ott ehhez csak a "jobb" éleket használjuk fel. Ebben a feltételrendszerben az "egyforma" kétirányú éleket is, ezért általánosítása a Davidson feltételrendszernek. A harmadik feltétel pedig a Davidson-féle feltételrendszer esetén azonnal teljesül, ha összefüggő a "jobb" élek mentén a gráf. A Mihálykó-Orbán feltételrendszerben a gráfösszefüggőség feltétele egy speciális esete az előbb ismertetett feltételrendszer összefüggőségi feltételének. A Mihálykó-Orbán harmadik feltételének pedig általánosítása a korábban ismertetett csak "jobb" éleket tartalmazó kör, ugyanis az nem csak kettő, hanem több elemből is állhat.

Ez a feltételrendszer jelentős előrelépést jelent a korábbi feltételrendszerekhez képest, de még nem szükséges és elégséges feltételrendszer. Vannak olyan esetek, amikor a feltételrendszer nem képes detektálni, hogy létezik és egyértelmű a maximumhely. Ilyen esetet láthatunk a 4. ábrán.

A 3. ábra egy példa arra, hogy az új feltételrendszer olyan esetekben is képes detektálni, hogy kiértékelhető az adathalmaz, amikor a másik kettő nem. A 3. ábrán látható példa esetén sem a Davidson-féle feltétel, sem a Mihálykó-Orbán feltétel nem teljesül, de a 3. Tétel feltételei igen. Ez azt jelenti, hogy az adatok biztosan kiértékelhetőek. A 3. ábra egy olyan esetet mutat be, amikor kiértékelhető a relációs adathalmaz, de a szigorú körös feltételrendszer sem képes detektálni, hogy kiértékelhetőek az adatok. A 4. ábra egy olyan példa, amely esetén a likelihood függvény felveszi a maximumát és a maximumhely egyértelmű. Ez a 3-opciós Bradley-Terry modell esetén analitikus számításokkal ellenőrizhető. Ennek ellenére egyik eddig ismert feltételrendszer sem teljesül. Vagyis ezek a feltételrendszerek nem szükségesek a kiértékelhetőséghez.



3. ábra. Példa olyan összehasonlítási eredményekre, amikor létezik csupa győzelmekből álló kör



4. ábra. Példa olyan összehasonlítási eredményekre, amikor nem létezik csupa győzelmekből álló kör, az adatok mégis kiértékelhetőek

4.4. Az adatok kiértékelhetőségének szükséges és elégséges feltételrendszere 3-opciós modellek esetén

Ezen fejezetben bemutatom a szükséges és elégséges feltételrendszert az $A^{(3)}$ adathalmaz kiértékelhetőségére a 3-opciós Thurstone-motivált modellek és a Davidson modell esetén. Az eredményt a [39] D1-es cikkben publikáltuk.

A feltételrendszerhez a gráfot a 4. Definícióban megadott módon definiálom.

A kiértékelhetőség szükséges és elégséges feltétele a következő:

4. Tétel. *Tekintsünk egy 3-opciós Thurstone-motivált modellt $F \in \mathbb{F}$ eloszlásfüggvénnyel, vagy a Davidson modellt. Akkor és csak akkor létezik a likelihood függvény egyértelmű maximumhelye az $m_1 = 0$ és $0 < d = -d_1$ megkötéssel a 3-opciós Thurstone-motivált modellben, vagy $0 < \pi, 0 < \nu$, $\sum_{i=1}^n \pi_i = 1$ megkötéssel a Davidson modellben, ha az $A^{(3)}$ mátrix kielégíti az alábbi 3 feltételt:*

1. Van legalább egy olyan (i, j) indexpár amelyre $0 < A_{i,j,2}^0$.
2. A 4. Definícióban megadott $\overleftrightarrow{G}^{(3)}$ gráf erősen összefüggő.
3. Létezik olyan irányított kör a $\overleftrightarrow{G}^{(3)}$ gráfban, amelyben több a "jobb" élek száma, mint az "egyforma" élek száma.

Ezzel a tétellel választ adtunk egy több, mint 70 éve nyitott kérdésre. A matematikailag bizonyított állítást úgy sejtettem meg, hogy igen nagy számú szimulációt végeztem és elemeztem a szimulációk eredményeit.

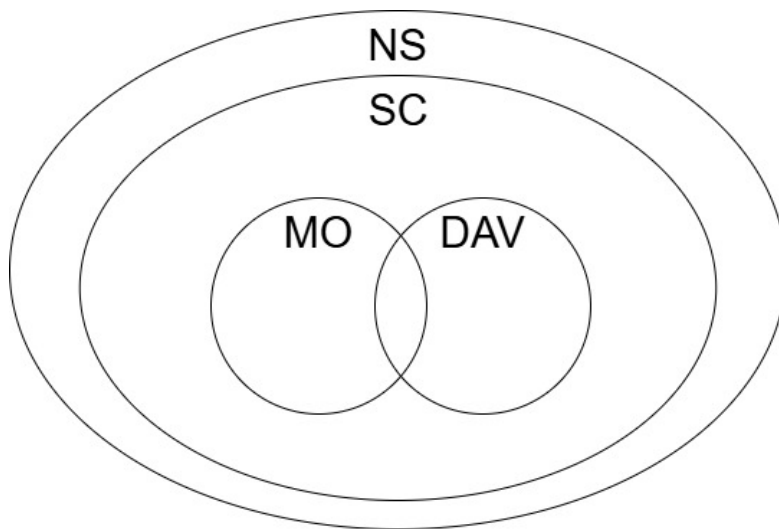
A feltételrendszer jól elkülönülő 3 feltételből áll. Az első feltétel azt biztosítja, hogy létezik a d -nek pozitív alsó korlátja. A második feltétel az egyes objektumok átlagos erősségét, azaz az m_i várható értéket korlátozza alulról és felülről. A harmadik feltétel következtében pedig a d -nek van felső korlátja. Így minden egyes paraméterében korlátos és zárt halmazra van korlátozva a likelihood függvény. Mivel a sűrűségfüggvény log-konkáv ezért a maximum helye egyértelmű. A bizonyítás a [39] publikáció Appendixében található.

Érdeemes megvizsgálni, hogy az egyes feltételrendszerek hatékonysága mennyire tér el egymástól. Ez szimulációk alkalmazásával lehetséges. Bradley-Terry és Thurstone modelleket alkalmaztam a szimulációim során.

4 különféle szituációt generáltam (lásd 6. táblázat), ahol az "egyforma" és "jobb"/"rosszabb" döntések aránya különböző. A döntetlen határát reprezentáló $d = -d_1$ két értéket vehetett fel 0,1-t és 0,5-t. Az erősségeket, az $\underline{m}$ vektort így definiálom: $\underline{m} = (0, h, 2h, \dots, (n-1) \cdot h)$. A h két értéket vehetett fel 0,1-t és 0,25-t. Ezen kiindulási adatok segítségével véletlenszerűen összehasonlításokat generáltam.

Négy feltételrendszerrel értékeltük ki az így generált összehasonlításokat. A [39] cikkből vett jelöléseket használom itt főleg. Az MO a Mihálykó-Orbán feltételrendszer (cikkben MC), DAV a Davidson-féle feltételrendszert, míg az NS a szükséges és elégséges feltételrendszert jelenti, míg az SC a szigorú körös feltételrendszert a [45]-ös cikkből.

Az egyes feltételrendszerek egymáshoz való viszonyát az 5. ábra mutatja be. A feltételrendszereknél az MO és a DAV két jelentősen különböző feltételrendszer, de van közös metszetük. Az SC általánosítása ennek a két korábbi feltételrendszernek, míg az NS az SC általánosítása. Mivel az NS szükséges és elégséges feltételrendszer, így akkor és csak akkor van egyértelmű maximumhelye a likelihood függvénynek, ha az NS feltételrendszer teljesül.



5. ábra. Az egyes feltételrendszerek egymáshoz való viszonya

A hatékonyság illusztrálása érdekében az egyes esetekben 10^8 darab szimulációt végeztünk el és megvizsgáltuk minden egyes szimulációban a különféle összehasonlítási számokra (20, 40, 80) melyik feltételrendszer teljesült.

6. táblázat. Paraméterválasztás a szimulációkhoz

Szituáció	h	d	Az "egyforma" aránya	A "jobb" és "rosszabb" aránya
I.	0,1	0,5	nagy	nagy
II.	0,1	0,1	kicsi	nagy
III.	0,25	0,5	nagy	kicsi
IV.	0,25	0,1	kicsi	kicsi

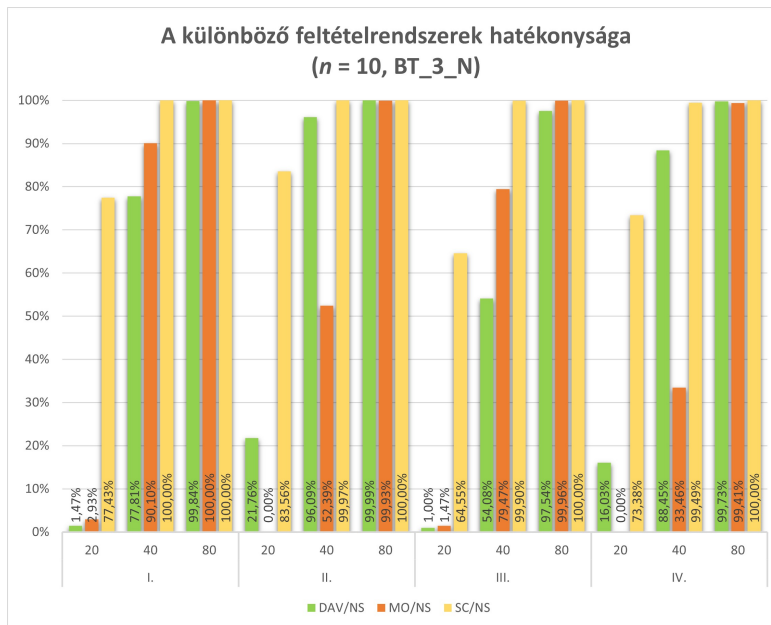
Bradley-Terry modellre 10 objektum esetén az eredmények a 6. ábrán láthatóak. Megfigyelhető, hogy ahogy egyre több összehasonlításunk van, úgy csökken az NS előnye a többihez képest. Az SC 80 összehasonlítás esetén azonos eredményt ér el, mint az NS, 40 összehasonlításnál majdnem 100%-t, de kevés (20) összehasonlítás esetén jelentősen lemaradt tőle. A DAV és MO feltételrendszerek, mivel csak egy részét fedik le a korábban említett NS és SC feltételrendszerek által lefedett eseteknek, jelentősen rosszabbul szerepelnek, különösen kevés összehasonlítás esetén. A legjobb esetben is 20 összehasonlítás esetén nem haladja meg az MO és DAV a 21.76%-ot. Míg az MO jól fel tudja használni a nagyszámú "egyforma" eredményt, addig a DAV mindössze egy "egyforma" döntést tud felhasználni. Ellenben a "jobb"/"rosszabb" együttes előfordulásának nagy aránya inkább a DAV-nak kedvez. Az "egyforma" eredmények felhasználásából következik, hogy a II. és IV. szituációban a DAV, míg az I. és III. szituációban az MO teljesít jobban.

20 objektum esetére is hasonló eredményeket kapunk a Bradley-Terry modellel kiértékelve, amelyek a 7. ábrán láthatóak. Itt is az objektumszám 2-, 4-, 8-szorosa számú összehasonlítást végeztünk el. Az NS előnye jelentősebb mint a 10 objektum esetén, ugyanis a többi feltételrendszer, SC, MO, DAV kevésbé képes detektálni a kiértékelhetőséget.

Thurstone modell esetére is elvégeztük a szimulációkat 10 és 20 objektumra. Az eredményeket a 8. és 9. ábrán láthatjuk. A Bradley-Terry modellben elért eredményekhez hasonló eredmények születtek kis eltérésekkel.

Nagy objektumszám esetén is ($n = 50$) végeztem szimulációkat a Bradley-Terry modellt használva, lásd 10. ábra. Itt az összehasonlítások számát 200, 400, 800-ra kellett növelnem. Az eredmények az NS jelentős előnyét mutatják a többi feltételrendszerrel szemben. Az MO és a DAV a 200, 400 összehasonlítás esetén alkalmazhatatlanok (0% körüli értékűek), még az SC is használhatatlan a szituációk felében 200 összehasonlítás esetén.

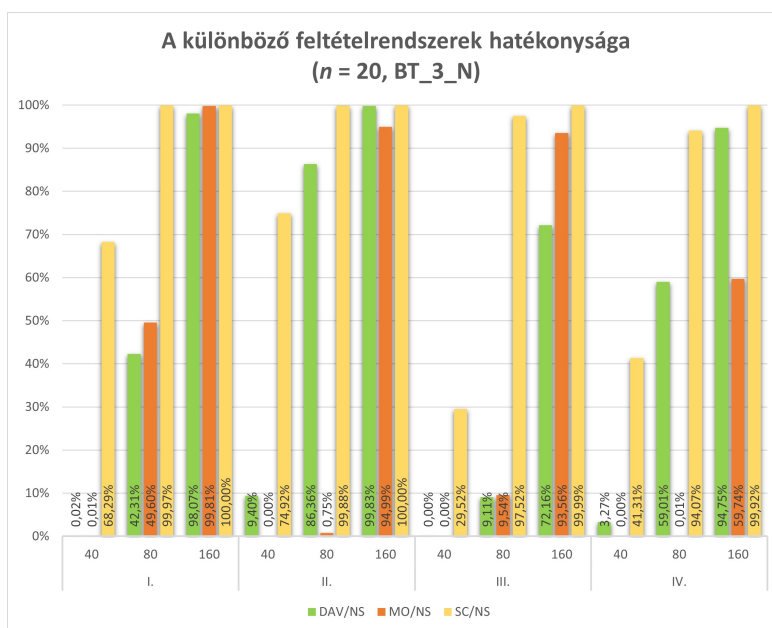
Az eredmények azt mutatják, hogy jelentős az NS előnye az SC-vel, az MO-



6. ábra. Az egyes feltételrendszerek hatékonysága a szükséges és elégséges feltételrendszerhez viszonyítva a Bradley-Terry modellben ($n = 10$)

vel, a DAV-val szemben. Akkor a legnagyobb, ha kevés összehasonlításunk van és sok objektumunk, ekkor sokszor csak az NS segítségével lehet meghatározni, hogy kiértékelhető-e egy adathalmaz vagy sem.

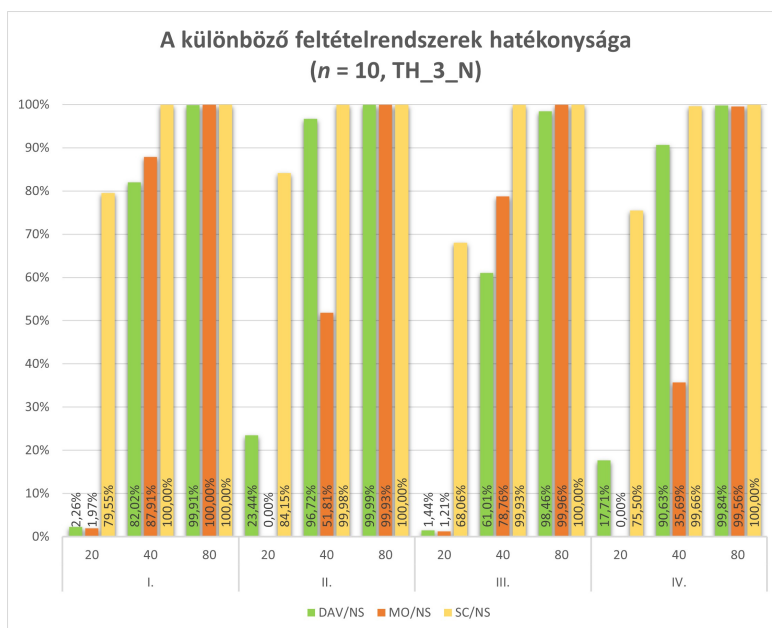
Egy valós életből vett példával is alátámasztom a szükséges és elégséges feltételrendszer fontosságát. A 2024/25-ös Premier League szezonot követtem végig fordulóról fordulóra és vizsgáltam, hogy kiértékelhető-e az adatok az egyes feltételrendszerek segítségével. Meg is tettem a kiértékelést, amennyiben lehetséges. Az angol első osztályban 20 csapat játszik 38 fordulón keresztül. A 11. ábra tartalmazza az első 22 forduló esetén a Spearman rangkorrelációs indexet [48] az utolsó forduló utáni rangsorral összevetve. Zöld színnel jelöltem azon fordulók rangkorrelációs értékeit, amikor egy új feltételrendszer teljesült az adatokra: az 5-dik fordulóban az NS, 6-dikban az SC, 11-dikben a DAV és legvégül a 22-dikben az MO. Az MO legjobb esetben is csak a 20-adik fordulóban teljesülhetett volna, mert akkor volt az első visszavágó, és ezen feltételrendszer esetén szükséges olyan két csapat, amelyiknél az egyik alkalommal az egyik, a másik alkalommal a másik csapat győzött. A korrelációs indexekből leolvasható, hogy már az elején is igen jó eredményt ér el a modell az igen kisszámú (50) összehasonlítás ellenére, a végső eredményhez



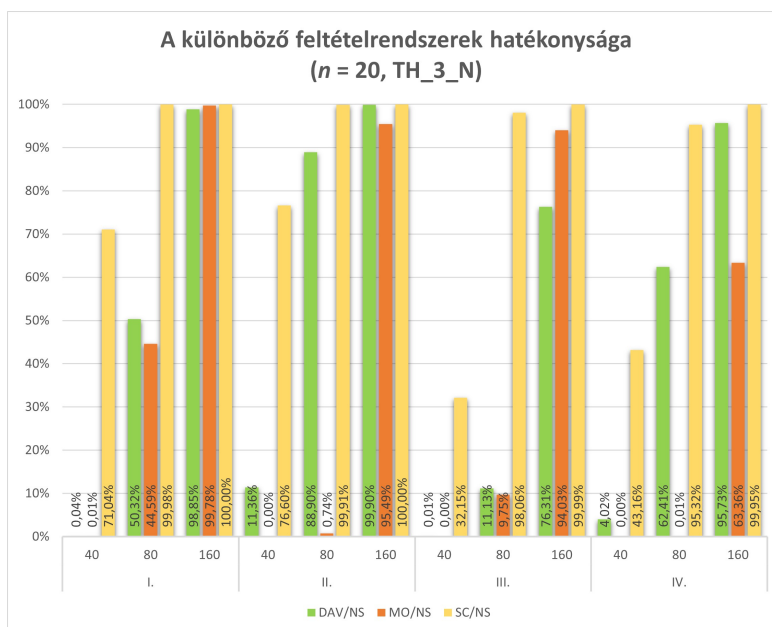
7. ábra. Az egyes feltételrendszerek hatékonysága a szükséges és elégséges feltételrendszerhez viszonyítva a Bradley-Terry modellben ($n = 20$)

viszonyítva. Idővel nő a korrelációs index, de több mint négyszer annyi mérkőzés (22-es forduló, 220 összehasonlítás) után is csak 0,1-el javult a Spearman rangkorrelációs index.

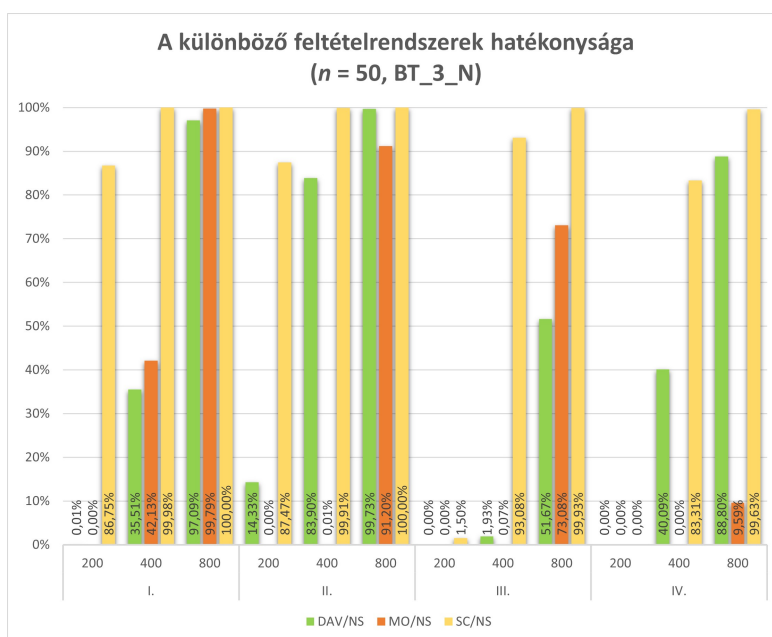
Az eredmények alátámasztják, hogy a gyakorlatban is hasznos tud lenni a szükséges és elégséges feltétel, különösen kisszámú összehasonlítás esetén. Ezen felül ez az információ hozzá tud járulni ahhoz, hogy amennyiben nem lehet kiértékelni az adathalmazt, akkor milyen összehasonlításokat lenne érdemes még elvégezni, hogy azok segítségével kiértékelhetővé váljon az adathalmaz. Az adott információ segítségével meg lehet tervezni összehasonlítási struktúrákat, mely a lehető legkevesebb összehasonlítással a lehető legtöbb információ visszanyerését teszi lehetővé a modell által.



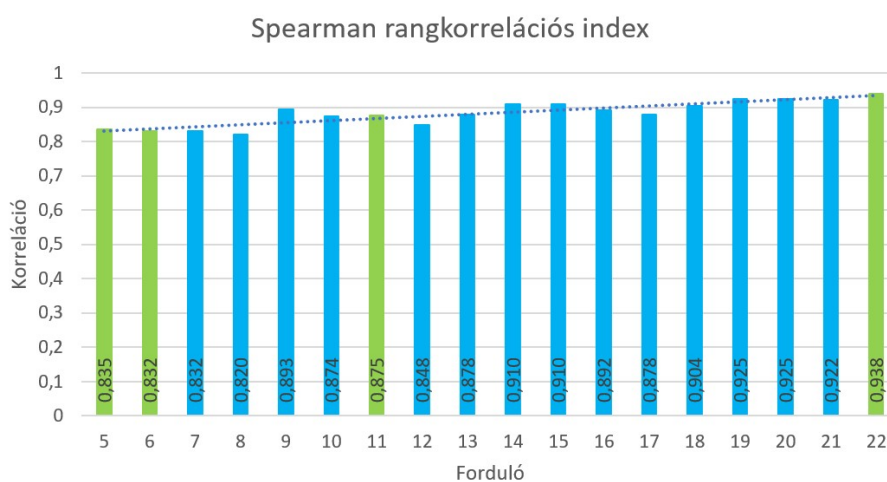
8. ábra. Az egyes feltételrendszerek hatékonysága a szükséges és elégséges feltételrendszerhez viszonyítva a Thurstone modellben ($n = 10$)



9. ábra. Az egyes feltételrendszerek hatékonysága a szükséges és elégséges feltételrendszerhez viszonyítva a Thurstone modellben ($n = 20$)



10. ábra. Az egyes feltételrendszerek hatékonysága a szükséges és elégséges feltételrendszerhez viszonyítva a Bradley-Terry modellben ($n = 50$)



11. ábra. Spearman rangkorreláció az első 22 forduló alapján kialakított rangsor és az összes mérkőzés eredménye alapján kialakított rangsor között. A zöld oszlopok egy-egy új feltételrendszer teljesülését jelzik (5-NS, 6-SC, 11-DAV, 22-MO)

4.5. Optimális opciószám választása

Vannak olyan esetek, amikor az összehasonlítások eredményei különböző opciószám megengedése esetén is értelmezhetőek. Például lehet, hogy csak "jobb"/"egyforma"/"rosszabb" kategóriákba soroljuk az összehasonlítások eredményeit, de az is lehet, hogy a jobb opciót tovább bontjuk "kicsit jobb" és "sokkal jobb" kategóriákra. A megfelelő opciószám választása igen fontos területe a Thurstone-motivált modelleknek. Megfelelő opciószám mellett korábban ki nem értékelhető adatok kiértékelhetővé válhatnak, vagy akár többletinformációt vagyunk képesek visszanyerni, azonos számú összehasonlítás felhasználásával. Az adott témában két publikációm jelent meg [49, 50]. Ezek a 2- illetve 4-opciós modellek, a 3- illetve 5-opciós modellek, valamint az általános s - és $(s+2)$ -opciós modellek kapcsolatával foglalkoznak. Ebben a témában elért eredményeimet foglalom össze ebben az alfejezetben.

Legyen $2 \leq s$ rögzített. A felső indexben jelezzük, hogy hány opciós modellről beszélünk. Álljanak rendelkezésünkre az $(s+2)$ -opciós modellben adatok oly módon, hogy minden $k = 1, \dots, s+2$ esetén létezzen olyan (i, j) pár, melyre $0 < A_{i,j,k}^{(s+2)}$ teljesül. Ezeket az adatokat konvertáljuk át s -opciós adatokká. Az s - és $(s+2)$ -opciós modellek adatai közt a kapcsolat legyen a következő:

$$A_{i,j,1}^{(s)} = A_{i,j,1}^{(s+2)} + A_{i,j,2}^{(s+2)} \quad (37)$$

$$A_{i,j,s}^{(s)} = A_{i,j,s+1}^{(s+2)} + A_{i,j,s+2}^{(s+2)} \quad (38)$$

$$A_{i,j,k}^{(s)} = A_{i,j,k+1}^{(s+2)}, k = 2, 3, \dots, s-1. \quad (39)$$

Fordított esetben, amennyiben az s -opciós modell adataiból $(s+2)$ -opciós modell adatait szeretnénk létrehozni, akkor szétbonthatjuk az s -opciós két szélső döntését, így képezve $(s+2)$ -opciós modellhez adatokat. A szétbontás oly módon végezzük el, hogy $0 < A_{i_1,j_1,s+1}^{(s+2)}$ és $0 < A_{i_2,j_2,s+2}^{(s+2)}$ teljesüljön valamely (i_1, j_1) illetve (i_2, j_2) párok esetén, és a (37) és a (38) összefüggések maradjanak igazak.

Vizsgáljuk meg, milyen kapcsolat állhat fent az s - és $(s+2)$ -opciós modellekben az adatok kiértékelhetősége között. Négy elméleti lehetőséget különböztethetünk meg:

- A) Az adatok csak az s -opciós modellben kiértékelhetők.
- B) Az adatok csak az $(s + 2)$ -opciós modellben kiértékelhetők.
- C) Az adatok mindkét modellben kiértékelhetők.
- D) Az adatok egyik modellben sem kiértékelhetők.

A B) eset létezik, ezt szemlélteti az a példa, amelyhez tartozó adatok a következők:

$$A_{1,2,1}^{(4)} = A_{2,1,4}^{(4)} = A_{1,2,3}^{(4)} = A_{2,1,2}^{(4)} = A_{2,3,2}^{(4)} = A_{3,2,3}^{(4)} = 1, \quad (40)$$

a többi adat pedig 0. Az adatokhoz tartozó irányított gráf a 12. ábrán látható. 4-opciós modellben gyűjtöttük az adatokat, a szemléletesség kedvéért dupla nyíllal szerepelnek a "sokkal jobb", míg szimpla nyíllal a "kicsit jobb" összehasonlítási eredményekhez tartozó élek.



12. ábra. Az (40) példának megfeleltetett irányított gráf a 4-opciós modellben

Ezen példa csak a 4-opciós modellben kiértékelhető. A 4-opciós modellben való kiértékelhetőség alátámasztása érdekében alkalmazhatjuk az [51] publikáció 3.2. Tételét. A 2-opciós modellben viszont az adatok nem kiértékelhetők, mert 2-opciós modellben a $\vec{G}^{(2)}$ irányított gráf nem erősen összefüggő, így nem teljesül a kiértékelhetőség szükséges és elégséges feltétele. (A 3-as objektum erősségét nem tudjuk megbecsülni, ugyanis a likelihood függvény akkor konvergál a szuprémumához, ha ezen objektum erőssége tart a $+\infty$ -hez.)

C) eset létezik, melyet illusztrál a következő 4-opciós modellbeli példa. Az adatok legyenek a következők:

$$\begin{aligned} A_{1,2,1}^{(4)} &= A_{2,1,4}^{(4)} = A_{1,2,3}^{(4)} = A_{2,1,2}^{(4)} = \\ &= A_{2,3,2}^{(4)} = A_{2,3,3}^{(4)} = A_{3,2,2}^{(4)} = A_{3,2,3}^{(4)} = 1, \end{aligned} \quad (41)$$

a többi érték pedig legyen 0. A (41) által generált irányított gráf a 13. ábrán látható. Ezen példa esetén mindkét modell segítségével ki lehet értékelni az

adatokat, ugyanis teljesül a 2-opciós esetben a szükséges és elégséges feltétel, valamint a 4-opciós esetben most is alkalmazhatjuk az [51] publikáció 3.2. Tételét.



13. ábra. A (41) példának megfeleltethető gráf a 4-opciós modellben

D) eset is létezik, példa rá az alábbi adathalmaz:

$$A_{1,2,3}^{(5)} = A_{2,1,3}^{(5)} = A_{1,2,2}^{(5)} = A_{2,1,4}^{(5)} = A_{3,2,1}^{(5)} = A_{2,3,5}^{(5)} = 1, \quad (42)$$

a többi érték pedig 0. A (42) adathalmaz által generált irányított gráfot a 14. ábra mutatja. A szemléletesség kedvéért dupla nyíllal szerepelnek a "sokkal jobb", míg szimpla nyíllal a "kicsit jobb" és oda-vissza nyíllal az "egyforma" összehasonlítási eredményekhez tartozó élek. Nem teljesül a kiértékelhetőség szükséges és elégséges feltétele 3 opció esetén, és 5 opció esetén sem lehet kiértékelni az adatokat, mert ha a 3-as objektum erőssége tart $-\infty$ -be, a log-likelihood függvény értéke szigorúan monoton nő.



14. ábra. A (42) példának megfeleltethető gráf az 5-opciós modellben

Az A) eset azonban nem fordulhat elő, ugyanis igaz az alábbi állítás:

5. Tétel ([49, 50]). *Legyen $s = 2, 3$, és legyen $A^{(s)}$ olyan adathalmaz, amelyben minden opcióba esik legalább egy döntés. Ha az s - és az $(s+2)$ -opciós modellek adatai között teljesül (37), (38) és (39), valamint az $(s+2)$ -opciós modellben is létezik minden opcióba eső döntés, továbbá $A^{(s)}$ kiértékelhető az s -opciós modellben, akkor az $A^{(s+2)}$ kiértékelhető az $(s+2)$ -opciós modellben.*

A 5. Tétel bizonyítása azon alapszik, hogy $s = 2, 3$ esetén a kiértékelhetőségnek ismerjük a szükséges és elégséges feltételeit. Megjegyezzük, hogy ha $3 < s$, akkor mivel nem ismerjük a szükséges és elégséges feltételeket

az általános s -opciós modellben, ezért ilyen állítást matematikailag nem tudunk bizonyítani. Mivel azonban szimulációkat tudtunk végezni, az állítás a nagyszámú szimulációk alapján igaznak tűnik $3 < s$ esetén is.

Összefoglalva az eddigieket, a négy elméleti lehetőség közül csak három valósulhat meg: egyik modellel sem vagyunk képesek kiértékelni; csak a nagyobb opciószámú modellel vagyunk képesek kiértékelni; illetve mindkét modell segítségével elvégezhetjük a kiértékelést. Ezen utolsó eset a legérdekesebb, mert ekkor döntést kell hoznunk, melyik modellel akarjuk kiértékelni az adatokat. A döntést aszerint célszerű meghozni, hogy melyikkel vagyunk képesek több információt visszanyerni.

Erre a kérdésre elméleti eszközökkel nem tudunk válaszolni. Szimulációk segítségével azonban megmutattam, hogy melyik opciószámú modell alkalmazása ajánlott különböző esetekben. Az s és $(s + 2)$ -es modellek közül az $s = 2$ és $s + 2 = 4$, illetve az $s = 3$ és $s + 2 = 5$ opciószámú modelleket választottam ki, mint leggyakrabban alkalmazott modelleket. Véletlen várható értékeket generáltam a 4- és az 5-opciós modell segítségével, ehhez előre fixált d_1 -et és d_1, d_2 -t használtam, vagyis így átlagosan rögzítettem az "egyforma", "kicsit jobb", "sokkal jobb" döntések arányát. Véletlen összehasonlításokat generáltam a Bradley-Terry modellben, majd azokat kiértékeltem mind az s -, mind $(s + 2)$ -opciós modellel ($s=2$, illetve $s = 3$). A szimulációszám 10^5 volt a különféle szituációkban. Egy-egy szituációt mutatok be a következő oldalakon.

Ahhoz, hogy mérni tudjuk a hibát, a következő mérőszámot vezetem be: néztem a kezdeti, és az adatok alapján becsült várható értékek négyzetes eltéréseinek összegét az s -, és $(s + 2)$ -opciós modellekben, azaz az

$$e^{(s*)} = \sum_{i=1}^n (\hat{m}_i^{(s*)} - m_i)^2 \quad (43)$$

mérőszámot. A hiba értéke $e^{(s*)}$, ahol s^* jelenti az opciószámot. Minél kisebb ez az érték, annál jobban illeszkedik a modell. A legkisebb e^{s^*} értéke 0, ami azt jelenti, hogy nincs eltérés az igazi és a becsült várható értékek között. Mivel sok szimulációt hajtottam végre, így átlagot képeztem, amit $\overline{e^{(s*)}}$ -gal jelöltem.

Amikor két különböző opciószámú modellt hasonlítok össze, akkor a relatív hibájuk az érdekes számomra. Például a 2-, és 4-opciós modell összehasonlítása esetén:

$$R^{(2),(4)} = \frac{\overline{e^{(2)}} - \overline{e^{(4)}}}{\overline{e^{(4)}}}. \quad (44)$$

Ez az érték ha negatív, akkor a kisebb opciószámú modell teljesít jobban, míg ha pozitív, akkor a nagyobb, 0 esetén pedig a két modell egyformán teljesít.

A bemutatott példában 2-, és 4-opciós modellek esetén 8 objektumot hasonlítottam össze. Az m_i -ket a $[0,2]$ intervallumból választottam ki egyenletes eloszlás szerint. A d_1 -t $-1,35$ -nek választottam. Az eredményeket az Appendixben található 31. táblázat tartalmazza. Az összehasonlítások számát $|A|$ -val jelölöm, a következő oszlopban azon esetek száma látható, amelyeket mindkettő modell segítségével ki tudtam értékelni, az ezt követően azok száma, amikor csak a 4-opciós modellel, ezután pedig azon esetek számát, amikor egyikkel sem, ezután következnek az egyes modellek esetén a hibák és a relatív hiba. A 21-edik összehasonlítástól számítva a 4-opciós modell teljesít jobban ($R^{(2),(4)} > 0$). Több más szituációban is végeztem teszteket, úgy kerestem d_1 paramétert, hogy a "kicsit jobb", "sokkal jobb" összehasonlítások relatív gyakoriságának aránya minden lehetséges értéket felvehessen egy tizedesre kerekítve. Azt tapasztaltam, hogy a különféle szituációkban, ahol különbözőek a "kicsit jobb", "sokkal jobb" összehasonlítások arányai, idővel mindig a nagyobb opciószámú modell teljesít jobban. Ez azt jelenti, hogy többtinformációt képes a nagyobb opciószámot megengedő modell felhasználni, amennyiben az összehasonlítások száma nem túl kevés. Egyértelmű összefüggést tapasztaltam az objektumok száma és az összehasonlítások száma között, amitől kezdve jobb a nagyobb opciószámú modell. Ha az összehasonlítások száma eléri a legalább $2n + 5$ értéket (n az objektumok száma), akkor érdemesebb a 4-opciós modellt alkalmazni.

A 3- és 5-opciós modellek előszeretettel alkalmazott modellek, legtöbb modell esetén szükséges az "egyforma" megengedett opció, (például labdarúgás esetén döntetlen kimenetele is lehet a mérkőzésnek). Az ilyen összehasonlítások esetén a gólkülönbségekkel különbséget tudunk tenni a "kicsit jobb" és a "sokkal jobb" között. Sokszor elengedhetetlen a nagyobb opciószám alkalmazása, mert a kisebb opciószámú modellel nem lehet kiértékelni az adatokat. Pontosan ez történt a 2024 UEFA Európa Bajnokságban is, ahol Spanyolország minden meccsét megnyerte [52]. Azonban ha 3-opciós modell helyett 5-opciós modellt használunk, az adatok kiértékelhetővé válnak.

$s = 3$ esetén hasonló szimulációkat végeztünk el mint a 2- és 4-opciós modellek esetében. Most is bemutatok egy példát. Az összehasonlított

objektumok száma itt is 8 volt. A várható értékeket egyenletes eloszlással az $[0, 2]$ intervallumból választottuk, d_1 és d_2 rögzített értékek voltak, $d_1 = -1,35$, $d_2 = -0,1$. Ebben az esetben egy extrém helyzetet mutatok be, ahol nagyon ritka az "egyforma" döntés, miközben a "kicsit jobb" és "sokkal jobb" döntések aránya közel azonos. Az "egyforma", "kicsit jobb" és "sokkal jobb" döntések aránya rendre $0,0540 / 0,4768 / 0,4692$. A szimuláció eredményeit az Appendixben található 32. táblázat tartalmazza.

A 32. táblázat szerkezete a 31. táblázat struktúráját követi; az oszlopok elnevezései és jelentése a modellek opciószámai szerint változnak. Itt, hasonlóan az előző esethez, a várható értékek becslésének pontossága jobb az egyszerűbb (3-opciós) modell esetén, ha az összehasonlítások száma kicsi, míg nagyobb összehasonlítás szám esetén a bonyolultabb (5-opciós) modell teljesít jobban. Ahogy egyre több összehasonlítást hajtanak végre az objektumok között véletlenszerűen, a hibák egyre kisebbek lesznek.

Nagyszámú szimuláció alapján megállapítható, hogy 8 objektum kiértékelése esetén legalább 20 összehasonlítás kell ahhoz, hogy az 5-opciós modell átlagosan pontosabban becsülje a várható értékeket. Az 5-opciós modell használata ajánlott a 3-opciós modellel szemben 19 összehasonlítás felett.

Természetesen az összehasonlítások minimális száma, amelytől érdekesebb a nagyobb opciószámú modellt alkalmazni, függ az összehasonlításban résztvevő objektumok számától. Különböző objektumszámok esetén végzett nagyszámú szimuláció alapján arra a következtetésre jutottam, hogy az 5-opciós modell jobb teljesítményt nyújt, mint a 3-opciós modell, ha az összehasonlítások száma legalább $2n + 4$. Ez a határ nagyon hasonló az $s = 2$ esetén tapasztalt határhoz.

Annak érdekében, hogy általánosan ki tudjuk választani, melyik modellel érdemes kiértékelni az adatokat, opciószám-választó algoritmust terveztem. Az opciószám-választó algoritmus folyamatábrája az Appendixben található a 20. ábrán.

Következtetésként elmondhatjuk, hogy ha nagyobb opciószámú modellel végezzük el a kiértékelést, akkor pontosabban kaphatjuk meg az objektumok erősségét. Ez egybevág azzal a heurisztikus gondolattal, hogy több opció megengedése esetén több információ áll rendelkezésünkre, így a látni valószínűségi változók várható értékeit a többi információ alapján pontosabban becsülhetjük. Ez azonban kevés összehasonlítás esetén nem áll fenn. Amennyiben azonban elegendő adat áll rendelkezésünkre, akkor érdekesebb a nagyobb opciószámú modellt alkalmazni kiértékelésre. Maga ez az érték megközelítőleg az objektumszám kétszerese plusz egy kis konstans az $s + 2 = 4$ - és $s + 2 = 5$ -opciós modellek esetében.

5. Az adatok konzisztenciája és inkonzisztenciája a modellekben

A páros összehasonlítási mátrixok esetén a PC mátrix konzisztenciájának vizsgálata központi kérdés. Cikkek garmadája foglalkozik a mátrix (in)konzisztencia mérőszámának definíciójával, a mérőszámok összehasonlításával, az inkonzisztencia csökkentésének módszereivel [53, 54, 55, 56, 57, 58, 59, 60]. A Thurstone-motivált módszerek esetén az adatok konzisztenciájával nem foglalkoznak. A konzisztenciát a 2-opciós Bradley-Terry modellben vizsgáltam, itt tudtam elméleti kapcsolatot teremteni a PCM-en alapuló módszerek konzisztencia fogalma és a sztochasztikus háttérű módszerek esetén az adatok konzisztenciájának fogalma között.

A konzisztencia azt mutatja meg számunkra, hogy az összehasonlítások eredményei mennyire következetesek, van-e bennük egymásnak ellentmondó információ. Relációk esetén vegyünk egy nagyon egyszerű példát, ha egy értékelő véleménye szerint az a objektum "jobb" mint b , és b objektum "jobb" mint c , és a "jobb" mint c akkor a véleménye konzisztens, amennyiben a -nál "jobb" c , abban az esetben ellentmondásba ütközünk ennél a véleménynél.

Azonban sok véleményező lehet. Lehet, hogy az egyik véleményező konzisztensen az egyik irányú, a másik véleményező konzisztensen a másik irányú véleményen van? Mit jelent vajon az A adathalmaz konzisztenciája? Ideális esetben konzisztens az adathalmaz, ha nincsenek benne inkonzisztens összehasonlítások a véleményezők által adott összehasonlítási eredményekben.

Az inkonzisztens összehasonlítások kiszűrése igen fontos feladat, ugyanis konzisztens adatokból jobb információ visszanyerést lehet végrehajtani, azok segítségével pedig jobban lehet döntéseket támogatni. Tehát a konzisztencia vizsgálata nem csak egy lehetőség a modell kiértékelésénél, hanem kiemelt fontosságú feladat, hogy megtudjuk, hogy az általunk kiértékelni kívánt adathalmaz alkalmas-e erősségek meghatározására vagy sem, a jelentős inkonzisztencia következtében.

Az általam vizsgált modellek esetén, ha az adathalmaz konzisztens, akkor ha elhagyok összehasonlításokat az adathalmazból úgy, hogy kiértékelhető maradjon az adathalmaz, akkor ugyanazon erősségeket fogom kapni, mintha nem hagytam volna el összehasonlításokat. Az esetek nagy részében az adathalmazok nem konzisztensek és a teljes összehasonlítás nem áll rendelkezésre. Hogy érdemes ilyen esetben megtervezni az összehasonlításokat, melyeket érdemes elvégezni, hogy a lehető legtöbb információt nyerjük vissza? Lehet-e definiálni általánosan a legjobb összehasonlítási struktúrákat?

Ezekre és hasonló kérdésekre keresem a választ ebben a fejezetben.

5.1. Konzisztencia a 2-opciós Bradley-Terry modellben

A 2-opciós Bradley-Terry modellben az adatok konzisztenciájával kapcsolatos eredményeinket a [61] cikkben mutattuk be.

A PCM-en alapuló módszerekhez hasonlóan akartuk definiálni a konzisztencia fogalmát. A Thurstone-motivált modellek adataiból kiindulva bizonyos PCM-ek esetén a mátrix egyes elemeit $a_{i,j} = \frac{A_{i,j,2}}{A_{i,j,1}}$ módon definiáljuk. Ebben az esetben a konzisztencia definíciója $\prod a_{i,j} = 1$, ahogy korábban a 2.4 fejezetben részletesen ismertettem. Ennek analógjára definiáljuk az adatok konzisztenciáját a 2-opciós Bradley-Terry modellben $h_{i,j} = \frac{A_{i,j,2}}{A_{i,j,1}}$ mennyiségek segítségével a $\prod h_{i,j} = 1$ tulajdonság megkövetelésével.

A 2-opciós Bradley-Terry modell esetén a konzisztencia definíciója a következő: Legyen $K = \{(i, j) : 1 \leq i, j \leq n, i \neq j\}$ az összes összehasonlítási pár, és legyen I olyan részhalmaza K -nak, ahol az I -ben lévő (i, j) pár esetén teljesül, hogy $0 < A_{i,j,1}^2$ és $0 < A_{i,j,2}^2$.

5. Definíció. *Tegyük fel, hogy a $G^{(2)}(I)$, az I által generált irányítatlan gráf összefüggő. Vezessük be a*

$$h_{i,j} = \frac{A_{i,j,2}^{(2)}}{A_{i,j,1}^{(2)}}, (i, j) \in I \quad (45)$$

jelölést. Az $A^{(2)}$ adatmátrix konzisztens I -ben, ha bármely $(i_1, i_2, \dots, i_k, i_1)$ $G^{(2)}(I)$ -beli körben, azaz $(i_l, i_{l+1}) \in I$ és $(i_k, i_1) \in I$ esetén fennáll, hogy

$$h_{i_1, i_2} \cdot h_{i_2, i_3} \cdot \dots \cdot h_{i_k, i_1} = 1. \quad (46)$$

Az A^2 adatmátrix inkonzisztens I -ben, ha a nem konzisztens I -ben.

Az adathalmaz konzisztenciáját úgy definiáljuk, mint nemteljes PCM esetén a konzisztenciát [3]. Konzisztens PCM-ek esetén igaz az is, hogy a kiértékelések eredményeként kapott súlyvektor koordinátáinak hányadosa megegyezik a mátrix megfelelő elemeivel. Ezzel analóg a következő állítás Thurstone-motivált modellek esetén.

6. Tétel ([61]). *Legyen F logisztikus eloszlásfüggvény. Legyen a kiinduló vektorunk $\underline{m}^{(0)} = (0, m_2^{(0)}, \dots, m_n^{(0)})$, tetszőleges rögzített várható értékekkel, ahol az első koordináta nullának rögzített. Definiáljuk a teljes összehasonlításhoz tartozó adathalmazt a kezdővektor segítségével a következőképpen:*

$${}^{(m^{(0)})}A_{i,j,1}^{(2)} = F(m_j^{(0)} - m_i^{(0)}), \quad (47)$$

és

$$^{(m^{(0)})}A_{i,j,2}^{(2)} = F(m_i^{(0)} - m_j^{(0)}). \quad (48)$$

Hagyjuk el az adatok egy részét. Legyen a nemteljes adatmátrix $^{(I)}A^{(2)}$ a következő: $^{(I)}A_{i,j,1}^{(2)} = ^{(m^{(0)})}A_{i,j,1}^{(2)}$ és $^{(I)}A_{i,j,2}^{(2)} = ^{(m^{(0)})}A_{i,j,2}^{(2)}$ ha $(i, j) \in I$, különben 0. Ha az összehasonlításokhoz tartozó irányítatlan $G^{(2)}(I)$ gráf összefüggő, akkor az

$$\ln L(^{(I)}A|\underline{m}) = \sum_{i < j, (i,j) \in I} \left(^{(I)}A_{i,j,1}^{(2)} \cdot \ln p_{i,j,1}^{(2)} + ^{(I)}A_{i,j,2}^{(2)} \cdot \ln p_{i,j,2}^{(2)} \right) \quad (49)$$

likelihood függvény maximumhelye létezik és egyértelmű az $m_1 = 0$ feltétel mellett, valamint a becsléssel kapott paramétervektor megegyezik a kiindulási vektorral, azaz $\hat{\underline{m}} = \underline{m}^{(0)}$.

A 6. Tétel kimondja, hogy ha az $A_{i,j,k}^{(2)}$ -k egy várható érték vektorhoz tartozó pontos valószínűségek, akkor a maximum likelihood becslés az eredeti várható érték vektort adja vissza eredményül. Ha $A^{(2)}$ konzisztens I -ben és $G^{(2)}(I)$ összefüggő gráf, akkor $\underline{m}^{(0)} = (0, m_2^{(0)}, \dots, m_n^{(0)})$ egyértelmű maximumhelye (49)-nek.

A várható érték vektorból kiindulva a pontos valószínűségek kiszámítása a (47) és (48) képlet segítségével a Bradley-Terry modell és $a_{i,j} = p_{i,j,2}/p_{i,j,1}$ arányok segítségével felépített konzisztens, nemteljes PCM mátrix akár az LLSM-mel, akár a sajátvektor módszerrel kiértékelve azonos súlyvektort adnak. Ezekben az esetekben a három módszer eredményei még nemteljes összehasonlítások esetén is megegyeznek.

A 6. tétel alapján bizonyítható az alábbi állítás:

7. Tétel ([61]). Legyen $G^{(2)}(I)$ összefüggő gráf. Az $A^{(2)}$ adathalmaz akkor és csak akkor konzisztens I -n, ha létezik olyan $\underline{m}^{(0)} \in \mathbb{R}^n$, $\underline{m}^{(0)} = (0, m_2^{(0)}, \dots, m_n^{(0)})$ amelyre

$$\frac{A_{i,j,2}^{(2)}}{A_{i,j,1}^{(2)}} = \frac{F(m_i^{(0)} - m_j^{(0)})}{F(m_j^{(0)} - m_i^{(0)})}, (i, j) \in I, \quad (50)$$

ahol F a logisztikus eloszlásfüggvény.

Megjegyezzük, hogy konzisztens adathalmazoknál 2-opciós Bradley-Terry modell esetén

$$\frac{A_{i,j,2}^{(2)}}{A_{i,j,1}^{(2)}} = \frac{\pi_i}{\pi_j} = \frac{p_{i,j,2}^{(2)}}{p_{i,j,1}^{(2)}} = \frac{P(i \text{ "jobb" mint } j)}{P(j \text{ "jobb" mint } i)}. \quad (51)$$

5.2. Optimális összehasonlítási struktúrák 2-opciós Bradley-Terry modellek esetén inkonzisztens esetben

A konzisztencián keresztül már láttuk, hogy kapcsolat van az egyes páros összehasonlítási modellek között. A PCM alapú és 2-opciós Bradley-Terry modellek konzisztenciájának definíciója egymással analóg módon értelmezhető. Az 5.1 fejezet 6. Tételének értelmében konzisztens esetben a teljes és a nemteljes összehasonlításkor kapott prioritási vektorok (súlyvektorok) megegyeznek.

Inkonzisztens esetben van-e kapcsolat a két modell között? Inkonzisztens esetben a 2-opciós Bradley-Terry modell és a PCM alapú kiértékelések más-más eredményt adnak. Szádóczki és Bozóki [62] cikkükben elemezték az optimális összehasonlítási struktúrákat inkonzisztens PC mátrixok esetén és meghatározták, milyen struktúrájú nemteljes összehasonlítások esetén lehet a legközelebb kerülni a teljes összehasonlításokból kapott becslések eredményeihez, azaz mely összehasonlítási struktúrák biztosították a legtöbb információ kinyerését. Mi is elvégeztük ezt az elemzést inkonzisztens adatok esetén 2-opciós Bradley-Terry modell esetén.

A [61] cikkben, 2-opciós Bradley-Terry modell esetén vizsgáltuk információ visszanyerés szempontjából az optimális struktúrákat inkonzisztens esetben.

Monte-Carlo szimulációt alkalmazunk a modellek esetében a hiányos összehasonlításokból történő információ visszanyerés vizsgálatára.

A 2-opciós Bradley-Terry modell esetén a szimuláció lépései a következők voltak:

1. n darab véletlen egész számot generáltunk 1 és 9 között egyenletes eloszlással. Normáltuk ezeket, így megkaptuk a kiindulási súlyvektorunkat ($\underline{\pi}$).
2. Ezután kiszámítottuk a kezdeti erősségvektort az alábbi módon $m_i = \ln(\pi_i) - \sum_{i=1}^n \ln(\pi_i)/n$.
3. Kiszámítottuk az erősségek alapján az összes $p_{i,j,k}^{(2)} = A_{i,j,k}^{(2)}$ értéket. Tudjuk, hogy ezen adathalmaz konzisztens.
4. Konzisztens esetben teljes és nemteljes összehasonlítások esetére is ugyanazt a prioritási vektort kapnánk a maximum likelihood becslés alapján, így perturbáljuk az $A_{i,j,k}^{(2)}$ értékeket, hogy inkonzisztens adathalmazt kapjunk. A perturbáció mértéke egy $(-l; l)$ közötti érték, ahol $0 < l < 1$. Egyenletes eloszlással generálunk perturbációs értékeket ezen $(-l; l)$ intervallumon, ezzel növeljük majd a kezdeti $A_{i,j,2}^{(2)}$ értékét, míg csökkentjük a kezdeti $A_{i,j,1}^{(2)}$ -t. Az így kapott adathalmaz az $\tilde{A}_{i,j,k}^{(2)}$,

amelyre igaz továbbra is, hogy $\sum_{k=1}^2 \tilde{A}_{i,j,k}^{(2)} = 1$ és $0 < \tilde{A}_{i,j,k}^0 < 1$, $\forall i, j, k$.
Ha nem teljesülne a $0 < \tilde{A}_{i,j,k}^{(2)} < 1$, $\forall i, j, k$, akkor új értéket generálunk $(-l; l)$ intervallumon, mindaddig, amíg megfelelünk ennek a feltételnek.

5. Az újonnan kapott inkonzisztens adathalmazt kiértékeljük, elvégezzük az optimalizálást, s így megkapjuk a becsült $\underline{\hat{w}}$ súlyokat és az $\underline{\hat{m}}$ becsült erősségeket.

6. A különféle mérőszámokkal összehasonlítjuk az egyes gráfstruktúrákat.

Teljes összehasonlítás esetén jelöljük az eredményeket, $\underline{\hat{m}}^T$ -vel és $\underline{\hat{w}}^T$ -vel, nemteljes összehasonlítás esetén az egyes esetekben pedig $\underline{\hat{m}}^R$ -rel és $\underline{\hat{w}}^R$ -rel. Munkánkban különféle mérőszámokat alkalmaztunk, amelyek a következők:

1. Várható értékek euklideszi távolsága:

$$EU_M = EU(\underline{\hat{m}}^T, \underline{\hat{m}}^R) = \sqrt{\sum_{i=1}^n (\hat{m}_i^T - \hat{m}_i^R)^2}, \quad (52)$$

2. Súlyok, prioritási vektor euklideszi távolsága:

$$EU_W = EU(\underline{\hat{w}}^T, \underline{\hat{w}}^R) = \sqrt{\sum_{i=1}^n (\hat{w}_i^T - \hat{w}_i^R)^2}, \quad (53)$$

3. Pearson rangkorreláció, várható erősségekre [63]:

$$PE_M = PE(\underline{\hat{m}}^T, \underline{\hat{m}}^R) = \frac{\frac{\sum_{i=1}^n \hat{m}_i^T \cdot \hat{m}_i^R}{n} - \overline{\underline{\hat{m}}}^T \cdot \overline{\underline{\hat{m}}}^R}{\sqrt{\frac{\sum_{i=1}^n (\hat{m}_i^T - \overline{\underline{\hat{m}}}^T)^2}{n}} \sqrt{\frac{\sum_{i=1}^n (\hat{m}_i^R - \overline{\underline{\hat{m}}}^R)^2}{n}}}, \quad (54)$$

ahol $\overline{\underline{\hat{m}}}^T$ a teljes összehasonlítás esetén az erősségek átlagos értéke.

4. Pearson rangkorreláció, prioritási vektorra:

$$PE_W = PE(\underline{\hat{w}}^T, \underline{\hat{w}}^R) = \frac{\frac{\sum_{i=1}^n \hat{w}_i^T \cdot \hat{w}_i^R}{n} - \overline{\underline{\hat{w}}}^T \cdot \overline{\underline{\hat{w}}}^R}{\sqrt{\frac{\sum_{i=1}^n (\hat{w}_i^T - \overline{\underline{\hat{w}}}^T)^2}{n}} \sqrt{\frac{\sum_{i=1}^n (\hat{w}_i^R - \overline{\underline{\hat{w}}}^R)^2}{n}}}. \quad (55)$$

5. Spearman rangkorreláció [48]:

$$\rho = \rho(\underline{\hat{m}}^T, \underline{\hat{m}}^R) = 1 - \frac{\sum_{i=1}^n z_i^2}{6n(n^2 - 1)} \quad (56)$$

ahol z_i a rangsorban mért különbség az i -edik objektum esetén.

6. Kendall τ rangkorreláció [64]:

$$\tau = \tau(\hat{\underline{m}}^T, \hat{\underline{m}}^R) = \frac{2}{n(n-1)} \sum_{i < j} \text{sign}(\hat{m}_i^T - \hat{m}_j^T) \cdot \text{sign}(\hat{m}_i^R - \hat{m}_j^R). \quad (57)$$

Az egyes esetekben 10^6 darab szimulációt futtattunk, hogy megtudjuk melyek az optimális összehasonlítási struktúrák az egyes esetekben. A mérőszámok átlagait számoltuk ki az egyes esetekben, és azt vizsgáltuk, mely páros összehasonlítások során a legkisebbek/legnagyobbak ezen mérőszámok. Euklideszi távolság esetén a kisebb eltérés a jobb, korrelációk esetén a nagyobb érték számít jobbnak információ kinyerésnél.

Különféle 4, 5 és 6 objektumból álló struktúrák segítségével vizsgáltuk teljes és nemteljes esetekben inkonzisztens adathalmazokra melyek az optimális összehasonlítási struktúrák. Feltételezzük, hogy azon összehasonlítási struktúrákat tekinthetjük jobbnak adott objektumszám és élszám esetén, amelyeknek mérőszámai a lehető legjobbak, tehát az erősségek és prioritási vektor a lehető legközelebb van a teljes összehasonlításához tartozó erősségekhez és prioritási vektorhoz.

$n = 4$ esetén 6 egymástól különböző struktúra van, $n = 5$ esetén 21, míg $n = 6$ esetén 112. Különféle perturbációs l -ekre elvégeztük a méréseket kiértékelésekkel.

A futtatások számítási eredményei, a gráfok leírásai és hozzá tartozó vizuális megjelenítés az alábbi helyről tölthetőek le [65].

$n = 5$ esetén az egyes struktúrák és azok mérőszámai az Appendixben található 33. táblázatban láthatóak. Azon struktúrák sorait, melyek adott élszám esetén optimálisak, kivastagítva láthatjuk. Ezek teljesítenek legjobban adott élszám esetén.

4 objektum esetén az egyes gráfokhoz tartozó eredményeket az Appendixben a 21. ábrán láthatjuk. A gráfok sorra g1-g5-tel vannak jelölve. Leolvasható, hogy ahogy a gráfok élszáma nő, úgy egyre jobb eredményeket kapunk. Az EU_M és EU_W esetén a kisebb értékek, míg a többi 4 mérőszám esetén a nagyobb értékek a jobbak. Ugyanezen esetben a szórások pedig az Appendixben a 22. ábrán láthatóak.

6 objektum esetén az egyes élszámok esetén az átlagos távolságokat lásd az Appendixben lévő 23. ábrán, az átlagos korrelációkat pedig Appendixben lévő 24. ábrán. Ahogy egyre nagyobb lesz az élszám, úgy egyre jobban közelíti a nemteljes összehasonlítás eredménye a teljes összehasonlítás eredményét.

Az optimális struktúrák esetén megvizsgáltuk milyen kapcsolat van az euklideszi távolság, a Pearson rangkorreláció index és az élszám között, az eredményeket sorban lásd az Appendixben lévő 25. és 26. ábrán. Itt is a

korábbiakhoz hasonló eredményeket láthatunk, vagyis ahogy az élszám nő, úgy egyre jobb értékeket kapunk.

Egy példát is bemutatok az Appendix 27. ábráján, amikor 6 objektum esetén a 7 élszámú optimális struktúra (g30) és a 8 élszámú nem optimális összehasonlítási struktúra (g49) értékeit hasonlítom össze. Ezen példa is szemléletesen mutatja be, hogy igenis fontos a minél jobb összehasonlítási struktúra alkalmazása információ visszakeresés esetén.

A perturbáció mértékétől jelentősen függnének az egyes mérőszámok értékei, de az optimális struktúrák nem. 4 objektum esetén az Appendix 28. ábráján láthatóak az EU_M és EU_W távolságok az egyes perturbációs mértékek tekintetében a g1-es gráf (csillaggráf) esetében. Az Appendix 29. ábráján láthatóak a korrelációs mérőszámok eredményei. Az Appendix 30. ábráján pedig a szórások alakulása látható. Könnyedén leolvasható, hogy egyre nagyobb perturbációs értéket választva egyre rosszabbak lesznek az értékek adott azonos struktúrák esetén és lineáris kapcsolat van a mérőszámok értéke és a perturbáció mértéke között.

A kiértékeléseket nem csak a 2-opciós Bradley-Terry modellel, hanem a 2-opciós Thurstone modellel is elvégeztük, azt kaptuk, hogy mindegyik esetben megegyeztek az optimális struktúrák. Az optimális struktúrák pedig mind a PCM [62], mind a 2-opciós Bradley-Terry modellben azonosak voltak.

Általánosan azt kaptuk, hogy azon összehasonlítási gráfok, struktúrák a legjobbak, melyek esetén az egyes csúcsokban vett élek száma a legkiegyenlítettebb. Tehát ha az élszám megegyezik az objektumszámmal akkor ez pont a körgráf, míg később a minél egyenletesebb gráfok ezek. Ezen struktúrák alól egy kivétel van, amikor a lehető legkisebb az élszám, vagyis $n - 1$, akkor a csillaggráf teljesít a legjobban, vagyis ha egy objektumhoz van hasonlítva az összes többi.

Minden mérőszám esetén ugyanazon struktúrák voltak a legjobban teljesítők. A Thurstone-motivált modelleknél optimális struktúra optimális a PCM-eknél is. Továbbá a struktúra optimalitása nem függ a perturbáció mértékétől sem.

Összefoglalásképpen elmondható, hogy igen hasznos megtervezni, hogy mely összehasonlításokat végezzük el, annak érdekében, hogy a lehető legtöbb információt nyerhessük ki az adatokból.

Az adatok konzisztencia fogalmával és az optimális összehasonlítási struktúrák vizsgálatával kapcsolatos kutatások többopciós modellek esetén tovább folytathatók, mint tették ezt a szerzők a Davidson modell esetén a [66] publikációban.

6. A modellek alkalmazása

A modellek hasznossága az alkalmazások révén megerősítést nyerhet. A különféle döntéstámogatási helyzetekben előszeretettel alkalmazott modellek a páros összehasonlítási modellek, legyen szó az élet különböző területeiről: oktatás [67], pszichológia [68], biztonság [69], technológia [70], fényforrások [71], biznisz [32] és döntéstámogatás [72]. Különféle sportágakban csapatok, versenyzők rangsorolására is alkalmas, használták autóversenyzésben [73], kézilabdában [44], teniszben [74], sakkbán [75, 76], e-sportban [77] és labdarúgásban [78].

A Thurstone-motivált modellek több lehetőséget kínálnak a PCM alapú modellekkel szemben, mivel valószínűségi háttérrel rendelkeznek, így többek között jövőbeli összehasonlítások esetére valószínűségeket képesek adni. Emellett a modellek kevés adat esetén is kiválóan teljesítenek erősségek, rangsorok meghatározásakor. Mivel egyre nagyobb igény van a különféle döntéstámogatási modellekre és azok alkalmazásaira napjainkban, így ebben a fejezetben a Thurstone-motivált modellek különféle alkalmazásait mutatom be, és a modellek által szolgáltatott eredményeket hasonlítom össze más modellek eredményeivel.

A fejezet 5 alfejezetre oszlik, amelyek segítségével 5 cikk eredményeit ismertetem. Az első alfejezetben a [79] cikk segítségével mutatom be, hogy hogyan lehet nagyobb adathalmazokat kevés összekötő információ segítségével összefűzve egyetlen nagy adathalmazként kiértékelni és az eredményekből következtetéseket levonni. TH_3_N modell segítségével végzem el az összefűzéseket. A második alfejezetben fényforrások csoportosításával, mérésével kapcsolatos kutatás eredményeit vázolom, melyet egy BT_7_N modell segítségével értékelek ki. A harmadik alfejezetben egy többkritériumos saját döntéshozási modellt mutatok be, melynek alapja szintén egy BT_7_N modell. Ezen eredmények megjelentek a [80, 81]-ben. A negyedik alfejezetben előrejelzésre alkalmazom a Thurstone-motivált modelleket. Bemutatom a [44]-es cikk eredményeit. A DELO EHF Női Bajnokok Ligája 2020/2021-es szezonjának eredményeit követem végig és jelzem előre egy TH_3_N modell segítségével. Végül az ötödik alfejezetben a [82] cikk eredményeit ismertetem. A legjobb 5 európai elsőosztályú labdarúgó bajnokság eredményeit jelzem előre és elemzem különféle mutatók segítségével és vetem össze más sikeres modell eredményeivel. Az általam alkalmazott modellek itt a BT_5_N, BT_5_K és BT_5_K_V modellek.

6.1. Csapatok rangsorolása sportmérkőzések eredményei alapján

Sportmérkőzések eredményeit is tekinthetjük páros összehasonlítások eredményeinek. 2 csapat vagy játékos mérkőzése alapján kapunk egy páros összehasonlítási eredményt. A győztest tekintjük a jobbnak az összehasonlításban. Így a sportmérkőzések eredményeit páros összehasonlítási módszerekkel is kiértékelhetjük, mint tették ezt a [74, 75, 76, 78] PC mátrixok esetén. Mi is elvégeztünk ilyen kiértékelést, sőt össze is hasonlítottuk publikációkban a különféle módszerek eredményeit. Munkánkat a [79] cikkben publikáltuk.

Az európai labdarúgás 5 legerősebb elsőosztályú bajnokságának eredményeit értékeltem ki az UEFA Bajnokok Ligája és az UEFA Európa Liga és az UEFA Szuperkupa eredményeit is figyelembe véve. Ezáltal egy olyan összesített rangsor jön létre az 5 ország csapataiból, ahol minden egyes csapat összehasonlítható minden más csapattal annak ellenére, hogy sosem játszottak egymással. A Bajnokok Ligájában a csapatok nemzetiségi megoszlása alapján lett kiválasztva a legjobb 5 nemzeti bajnokság, mely az angol, a francia, a német, az olasz és a spanyol. A 2020/21-es szezonban elért eredményeket mutatom be három modell segítségével, TH_3_N, AL (PCM modell LLSM kiértékeléssel), RN (RankNet) előrecsatolt mély neurális hálózat (lásd [79]), melyet rangsorolásra fejlesztettek ki [83]. Az eredményeket összevetem a Transzfermarket (TM) értékeivel [84], mely a csapatok pénzbeli értékét fejezi ki.

Az angol, francia, olasz, spanyol nemzeti bajnokságban 20-20 csapat van, míg a német bajnokságban 18 csapat. Mivel az egyes bajnokságokban mindegyik csapat játszik mindegyik másik csapattal kétszer, egyet odahaza, egyet idegenben, így rendre 380, illetve 306 mérkőzést játszanak le az egyes bajnokságokban. 1826 mérkőzést játszott le a 98 csapat az 5 nemzeti bajnokságban összesen. A Covid járvány ellenére mindegyik mérkőzést le tudták játszani. A nemzetközi szinten találkoztak egymással a csapatok, ahol az UEFA Bajnokok Ligájában 53 mérkőzés volt, 19 az UEFA Európa Ligában és 1 UEFA Szuperkupa mérkőzés (csak azon mérkőzéseket vettük figyelembe, ahol mindkét csapat a 98 csapat közül került ki), tehát nemzetközi mérkőzésből összesen 73 állt rendelkezésre. Az összes mérkőzésnek ez mindössze 3,84%-a, ami elég kevés adat. A modellek összefésülésre, összefűzésre vonatkozó teljesítménye kritikus lesz ebben az alfejezetben az eredmények szempontjából.

Bemutatom az összefűzött rangsort ebben az alfejezetben, részletezem, hogy melyek a legjobb csapatok a 2020/21-es szezonban az eredmények alapján, továbbá meghatározom a legerősebb nemzeti bajnokságot és következtetéseket vonok le a TM (csapatok pénzbeli) értékeit és a csapatok által elért

eredményeket összehasonlítva.

Mivel több különálló adathalmaz van kevés összekötő információval, így egy saját összefűzési feltételt dolgoztam ki és alkalmaztam ezen esetekben. A rangsorok összefűzését a 2024-ben megjelent [43] cikkben mutattam be.

Legyen két kiértékelhető adathalmazom. Az objektumok halmazai legyen S és $\bar{S}$, S legyen az egyik kiértékelhető adathalmaz csúcsainak halmaza, míg $\bar{S}$ a másiké. Legyenek $i_1, i_2, i_3 \in S$ nem feltétlenül különböző objektumok és $j_1, j_2, j_3 \in \bar{S}$ szintén nem feltétlenül különböző objektumok. Ekkor ha $A_{i_1, j_1, 2} > 0$ teljesül vagy ha teljesül $A_{i_2, j_2, 1} > 0$ és $A_{i_3, j_3, 3} > 0$, akkor a két kiértékelhető adathalmaz az összekötő információ (két adathalmaz objektumai közötti összehasonlítások) segítségével egybefüggő adathalmazként is kiértékelhető, azaz a likelihood becslésnek lesz egyértelmű maximumhelye.

A feltételrendszer segítségével sokszor sokkal gyorsabban eldönthető, hogy ki tudjuk-e értékelni az adatokat, mint az összesített eredmény alapján. Vannak olyan esetek mikor $S \cup \bar{S}$ kiértékelhetőségét sokkal nehezebb eldönteni, mint hogy S -ről és $\bar{S}$ -ről eldönteni, majd alkalmazni ezen feltételrendszert. Ennek következményeképpen sokszor érdemes lehet a kiértékelhetőséget lépésről, lépésre ellenőrizni, és így meggyőződni arról, hogy kiértékelhető-e az adathalmaz.

Most hely hiányában csak az angol bajnokság eredményeit mutatom be, amelyek a 7. táblázatban láthatóak. A bajnokságban a különféle modellek eredményei hasonlóak, a rangsorok csak kicsiben térnek el egymástól. Mindössze három helyen van helycsere: az 5., 6. helyen, a 12., 13. helyen és a 14., 15. helyen. A Hamming-távolság bármely két modell által adott eredmény között négy. A TH_3_N és az AL az 5., 6., 12. és 13. pozícióban tér el. A TH_3_N és az RN az 5., 6., 14. és 15. pozícióban tér el. Míg az AL és az RN a 12., 13., 14., 15. helyen tér el egymástól. Megjegyzendő, hogy a TH_3_N és az AL esetén, amikor ezek a különbségek vannak, akkor ott a csapatok közti súlykülönbség az adott rangsorokon belül körülbelül 0,001.

Megállapítható, hogy a vezető Manchester City (Manch City) csapata jelentősen erősebb a többihez viszonyítva, hasonlóan, a 2. helyezett csapat is jelentősen erősebb az öt követő két csapatnál, míg a 3. és 4. helyen lévő csapatok hasonló erősségűek.

A csapatok által a bajnokságban elért pontokat és pénzbeli, TM értékeiket a 8. táblázat tartalmazza. Itt már nagyobbak az eltérések a két rangsor között. Mindkét lista szerint első a Manchester City csapata, Az első négy csapatot mindkét rangsorban azonos csapatok alkotják. A csapatok pénzbeli értékei nagyjából meghatározzák a csapatok erősségeit, nagy eltérések nincsenek. A legnagyobb eltérés a Leeds esetén van, mely csapat 7 helyet javít a rangsorban a bajnokság szerint a pénzbeli értékéhez képest. Ami még igen szembetűnő, hogy a legdrágább és legolcsóbb csapat között majdnem 8-szoros szorzó van

7. táblázat. A Premier League csapatainak rangsorai és erősségei különböző módszerekkel, a nemzeti bajnokságban játszott mérkőzéseik alapján

helyezés	TH 3 N		AL		RN	
	név	súly	név	súly	név	súly
1	Manch City	0,10637	Manch City	0,08523	Manch City	0,34021
2	Manch Utd	0,07886	Manch Utd	0,07228	Manch Utd	0,14567
3	Liverpool	0,06772	Liverpool	0,06476	Liverpool	0,08614
4	Chelsea	0,06558	Chelsea	0,06301	Chelsea	0,07099
5	Leicester	0,06190	West Ham	0,05964	West Ham	0,06401
6	West Ham	0,06162	Leicester	0,05964	Leicester	0,05341
7	Tottenham	0,05732	Tottenham	0,05645	Tottenham	0,04191
8	Arsenal	0,05521	Arsenal	0,05492	Arsenal	0,03644
9	Everton	0,05320	Everton	0,05343	Everton	0,03279
10	Leeds	0,05314	Leeds	0,05199	Leeds	0,02987
11	Aston Villa	0,04717	Aston Villa	0,04921	Aston Villa	0,02275
12	Newcastle	0,03852	Wolverhampton	0,04173	Newcastle	0,01485
13	Wolverhampton	0,03844	Newcastle	0,04173	Wolverhampton	0,01182
14	Crystal Palace	0,03765	Crystal Palace	0,04060	Brighton	0,01073
15	Brighton	0,03707	Brighton	0,04060	Crystal Palace	0,01048
16	Southampton	0,03607	Southampton	0,03950	Southampton	0,00945
17	Burnley	0,03290	Burnley	0,03739	Burnley	0,00781
18	Fulham	0,02721	Fulham	0,03171	Fulham	0,00467
19	West Brom	0,02541	West Brom	0,03001	West Brom	0,00362
20	Sheffield Utd	0,01865	Sheffield Utd	0,02616	Sheffield Utd	0,00239

8. táblázat. Az egyes csapatok Transzfermarket értékei és az általuk elért pontok a Premier League esetén

helyezés	Transzfermarket		Hivatalos rangsor	
	név	érték	név	pontszám
1	Manch City	€ 1040m	Manch City	86
2	Liverpool	€ 969,65m	Manch Utd	74
3	Chelsea	€ 889,2m	Liverpool	69
4	Manch Utd	€ 770,05m	Chelsea	67
5	Tottenham	€ 703,5m	Leicester	66
6	Arsenal	€ 619,25m	West Ham	65
7	Everton	€ 511,2m	Tottenham	62
8	Leicester	€ 510,7m	Arsenal	61
9	Wolverhampton	€ 450,3m	Leeds	59
10	Aston Villa	€ 430,75m	Everton	59
11	West Ham	€ 334,9m	Aston Villa	55
12	Brighton	€ 290,4m	Newcastle	45
13	Southampton	€ 282,8m	Wolverhampton	45
14	Newcastle	€ 255,35m	Crystal Palace	44
15	Fulham	€ 239,75m	Southampton	43
16	Leeds	€ 238,1m	Brighton	41
17	Crystal Palace	€ 193,85m	Burnley	39
18	Sheffield Utd	€ 149,85m	Fulham	28
19	West Brom	€ 141,15m	West Brom	26
20	Burnley	€ 132,3m	Sheffield Utd	23

pénzbeli értéket tekintve.

A nemzetközi mérkőzések eredményeit figyelembe véve a három modell szerint elkészítettem az összesített rangsorokat illetve a TM rangsort, amelyek egy részlete a 9. táblázatban látható. Ebben az esetben 4-es súlyt alkalmaztam a nemzetközi mérkőzések esetében. Ezen kívül alkalmaztam 1-es súlyozást is, ahol nincs megnövelt súlya a nemzetközi mérkőzéseknek.

Mivel a nemzetközi mérkőzésekre a csapatok jobban felkészülnek, jobban teljesítenek, így ezért is érdemes a nemzetközi mérkőzéseket nagyobb súllyal figyelembe venni. Held és társai is súlyozást használtak a nemzetközi mérkőzések esetén, hasonló megfontolások mentén [38]. A Manchester City csapata első minden rangsor szerint. Jól megfigyelhető a rangsorokban, hogy a Chelsea csapata kiválóan teljesít, ami a nemzetközi eredményeit figyelembe véve nem meglepő, ugyanis mind az UEFA Európa Ligát, mind az UEFA-szuperkupát megnyerte. A legjobb csapatok között van még a

Bayern München, Real Madrid mindegyik rangsor szerint. Ezen rangsoroknál a TH_3_N modell esetén az angol csapatok közül 9, az AL szerint 11, míg az RN szerint 9 csapat van a legjobb 20 között.

9. táblázat. A rangsorok és súlyok a nemzetközi mérkőzések 4-es súllyal történő figyelembevételével

	TH_3_N		AL		RN		TM	
helyezés	név	súly	név	súly	név	súly	név	érték
1	Manchester City	0,03985	Manchester City	0,02575	Manchester City	0,14878	Manchester City	1040,00m
2	Chelsea	0,03276	Chelsea	0,02145	Lille	0,08444	Liverpool	969,65m
3	Real Madrid	0,02816	Manchester Utd	0,01967	Chelsea	0,07968	Chelsea	889,20m
4	Bayern München	0,02791	Liverpool	0,01902	Bayern München	0,06167	Bayern München	858,23m
5	Liverpool	0,02036	Real Madrid	0,01830	Real Madrid	0,05888	Barcelona	817,50m
6	Manchester Utd	0,01998	Bayern München	0,01830	Lyon	0,03717	Paris SG	805,00m
7	Atl. Madrid	0,01736	Leicester	0,01776	Monaco	0,03552	Real Madrid	787,30m
8	Inter	0,01728	West Ham	0,01776	Inter	0,02765	Atl. Madrid	773,10m
9	Leicester	0,01702	Tottenham	0,01682	Manchester Utd	0,02702	Manchester Utd	770,05m
10	Sevilla	0,01603	Atl. Madrid	0,01623	Juventus	0,02581	Tottenham	703,50m
11	Wolfsburg	0,01597	Everton	0,01592	Napoli	0,02260	Inter	663,90m
12	Tottenham	0,01592	Arsenal	0,01566	Liverpool	0,01896	Juventus	633,30m
13	Eintracht Frankfurt	0,01553	Leeds	0,01549	Eintracht Frankfurt	0,01752	Dortmund	628,40m
14	Paris SG	0,01535	Inter	0,01466	West Ham	0,01722	Arsenal	619,25m
15	Barcelona	0,01507	Aston Villa	0,01466	Wolfsburg	0,01659	RB Leipzig	574,95m
16	RB Leipzig	0,01477	Barcelona	0,01446	Bayer Leverkusen	0,01525	AC Milan	547,78m
17	Leeds	0,01473	Sevilla	0,01409	Leicester	0,01447	Napoli	530,20m
18	Lille	0,01469	Lille	0,01384	Sevilla	0,01232	Everton	511,20m
19	Everton	0,01463	RB Leipzig	0,01356	Barcelona	0,01217	Leicester	510,70m
20	Juventus	0,01447	Paris SG	0,01313	Tottenham	0,01114	Wolverhampton	450,30m
...	...	...	...	...	...	...	...	...
89	Angers	0,00478	Udinese	0,00614	Werder Bremen	0,00108	Spezia	71,68m
90	Strasbourg	0,00463	Torino	0,00614	West Brom	0,00107	Dijon	65,63m
91	Lorient	0,00461	Brest	0,00612	Valladolid	0,00105	Huesca	64,10m
92	Nantes	0,00459	Cagliari	0,00581	Benevento	0,00103	Valladolid	63,05m
93	Brest	0,00457	Benevento	0,00550	Eibar	0,00081	Nimes	60,40m
94	Schalke	0,00421	Nimes	0,00548	Sheffield Utd	0,00061	Arminia Bielefeld	56,73m
95	Nimes	0,00383	Schalke	0,00507	Dijon	0,00048	Crotone	52,75m
96	Parma	0,00334	Parma	0,00429	Parma	0,00039	Benevento	50,13m
97	Crotone	0,00318	Crotone	0,00429	Crotone	0,00038	Cadiz CF	47,65m
98	Dijon	0,00262	Dijon	0,00428	Schalke	0,00028	Elche	45,65m

10. táblázat. A rangsorok korrelációi a teljes rangsor esetén a különféle modellekkel és súlyokkal

		Nemzetközi súly=1			Nemzetközi súly=4			TM
		TH_3_N	AL	RN	TH_3_N	AL	RN	
Nemzetközi súly=1	TH_3_N	1	0,763	0,716	0,884	0,772	0,685	0,622
	AL	0,763	1	0,638	0,805	0,987	0,64	0,631
	RN	0,716	0,638	1	0,637	0,628	0,88	0,592
Nemzetközi súly=4	TH_3_N	0,884	0,805	0,637	1	0,814	0,641	0,621
	AL	0,772	0,987	0,628	0,814	1	0,637	0,633
	RN	0,685	0,64	0,88	0,641	0,637	1	0,552
	TM	0,622	0,631	0,592	0,621	0,633	0,552	1

A teljes rangsorok és azok közti rangkorrelációk a 10. táblázatban láthatóak. A különféle modellek közül azok állnak legközelebb egymáshoz amelyekben csak a nemzetközi mérkőzések súlya a különböző. TH_3_N és az AL modellek rangsorai hasonlítanak ezután legjobban egymásra, majd az RN rangsor hasonlít az előző kettőre. A TM tér el a legjobban a többitől, de ez nem meglepő, ugyanis egy csapat teljes értéke nem határozza meg teljesen a csapat erősségét, ugyanis a kispadon lévő játékosok is beleszámítanak a csapat teljes értékébe, illetve a kiemelkedőbb játékosok jelentősen drágábbak, mint amennyivel jobb teljesítményt nyújtanak.

A csapatok országokénti átlagos súlyai és helyezései sorra a 11. és a 12. táblázatban láthatóak. Mindkettő szerint az elsők az angolok a különféle nemzetközi súlyok figyelembevételével, őket követik a németek, majd a spanyolok, az olaszok és legvégül a franciák zárják a sort.

Pénzbeli átlagos értékeket tekintve a 13. táblázat mutatja be a rangsort. Az első és utolsó nemzeti bajnokság esetén egyértelmű a helyzet, míg a német, olasz és spanyol csapatok átlagos pénzbeli értéke nagyon hasonló. Tehát az erősségek és a pénzbeli érték között van összefüggés, ha ez nem is olyan erős az egyes bajnokságokra lebontva átlagosan.

Az angol csapatok esetén egy olyan elemzést is elvégeztem, melyben a csapatok pénzbeli értékét összehasonlítom a súlyukból számított pénzbeli értékükkel. Mivel a teljes bajnokságban lévő csapatok súlya összegezve 1, így a teljes bajnokság Transzfermarket értékét lebontom a súlyok alapján az egyes csapatok között, így megkapom a számított pénzbeli értéküket. Az eredmények a 15. ábrán láthatóak. Az ábrán az egyes csapatok tényleges és számított pénzbeli értékei láthatóak. A ténylegeset az x , a kalkuláltat az y tengely jelenti. Jól megfigyelhető, hogy az egyes csapatok valós és számított pénzbeli értékei egyes csapatok esetén hasonlóak. Tendenciaként leolvasható,

11. táblázat. Az átlagos erősségek országok szerint

név	Nemzetközi súly=1			Nemzetközi súly=4		
	TH_3_N	AL	RN	TH_3_N	AL	RN
angol	0,012	0,015	0,013	0,014	0,015	0,018
francia	0,008	0,008	0,011	0,007	0,008	0,010
német	0,012	0,010	0,008	0,011	0,010	0,008
olasz	0,009	0,008	0,011	0,008	0,008	0,007
spanyol	0,010	0,010	0,009	0,010	0,010	0,007

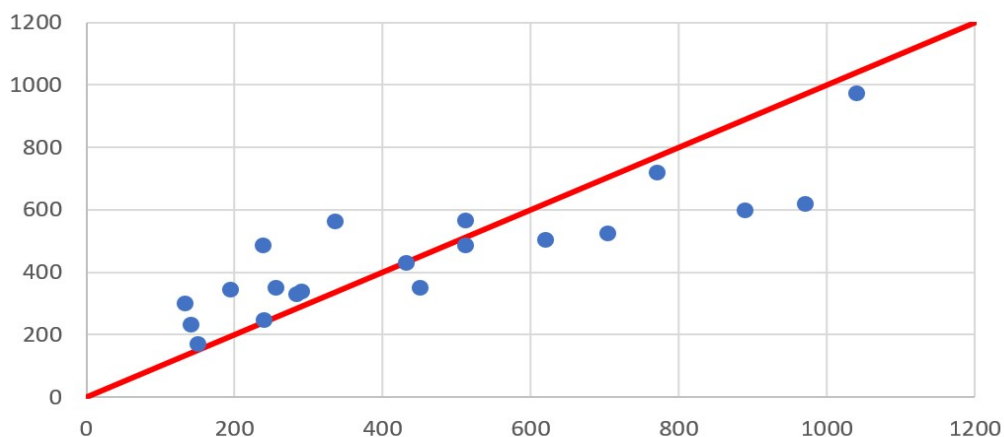
12. táblázat. Átlagos helyezések országoként

név	Nemzetközi súly=1			Nemzetközi súly=4		
	TH_3_N	AL	RN	TH_3_N	AL	RN
angol	39,000	20,800	43,700	31,600	20,850	38,400
francia	67,850	65,200	47,900	71,750	66,100	50,150
német	36,389	46,167	50,778	38,722	46,278	52,556
olasz	55,500	65,800	51,600	58,150	65,000	52,250
spanyol	47,450	49,200	53,650	46,200	48,950	54,450

13. táblázat. Az átlagos pénzbeli értékek a Transzfermarket szerint

név	Transzfermarket
angol	€457,653m
francia	€201,006m
német	€266,406m
olasz	€272,250m
spanyol	€261,403m

A kalkulált és a valós Transzfermarket értékek a
TH_3_N erősségei alapján.



15. ábra. Az angol csapatok becsült és valós Transzfermarket értékei

hogy míg a nagyobb csapatok számított értékei alul maradnak a tényleges pénzbeli értéküknek, addig a kisebb csapatoknál pont ellenkező helyzet áll fent, míg a rangsor közepén lévő csapatoknál kiegyenlített átlagosan a valós és számított pénzbeli érték. Ez a jelenség könnyen magyarázható, egy nagyobb csapatnak nagyobb kerete van, amelyből sok értékes játékos a kispadon van, míg a kisebb csapatok nem engedhetik meg maguknak, hogy sok jó csere álljon rendelkezésükre. Ennek következtében az erősebb csapatoknál több pénz megy el erősségbeli javulásra, mint a gyengébb csapatoknál. Továbbá, ha valaki az első, ha nem is nagy különbséggel, a pénzbeli értéke jelentősen nagyobb lehet, mint a második legjobb játékosé, így ez is megjelenhet ezen a grafikonon.

6.2. Fényforrások csoportosítása erősségük alapján

Ahogy korábban már említettem, nem csak sportágak esetén lehet kiválóan alkalmazni a Thurstone-motivált modelleket, hanem az élet szinte bármilyen területén [68, 69]. Ebben a fejezetben fényforrások elemzéséhez használok a Thurstone-motivált modellt, ahogy tették már előttem mások is [71].

Svájci kollégákkal (Jan Wienold, Chui Ling Yuen, École Polytechnique Fédérale de Lausanne, Laboratory of Integrated Performance in Design), az ő kérésükre kooperációban fényforrásokat elemeztünk. Az adatokat a Lausanne-i Egyetem munkatársai biztosították. A kiértékeléseket, pedig 7-opciós Bradley-Terry modellel tettük meg (BT_7_N).

A kutatásban résztvevőknek, meg kellett határozniuk, hogy két fényforrás közül melyik kellemesebb számukra és mennyivel: "sokkal rosszabb", "rosszabb", "kicsit rosszabb", "egyforma", "kicsit jobb", "jobb", "sokkal jobb". A kellemetlenségét a vakító tulajdonságon keresztül mérték. Összesen 11 (I, II, III, ..., XI) fényforrást hasonlítottak össze a kutatásban résztvevők. 34 résztvevője volt a kutatásnak. Minden egyes résztvevő 87 összehasonlítást végzett el, melyben volt ismétlés a párok között és volt a fényforrások önmagukkal való összehasonlítása is.

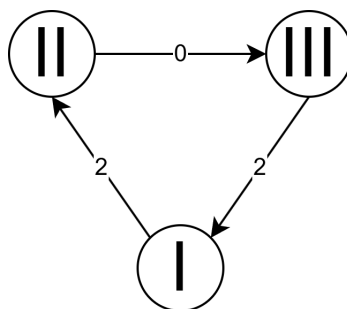
Az így kapott 2958 összehasonlításból értékeltük ki az egyes fényforrások kellemességét/vakítóosságát, melyiket találják a legjobbnak/legkevésbé jónak a résztvevők.

Mivel fontos az értékelések megbízhatósága, elvégeztük ezeket a méréseket. Ehhez kétfajta konzisztencia vizsgálatot végeztünk minden egyes résztvevő esetében.

Az első esetben az olyan összehasonlításokat vizsgáltuk, ahol azonosakat hasonlítottak össze a résztvevők, pl. I-t I-vel. Ezekben az esetekben azt vizsgáltuk, hogy az egyes személyek esetén mekkora volt az eltérés az értékelésben az "egyformá"-tól. Ezen összesített érték meghatározza, milyen megbízható az egyes résztvevő ítélete.

A másik módszer az volt, hogy vettünk minden 3 fényforrásból álló csoportot, pl. I-II-III és ezek között megvizsgáltuk, hogyha körben haladunk az élek mentén, akkor ellentmondásmentesek-e a döntések. Legyen a "sokkal rosszabb" -3, a "rosszabb" -2 és így tovább, a "sokkal jobb" 3. Példánkban (lásd 16. ábra), I "jobb" mint II, II "egyforma" III-vel és III "jobb" mint I ($A_{I,II,6} = A_{II,III,4} = A_{III,I,6} = 1$). A körben vizsgáljuk meg, hogy körbehaladva mennyit kapunk összegként az összehasonlítások mentén. Ellentmondásmentesnek mondhatunk egy ilyen hármast, ha 0-át kapunk eredményül a 3 szám összegeként. Amennyiben az összeg eltér 0-tól abban az esetben van ellentmondás a véleményekben. Kisebb eltérések megengedettek, mivel nagyon nehéz mindig jól megállapítani a fényforrások jóságát, de

nagyobb eltérések már arra utalnak, hogy nem megbízhatóak a véleményező által adott adatok. Példánk esetén jól látható, hogy a körbe haladás mentén az összeg 4, ami jelentős inkonzisztenciára utal. Egyszerű intuícióval is átláthatjuk, ha III "jobb" mint I és I "jobb" mint II, akkor vajon hogyan lehet II és III "egyforma".



16. ábra. Példa inkonzisztenciára 3 objektum esetén

A tesztek sikeresek voltak. Sikerült kiszűrniünk olyan felhasználót, akinek a döntései sok ellentmondást tartalmaztak, így az ellentmondásos véleményező adatait el tudtuk távolítani az adathalmazból.

Az így megtisztított adathalmazt kiértékeljük és azt kaptuk, hogy alig tér el az eredmény a nem megtisztított adathalmazzal történt kiértékelés eredményeitől. Ez arra mutat rá, hogy a modell robusztus, és kisebb hibák nem befolyásolják jelentősen az eredményt, ami igenis fontos, főleg, ha hasonló területeken kell döntéseket hozni, ahol nehezen definiálható a különféle opciók például az "egyforma" vagy a "jobb" közti különbség. A kapott súlyok alapján csoportosítottuk a fényforrásokat és javaslatokat tudtunk adni arra, mely fényforrásokat célszerű alkalmazni, ha el szeretnénk kerülni azt, hogy vakítóként érzékeljék a felhasználók.

A kutatás eredményeiből egy folyóirat publikáció készül, melynek tervezett címe "Absolute Category Rating versus Pairwise Comparison in Discomfort Glare User Experiments".

Ennek az alkalmazásnak a segítségével valós életből vett példán bemutattuk, hogy a modell kiválóan alkalmas objektumok összehasonlítására, rangsorolására, értékelésére, amikor nagyon nehéz az egyes objektumok jóságát egy skála értékeihez viszonyítani.

6.3. Többkritériumos döntéshozatal: atléták besorolása szakágakba

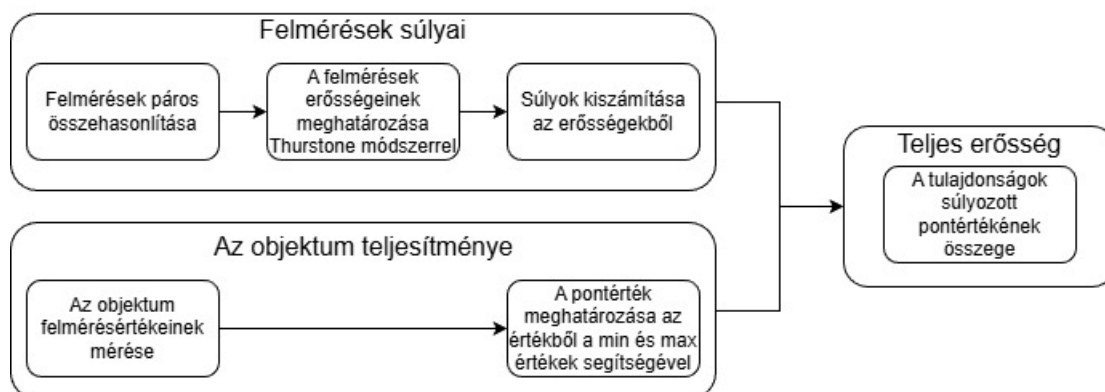
Vannak olyan szituációk, amikor nem lehet egyetlen egy paraméterrel leírni egy objektum erősségét. Például az, hogy egy telefon mennyire gyors, önmagában nem írja le mennyire jó az a készülék. Rengeteg további paraméter ismerete szükséges, hogy jónak nevezhessük, például kamera, akkumulátor, tárhely, minőség, stb. Azt meg lehet mondani, hogy valamelyik tulajdonság alapján melyik jó, melyik nem, de azt hogy összesítve melyik a legjobb készülék, nehéz meghatározni. Ezen problémákra előszeretettel alkalmazzák a többkritériumos döntéstámogatási modelleket [85, 86, 87]. Az ilyen és hasonló esetekre fejlesztettem ki a többkritériumos döntéstámogatási modellt, melynek alapja egy BT_7_N modell. Az eredmények a [80, 81] cikkekben jelentek meg.

Valós életből vett példán keresztül mutatom be ezen többkritériumos döntéstámogatási modellt. Sportolók esetén döntést kell hoznunk, hogy három szakág közül melyiket érdemes választaniuk. Hol képesek a legjobban teljesíteni? A három szakág: a dobás, a futás és a futás/ugrás. A futás hosszútávfutást, a futás/ugrás esetén a futás rövidtávfutást takar. Ezek a szakágak különböző típusú képességeket igényelnek: A hosszútávfutás lassú folyamatos terhelést, a dobás, az ugrás, és a rövidtávfutás pedig hirtelen erő kifejtést igényel. Szeretnénk azt is tudni mindössze egyetlen számba (értékbe) sűrítve, hogy mennyire jó az egyes sportoló az adott szakághoz tartozó képesség területeken a korosztályához, neméhez mérten. Ugyanis különféle korosztályok, nemek különféle teljesítményre képesek az egyes atlétikai versenyeken. Használjuk számításaink kiinduló pontjának a NETFIT felméréseit [88], melyek általánosan alkalmazott felmérések az iskolákban Magyarországon. A felmérések elnevezései a 18. ábrán láthatók. Jelen 6.3 fejezetben a saját modellünk által képzett rangsort és erősségeket NETFIT rangsornak és NETFIT erősségeknek nevezzük a későbbiekben.

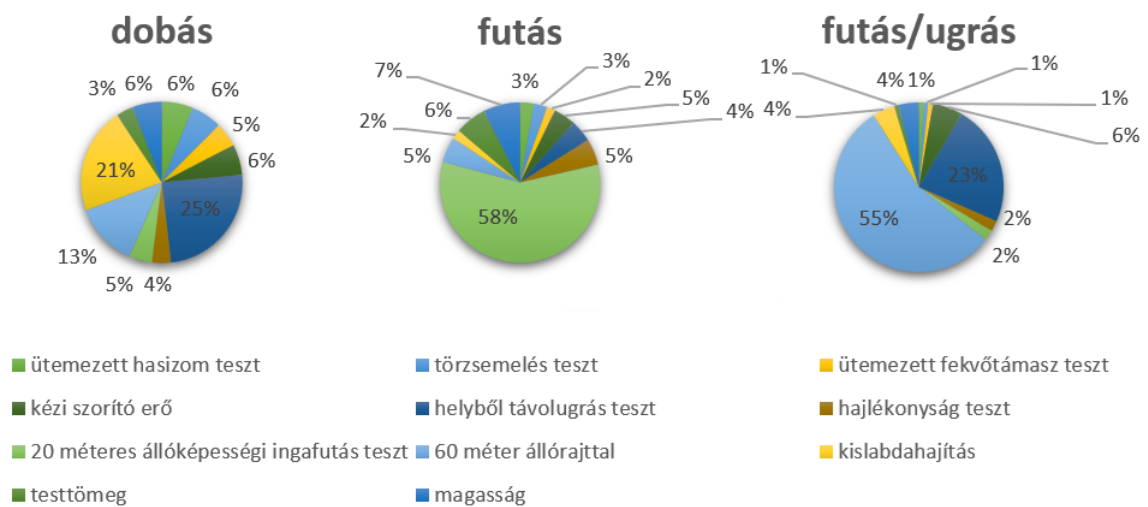
A modell két fő részből áll, először is meg kell határoznunk az egyes részképességek fontosságát, majd ezután meg kell határoznunk azt, hogy az egyes felmérésekben mennyire jó a sportolónk. A számítási módszer a 17. ábrán látható. Először meghatározzuk az egyes felmérések súlyát, mennyire is fontos számunkra az, hogy abban a sportoló jól teljesítsen. Például a futás szakág esetén valószínűsíthető, hogy a futással kapcsolatos felmérések lesznek fajsúlyosak, míg a dobás szakágnál a dobással kapcsolatosak. A felmérések súlyát így határozzuk meg: először páronként összehasonlítjuk szakáganként az egyes felméréseket, vajon melyik fontosabb: a kislabdahajítás vagy a 60 méter futás állórajttal? Ezen eredményeket Thurstone-motivált modellel kiértékeljük és súlyokat képeztünk. A fontosság tekintetében a véleményezést atlétaedzők végezték el 7-opciós modellben páronként összehasonlítva a részképességeket.

Tehát megkaptuk az egyes szakágakban a felmérések súlyát, jelentőségét. Az eredményeket lásd a 18. ábrán.

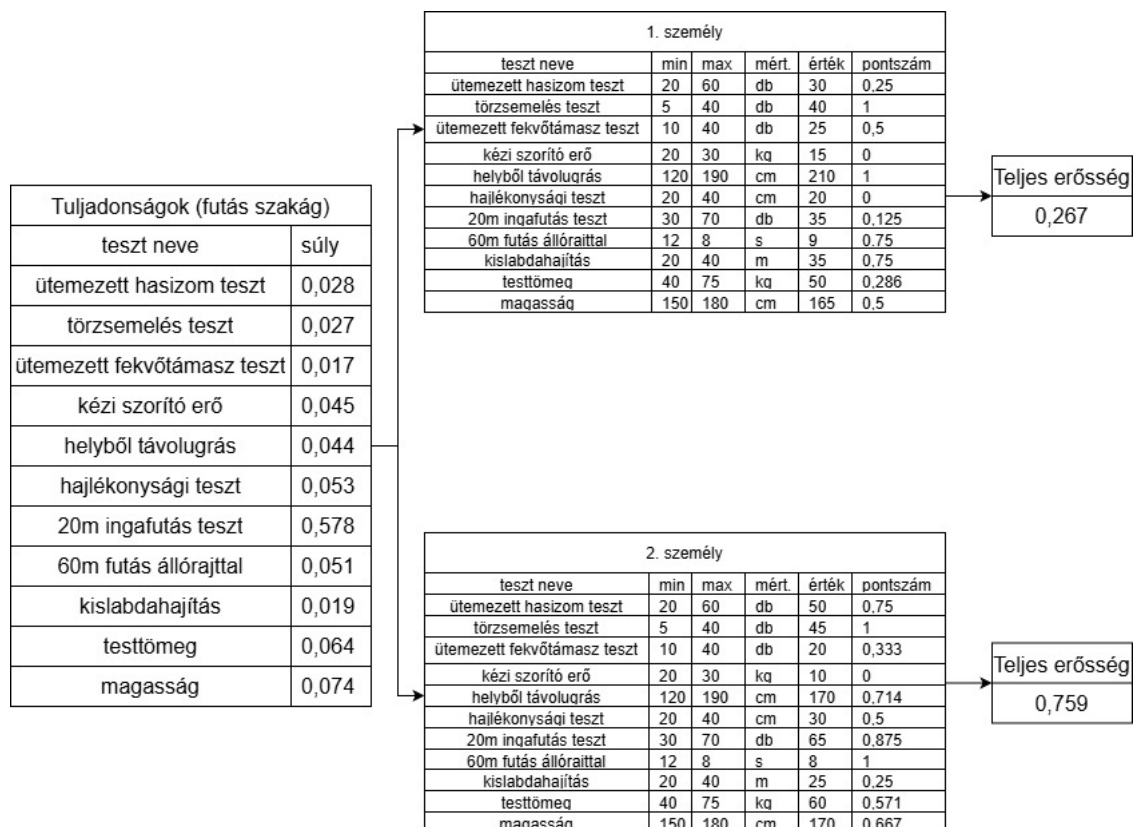
Minden egyes korosztályban van egy skálánk, aminek *min* és *max* értékei vannak az egyes felmérésekhez. Ezeket a *min* és *max* értékeket az edzők adják meg a korosztályra és nemre való tekintettel. Például 60m futás esetén 7 másodperchez 1-es értéket rendelünk (ez a max) és 17 másodperchez 0-ás értéket rendelünk (ez a min). A sportoló eredménye legyen *value*. Ha $min < value < max$, akkor a pontérték = $(value - min) / (max - min)$. Ha $value \leq min$ akkor pontérték = 0 és ha $max \leq value$ akkor pedig 1. A korábbi példára visszatérve a 12 másodperces értékhez $(12-17)/(7-17)=0,5$ -ös értéket rendelünk. Így megkapjuk az egyes pontértékeket az egyes felmérések esetén. Ezt követően vesszük a pontértékeknek a súlyozott átlagát. Az egyes felmérések esetén a súlyok a korábban kiszámított felmérések súlya. Egy teljes számításra példa a 19. ábrán látható. A végső érték így 0 és 1 közötti lesz; ha 0, akkor mindenből minimális a teljesítménye a sportolónak, míg 1-es érték esetén mindenből maximálisan teljesít.



17. ábra. A módszer számítási folyamata



18. ábra. A felmérések súlyai a szakágakban



19. ábra. Példa a teljes erősség kiszámítására

A 18. ábrán kitűnik az eredményekből, hogy az egyes szakágaknál az a fontos felmérés, ami a szakág jellemzője: futásnál a 20 méteres állóképességi ingafutás teszt, futás/ugrásnál a helyből távolugrás teszt és a 60 méter állórajttal, míg dobásnál a kislabdahajítás és a helyből távolugrás teszt. Ez utolsó igen meglepő eredménynek tűnik elsőre, de az edzőket megkérdezve, akik az összehasonlításokat elkészítették, amelyekből a súlyok lettek kiszámolva, ez nem meglepő. A dobás esetén nagyon fontos a hirtelen erő kifejtés, melyet a helyből távolugrás nagyon jól szemléltet, ugyanis súlylökésnél a sportoló az alsó végtagjait is használja jelentősen a súly eldobásához. A többi felmérés kis súllyal számít bele az összegzésbe.

Különbőféle alkalmazások segítségével akartuk a modell hasznosságát és használhatóságát tesztelni. Első ilyen alkalmazás volt, hogy szakági besorolásra váró sportolók eredményeit kiértékeljük és az edzői véleményeket is bekértük. Ezek alapján elkészítettük az ajánlásokat, amelyek esetén azt a szakágat javasoltuk, amiben a legjobb a sportoló. Az eredmények a 14. táblázatban láthatóak. Azt tapasztaltuk, hogy az edzői ajánlások és a modell által ajánlott szakág 89%-ban megegyezett, 18 esetből 16-szor sikerült a modellünk segítségével azonos besorolást adnunk, mint az edzők adnának a fiatal sportolóknak.

Másik módon is mérni akartuk a modell helyességét. Szakágakba besorolt sportolók eredményeit vettük figyelembe és azokat értékeltük ki, vizsgálva, hogy a NETFIT értékek, a rangsor és a tényleges eredmények között van-e összefüggés. Fiúknál és lányoknál is futás/ugrás szakág esetén vizsgáltuk ezt. Az eredmények a 15. táblázatban láthatóak. Azt kaptuk, hogy mindegyik sportoló a saját szakágában teljesített a legjobban, és a pontszáma segítségével létrehozott NETFIT rangsor megegyezik 6-ból 4 esetben, a lányoknál az utolsó kettő helyet cserélt a két rangsor szerint.

Több, szakágakba már besorolt sportoló esetén is elvégeztük a kiértékeléseket, ajánlásokat, lásd 16. táblázat. 13-ból 11-szer helyesen adta meg a modell a sportoló szakágát, a maradék két esetből egyik esetén sérülés állt fent (fiú IV), míg a másik sportoló (fiú III) több szakágban is kimagaslik a korcsoportjában.

A tesztek segítségével nem csak szakágakba tudjuk sorolni az egyes sportolókat, hanem azok előremenetelét nyomon tudjuk követni számszerűsítve. Lányok és fiúk esetén is 3-3 sportolónál elvégeztük az elemzéseket, melyek a 17. és 18. táblázatban láthatóak. A táblázatok esetén 14-15 éves sportolók adatait használtuk fel. Ebben a korban a lányok már nem sokat fejlődnek fizikailag, míg a fiúk nagy fejlődésen, erősödésen mennek át. Mindegyik sportoló szakága a futás/ugrás szakág. Mindenkinek, fiú I kivételével saját szakágában volt a legmagasabb az értéke a szeptemberi felméréskor. Minden sportolónak sikerült javítania korábbi helyből távolugrás teszt értékein, mely kulcsfontosságú ezen

14. táblázat. Szakágakba besorolandók eredményei és ajánlások a megfelelő szakágakra

születési év	név	dobás	futás	futás/ugrás	edzői ajánlás	módszer ajánlása
2011	atléta 1	0,506	0,483	0,637	futás/ugrás	futás/ugrás
2012	atléta 2	0,280	0,321	0,427	futás	futás/ugrás
2012	atléta 3	0,409	0,632	0,508	futás	futás
2012	atléta 4	0,447	0,529	0,605	futás/ugrás	futás/ugrás
2012	atléta 5	0,116	0,610	0,047	futás	futás
2013	atléta 6	0,423	0,592	0,576	futás	futás
2014	atléta 7	0,558	0,825	0,612	futás	futás
2014	atléta 8	0,515	0,794	0,568	futás	futás
2014	atléta 9	0,628	0,317	0,641	futás/ugrás	futás/ugrás
2014	atléta 10	0,258	0,232	0,259	futás/ugrás	futás/ugrás
2014	atléta 11	0,293	0,221	0,210	dobás	dobás
2015	atléta 12	0,362	0,191	0,492	futás/ugrás	futás/ugrás
2015	atléta 13	0,502	0,239	0,591	futás/ugrás	futás/ugrás
2015	atléta 14	0,124	0,097	0,220	futás/ugrás	futás/ugrás
2015	atléta 15	0,049	0,032	0,146	futás/ugrás	futás/ugrás
2016	atléta 16	0,097	0,096	0,032	dobás	dobás
2017	atléta 17	0,068	0,058	0,054	futás	dobás
2017	atléta 18	0,018	0,025	0,009	futás	futás

15. táblázat. A fiú és lány versenyzők eredményei a két rangsor szerint

név	dobás	futás	futás/ugrás	hivatalos rangsor	NETFIT rangsor
fiú 1	0,609	0,661	0,819	5,67m	1
fiú 2	0,671	0,570	0,683	5,05m	2
lány 1	0,634	0,495	0,754	13,57s	1
lány 2	0,562	0,624	0,666	13,68s	2
lány 3	0,452	0,536	0,572	14,33s	3
lány 4	0,439	0,432	0,550	13,90s	4

szakág esetén. Ami kimondottan érdekes, az az hogy a fiúknál a hajlékonyság területén (hajlékonyság teszt) jelentős fejlődés történt, ami a kiváló edzői munka eredménye, ugyanis ez könnyen csökken az idő előrehaladtával és minél

16. táblázat. Szakágakban lévő atléták eredményei

születési év	név	dobás	futás	futás/ugrás	edzői ajánlás	módszer ajánlása
2011	lány I	0,806	0,899	0,943	futás/ugrás	futás/ugrás
2011	lány II	0,917	0,921	0,936	futás/ugrás	futás/ugrás
2010	lány III	0,939	0,861	0,983	futás/ugrás	futás/ugrás
2009	lány IV	0,887	0,690	0,970	futás/ugrás	futás/ugrás
2008	lány V	0,826	0,724	0,965	futás/ugrás	futás/ugrás
2008	lány VI	0,691	0,587	0,934	futás/ugrás	futás/ugrás
2007	lány VII	0,693	0,692	0,871	futás/ugrás	futás/ugrás
2010	fiú I	0,827	0,895	0,960	futás/ugrás	futás/ugrás
2009	fiú II	0,901	0,937	0,980	futás/ugrás	futás/ugrás
2009	fiú III	0,876	0,956	0,899	futás/ugrás	futás
2008	fiú IV	0,753	0,885	0,874	futás/ugrás	futás
2008	fiú V	0,583	0,627	0,710	futás/ugrás	futás/ugrás
2008	fiú VI	0,932	0,963	0,966	futás/ugrás	futás/ugrás

rugalmatlanabb a sportoló, annál nagyobb a kockázata a sérülésnek.

17. táblázat. A régi és az új teszt eredményei a lányoknál

név	lány I	lány I	lány II	lány II	lány III	lány III
dátum	01.25	09.26	01.25	09.26	01.25	09.26
ütemezett hasizom teszt	44	56	44	61	45	80
törzsemelés teszt	17	20	9	13	16	19
ütemezett fekvőtámasz teszt	18	12	21	16	16	11
kézi szorítóerő	25	31,6	24	33,5	22	33,5
helyből távol ugrás	200	210	200	206	210	220
hajlékonysági teszt	28	25	25	25	33	30,5
20 méteres állóképességi inga futás	53	36	60	39	63	46
60 méter állórajttal	8,7	8,1	8,6	8,4	8,3	8,1
kislabdahajítás	25	20	28	26	30	27
testtömeg	55	53	52	51	57	55
magasság	164	170	171	173	165	172
dobás	0,717	0,691	0,774	0,826	0,819	0,881
futás	0,860	0,587	0,893	0,724	0,851	0,942
futás/ugrás	0,810	0,934	0,881	0,965	0,924	0,976

18. táblázat. A régi és új teszt eredményei a fiúknál

név	fiú I	fiú I	fiú II	fiú II	fiú III	fiú III
dátum	01.25	09.26	01.25	09.26	01.25	09.26
ütemezett hasizom teszt	70	80	80	80	80	80
törzsemelés teszt	35	25	17	16	19	16
ütemezett fekvőtámasz teszt	29	35	16	26	16	30
kézi szorítóerő teszt	30	42	32	38,8	30	49
helyből távolugrás teszt	230	250	200	220	230	270
hajlékonyság teszt	30	37	31	39	30	34,5
20 méteres állóképességi ingafutás teszt	90	80	60	52	80	90
60 méter állórajttal	8	8,1	8,5	8,1	7,1	7,1
kislabdahajítás	35	43	39	40	40	48
testtömeg	61	63	52	58	75	75
magasság	175	176	165	166	173	180
dobás	0,634	0,876	0,373	0,583	0,606	0,901
futás	0,873	0,956	0,771	0,627	0,816	0,937
futás/ugrás	0,797	0,899	0,367	0,710	0,790	0,980

A modellről különféle alkalmazások során bemutattam, hogy alkalmas többkritériumos döntéstámogatásra nagyon kevés információ segítségével is. Ezen döntéstámogatási modell az élet más területein is könnyen alkalmazható, ahol több kritérium együttesen adja meg az objektumnak a teljes erősségét.

Általánosan alkalmazva a modellt, a következőképpen érdemes eljárunk. Meg kell határoznunk azon tulajdonságokat, melyek fontosak számunkra a komplex döntéstámogatás folyamán. Ezután ezen tulajdonságokat kell párban összehasonlítani (jobb, rosszabb, ...), az így kapott páros összehasonlítások eredményeit pedig kiértékelni. A Thurstone-motivált modellek segítségével megkapott eredményekből súlyt képzünk. Ezen súlyok határozzák meg, hogy az egyes tulajdonságok milyen mértékben számítanak bele a döntésbe. Ezt követően meg kell határoznunk, minden egyes tulajdonsághoz egy minimum és maximum értéket, vagyis skálát. A skála segítségével pedig 0-1 közötti értékre normálható az objektum minden egyes tulajdonságának a pontértéke. Legvégül pedig vesszük a tulajdonságokhoz tartozó pontértékeknek a súlyozott számtani átlagát. Megjegyezzük, hogy az egyes tulajdonságok alapján meghatározott pontértékeket akár páros összehasonlítás segítségével is kiszámolhatjuk.

Ezen módszerrel így könnyen tudunk kevés információ segítségével összehasonlítani például telefonokat, tendereket.

6.4. Sportesemények előrejelzése

A különböző páros összehasonlítási modellek különböző lehetőségeket rejtenek magukban. Míg a PC mátrixon alapuló modellek nem, addig a sztochasztikus modellek lehetőséget biztosítanak arra, hogy jövőbeli események kimenetelére valószínűségeket adjunk meg. Modellek segítségével előre lehet jelezni sportesemények kimeneteleit, becsülni lehet a kimenetek valószínűségeit.

Ebben a fejezetben két cikk eredményeit részletezem [43, 44]. A DELO EHF Női Bajnokok Ligája 2020/2021-es eredményeit elemzem és jelzem előre és vetem össze a végleges eredményekkel.

A DELO EHF Női Bajnokok Ligája 2020/2021-es eredményeit elemző és előrejelző cikkeim a CJOR-ban és Knowledge-ben jelentek meg. Az elemzésnél különös jelentőséget adtak a COVID'2019 miatti lezárások, a tervezett mérkőzések egy része nem lett lejátszva.

A Bajnokok Ligájában a csoportkörben 2 csoport 8-8 tagja játszott egymással oda és vissza mérkőzést. Ezután a csoporton belüli rangsorok alapján egyenes ági kieséses szakasz következett. Végül a versenyt a Final4 zárta, melyben eldőlt az első 4 csapat helyezése.

Ebben a szezonban jó néhány meccs el lett halasztva, vagy el is maradt a Covid miatt. Sok mérkőzés így nem lett lejátszva, és a vétlen félnek ítelték meg a pontokat, ami torzította a csoportokban a rangsort. Az A csoport eredményei a 19. táblázatban láthatóak. Kétféle rangsort készítettem el, egyik a $TH_3_N^1$ alapján, ezen esetben minden meccs eredményét figyelembe vettük, az olyan eredményeket is, amikor a meccs elmaradt, de valakinek odaítélték a meccsért járó pontot. A $TH_3_N^2$ pedig csak a ténylegesen lejátszott mérkőzések alapján létrejött rangsort mutatja. Itt azért vannak nem egész értékű pontszámok, mert a le nem játszott mérkőzések esetén a pontszámot a két csapat között elosztottuk az alapján, hogy milyen valószínűséggel szerzi meg a pontot a csapat.

Az A csoport esetén csak az FTC-Rail Cargo Hungaria csapata játszott le minden mérkőzést, a többi csapatnak voltak le nem játszott mérkőzéseik, melyek eredményét és az érte járó pontokat az EHF határozta meg. Az odaítélt pontok nem csak az erősségeket, de a végleges rangsort is jelentősen befolyásolják az A csoportban. Két csapat helyet cserélt az így odaítélt pontoknak köszönhetően. A nagy győztese az odaítélt pontoknak a CSM Bucuresti, mely így a 3. helyen végzett. Feltehetőleg, ha lejátszotta volna a mérkőzéseit akkor csak az 5. helyen végez a csoportban. A nagy vesztese pedig a Vipers Kristiansand csapata volt, mely 5. helyen végzett az oda (nem) ítélt pontoknak köszönhetően, bár ha lejátszotta volna a mérkőzéseit, a 3. helyen végezhetett volna a többi eredménye alapján.

A B csoport eredményeit a 20. táblázat mutatja be. Ebben a csoportban is

19. táblázat. Az A csoport eredményei

Csapatok hivatalos rangsora	pont	$TH_3_N^1$		$TH_3_N^2$		
		helyezés	$\hat{m}_i$	helyezés	$\hat{m}_i$	várható pont
1. Rostov-Don	21	1.	0,455	1.	0,535	21,750
2. Metz Handball	20	2.	0,447	2.	0,360	19,178
3. CSM Bucuresti	17	3.	0,101	5.	-0,052	15,345
4. FTC-Rail Cargo Hungaria	16	5.	0	4.	0	16
5. Vipers Kristiansand	16	4.	0,016	3.	0,132	17,355
6. Team Esbjerg	12	6.	-0,335	6.	-0,397	11,241
7. RK Krim Mercator	7	7.	-0,788	7.	-0,712	7,758
8. SG BBM Bietigheim	3	8.	-1,389	8.	-1,311	3,372

voltak elmaradt mérkőzések, de ezen odaítélt pontok nem változtatták meg a végleges helyezéseket a csoportban. Az odaítélt pontoknak a "nagy" győztese az utolsó HC Podravka Vegeta csapata volt, de az így kapott többletpontok segítségével sem volt képes változtatni helyezését. Míg a Brest Bretagne Handball csapata kevesebb pontot szerzett az odaítéltekkel, mint szerezhethetett volna, ha minden mérkőzését lejátszta.

20. táblázat. A B csoport eredményei

Csapatok hivatalos rangsora	pont	$TH_3_N^1$		$TH_3_N^2$		
		helyezés	$\hat{m}_i$	helyezés	$\hat{m}_i$	várható pont
1. Győri Audi ETO KC	24	1.	1,129	1.	1,210	24
2. CSKA	23	2.	1,112	2.	1,172	23
3. Brest Bretagne Handball	17	3.	0,351	3.	0,536	18,629
4. Odense Handbold	13	4.	0	4.	0	13
5. Buducnost	12	5.	-0,165	5.	-0,179	12
6. SCM Ramnicu Valcea	10	6.	-0,404	6.	-0,400	10,286
7. BV Borussia Dortmund	9	7.	-0,480	7.	-0,525	8,811
8. HC Podravka Vegeta	4	8.	-1,155	8.	-1,546	2,273

A rájátszásban az A csoport legjobbját a B csoport legrosszabbjával játszotta, az A csoport másodikként a B csoport utolsó előttiével és így tovább, lásd a 23. táblázatban. Mivel a két rangsor különálló és nincs adat amivel összehasonlíthatnánk, így vettük a két legjobb és két legrosszabb közti mérkőzéseket. Nagy valószínűséggel állítható, hogy a legjobbak legyőzik a legrosszabb csapatokat. Vagyis a Győri Audi ETO KC megveri az SG BBM Bietigheim csapatát, míg a Rostov-Don győz a HC Podravka Vegeta csapata ellen, ahogy ez a valóságban meg is történt. Ha ezen összekötő mérkőzéseket behúzzuk,

amelyek nem jelentenek jelentős többlettudást, mivel szinte borítékolható eredmények, akkor a két rangsorból egy összefűzött rangsor készíthető, mert kiértékelhető az adathalmaz egyesített adathalmazként. Az így elkészült összefűzött rangsort a 21. táblázat tartalmazza. Ezen rangsor alapján már el tudtuk készíteni az előrejelzést a rájátszás meccseire.

21. táblázat. A legjobbak és a legrosszabbak egymás elleni mérkőzéseinek vélt (és valós) eredményeit felhasználó egységes rangsor

	TH_3_N	$\hat{m}_i$
1	Győri Audi ETO KC	2,689
2	CSKA	2,654
3	Rostov-Don	2,301
4	Metz Handball	2,108
5	Brest Bretagne Handball	2,032
6	Vipers Kristiansand	1,878
7	FTC-Rail Cargo Hungaria	1,743
8	CSM Bucuresti	1,688
9	Odense Handbold	1,521
10	Buducnost	1,343
11	Team Esbjerg	1,331
12	SCM Ramnicu Valcea	1,127
13	BV Borussia 09 Dortmund	1,011
14	RK Krim Mercator	1,010
15	SG BBM Bietigheim	0,390
16	HC Podravka Vegeta	0

A 22. táblázat tartalmazza a 8 előrejelzett továbbjutó csapatot, a szakértői előrejelzéseket [89] és a Thurstone modell szerinti előrejelzéseket. Azon csapatok neveit megvastagítottuk, melyek ténylegesen továbbjutottak a legjobb 8 csapat közé. A legjobb 8 közé jutó csapat közül 6-ot sikeresen előrejeleztek a szakértők, míg a mi módszerünk csak egy csapatot nem volt képes előrejelezni, hetet igen. Az FTC-Rail Cargo Hungaria továbbjutását mindkét lista előrejelezte, de végül nem jutott tovább. Az odaítélt pontok jelentős vesztesét, a Vipers Kristiansand csapatát mi továbbjutónak jeleztük, míg a szakértők nem. Végül azonban továbbjutott ez a csapat.

A 23. táblázat tartalmazza az egyes mérkőzéseket a rájátszások, negyeddöntők és döntők esetében.

Mivel az erősebb csapatok hajlamosak arra, hogy a gyengébb csapatok ellen nem teljes erővel játszanak és az erősebb ellenfelek ellen játszanak teljes erővel, így olyan rangsort is kialakítottunk, mely csak a legjobb 8 csapat

22. táblázat. A legjobb nyolc csapat esetén az előrejelzések (félkövér betűk: helyes előrejelzések)

Szakértői előrejelzés	TH_3_N alapján előrejelzés
Győri Audi ETO KC	Győri Audi ETO KC
CSKA	CSKA
Rostov-Don	Rostov-Don
Metz Handball	Metz Handball
FTC-Rail Cargo Hungaria	Brest Bretagne Handball
Brest Bretagne Handball	Vipers Kristiansand
CSM Bucuresti	FTC-Rail Cargo Hungaria
Odense Handbold	CSM Bucuresti

egymás elleni eredményeit tartalmazza és azon mérkőzések eredményeiből számoltuk ki az erősségeket. Ezen rangsor és az erősségek a 24. táblázatban láthatók. Meglepőek az eredmények: a korábbi rangsorokban a Győri Audi ETO KC csapata volt az első, míg itt a Vipers Kristiansand csapata vezeti a rangsort határozottan megelőzve a többi csapatot!

Mivel ezen eredmények fajsúlyosabb eredmények, mintha az összes meccset vennénk, így már ezen eredmények figyelembe vételével készítettük el a következő előrejelzéseket. A legjobb 4 csapat előrejelzett listája a 25. táblázatban látható. Míg a szakértők 2-t, addig a modellünkkel mind a 4 továbbjutót el tudtuk találni.

Végül a Final4 eredményei és azok modell általi valószínűségei a 26. táblázatban láthatóak. A Vipers Kristiansand és CSKA csapatok elődöntő mérkőzésének a kimenetele megegyezett az előrejelzettel. A Győri Audi ETO KC és a Brest Bretagne Handball mérkőzésén bár egyértelműen a győri csapat győzelme volt várható, a meccs végén döntetlen alakult ki, az utolsó percben szerzett Brest Bretagne Handball góllal, és a hosszabbítás is döntetlennel zárult. Végül a büntetők során a győri csapat alulmaradt. A döntő mérkőzés a vártak szerint alakult és a Vipers Kristiansand csapat győzelmét hozta el, míg a Győri Audi ETO KC csapata a 3. helyért vívott mérkőzést megnyerte. 3 mérkőzést sikeresen előrejeleztünk a modell segítségével, míg a 4. mérkőzés esetében csak egy hajszálon múlt az előrejelzett eredmény.

23. táblázat. A mérkőzések eredményei

Rájátszások			
SG BBM Bietigheim	Győri Audi ETO KC	20-37	28-32
HC Podravka Vegeta	Rostov-Don	20-29	24-42
RK Krim Mercator	CSKA	25-20	21-27
BV Borussia Dortmund	Metz Handball	0-10	0-10
Team Esbjerg	Brest Bretagne Handball	27-33	27-30
SCM Ramnicu Valcea	CSM Bucuresti	24-33	27-21
Vipers Kristiansand	Odense Handbold	35-36	30-26
Buducnost	FTC-Rail Cargo Hungaria	22-19	28-29
Negyedöntő			
Buducnost	Győri Audi ETO KC	19-30	21-24
Vipers Kristiansand	Rostov-Don	34-27	23-23
CSM Bucuresti	CSKA	32-27	19-24
Brest Bretagne Handball	Metz Handball	34-24	26-26
Elődöntő			
Győri Audi ETO KC	Brest Bretagne Handball	23(2)-23(4)	
Vipers Kristiansand	CSKA	33-30	

24. táblázat. A legjobb 8 csapat és erősségeik csak az egymás elleni mérkőzéseik alapján

Csapat	$\hat{m}_i$
Vipers Kristiansand	2,117
Rostov-Don	1,729
Győri Audi ETO KC	1,649
CSKA	1,462
Brest Bretagne Handball	0,764
Metz Handball	0,740
CSM Bucuresti	0,720
Buducnost	0

25. táblázat. A Final4 várható résztvevői (félkövér betűk: helyes előrejelzések)

Szakértői előrejelzés	TH_3_N alapján előrejelzés
Győri Audi ETO KC	Vipers Kristiansand
CSKA	Győri Audi ETO KC
Rostov-Don	CSKA
Metz Handball	Brest Bretagne Handball

26. táblázat. Az első oszlopban szereplő csapatok győzelmének valószínűsége (félkövér betűk: az eredmény megegyezik a valószínűséggel)

Csapat 1	Csapat 2	Valószínűség
Vipers Kristiansand	CSKA	0,744
Győri Audi ETO KC	Brest Bretagne Handball	0,812
Brest Bretagne Handball	Vipers Kristiansand	0,088
Győri Audi ETO KC	CSKA	0,574

Amint látható, sikerült az elmaradt mérkőzések eredményeit is pótolnunk a modell segítségével, olyan formában, amely helytállóbb az erősségek tekintetében, mint a pontok odaítélése. A modellel igen kevés információt felhasználva képesek voltunk rangsorokat összefésülni megbízhatóan és azokból előrejelzéseket elkészíteni. Ezen előrejelzések nem csak megbízhatóaknak, de kimondottan jóknak minősültek összevetve a szakértői előrejelzésekkel.

6.5. Sportesemények előrejelzésének pontossága

Az előrejelzések pontosságát más kiemelkedő modellekkel is össze akartuk vetni, melyet a [82] cikkben tettünk meg. Háromféle modellt alkalmaztunk, az 5-opciós modellt (BT_5_N -t), az előnyt is alkalmazó modellt (BT_5_K -t) és az időbeli súlyozást is alkalmazó modellt ($BT_5_K_V$ -t).

Lasek és Gagolewski a [90] cikkben alaposan elemeztek és összehasonlítottak több ígéretes modellt és azok előrejelző képességét labdarúgás eredményeire vonatkozóan. A korábban említett három modellünket (BT_5_N -t, BT_5_K -t, $BT_5_K_V$ -t) fogjuk összehasonlítani a cikkükben legjobbnak ítélt módszerrel. Ez a módszer az egyparaméteres Poisson-modell iteratív változata, amelyet a cikkben PR_1^I -gyel jelöltek, itt pedig a P_I_1 jelölést fogom használni rá.

Az 5 legerősebb elsőosztályú európai bajnokság meccseit jeleztük előre modelljeinkkel. Ezek az angol, a francia, német, olasz és spanyol nemzeti bajnokságok, melyek csapatai a legjobban szerepelnek a Bajnokok Ligájában és az Európa Ligában.

Négy általánosan elfogadott mérőszámot alkalmaztunk, hogy ezek által vessük össze a modellek előrejelzőképességét mások előrejelzéseivel. Ezek megadásához szükséges az alábbi jelölés:

A helyes eredményt az $\underline{y}$ vektor tartalmazza, ahol két csapat mérkőzésekor, összehasonlításakor az y_1 a hazai csapat győzelmét, a döntetlen kimenetelt az y_2 , a hazai csapat vereségét az y_3 reprezentálja. Amennyiben $y_i = 0$, nem az i -edik esemény következett be, míg $y_i = 1$ esetén az i -edik esemény következett be. Például ha a hazai csapat győzött akkor $y_1 = 1, y_2 = 0, y_3 = 0$. Az egyes modellek esetén a $\hat{\underline{p}}$ vektor tartalmazza a modell által adott valószínűségeket az egyes kimenetekre.

Az irodalomban általánosan használt négy mérőszámot alkalmaztunk [90, 91, 92, 93, 94]:

- A logaritmusos veszteségmutató (LogLoss) számszerűsíti, hogy a modell előrejelzései milyen mértékben térnek el a tényleges eredménytől. A mutatószám minél kisebb, annál jobb, kisebb hiba esetén jobban teljesít a modell. A LogLoss pozitív értékeket vehet csak fel.

$$LogLoss(\underline{y}, \hat{\underline{p}}) = - \sum_{s=1}^3 y_s \cdot \ln(\hat{p}_s). \quad (58)$$

- A rangsorolt valószínűségi pontszám (RPS) egy olyan mutató, amelyet a kettőnél több kimenetelű előrejelzések pontosságának értékelésére használnak. Kumulatív módon veszi figyelembe a módszer az eltéréseket. Vagyis az előrejelzéstől minél jobban eltér a bekövetkezett esemény, annál nagyobb lesz a hiba. Például, ha a tényleges eredmény "rosszabb" volt, de mi "jobbat" jósoltunk, az RPS magasabb lesz, mintha "egyforma"-t jósoltunk volna. Az RPS 0 és 1 közötti értékeket vesz fel, a kisebb érték a jobb.

$$RPS(\underline{y}, \underline{\hat{p}}) = \frac{1}{2} \sum_{s=1}^2 \left(\sum_{l=1}^s y_l - \sum_{l=1}^s \hat{p}_l \right)^2. \quad (59)$$

- A Brier-pontszám (BS) az egyes előrejelzések pontosságát jelzi azáltal, hogy méri, mennyire térnek el a becsült valószínűségek a tényleges eredményektől. Csak nemnegatív értékeket vesz fel a BS, a legkisebb értéke 0, a legnagyobb értéke pedig 2/3. Ebben a mutatóban is a kisebb értékek jobb előrejelzési képességet jelentenek.

$$BS(\underline{y}, \underline{\hat{p}}) = \frac{1}{3} \sum_{s=1}^3 (y_s - \hat{p}_s)^2 \quad (60)$$

- A pontosság (ACC) mérőszámmal azt mérhetjük, hogy hányszor sikerült helyesen eltalálni a helyes kimenetelt. A legnagyobb valószínűségű kimenetelt jelezzük előre az egyes esetekben. Itt az argmax függvény azzal a koordinátaindexszel tér vissza, ahol a legnagyobb az érték. Ha több maximum is van, akkor a legelsőt választjuk ki közülük. Akkor és csak akkor 1 a függvény értéke, ha az előrejelzés helyes, különben 0. Az összes előrejelzés átlagolásával megkapjuk a helyes előrejelzések arányát. Ez az átlag 0 és 1 között mozog, a nagyobb értékek jobb előrejelzési teljesítményt jelentenek. Az 1-es érték azt jelenti, hogy minden előrejelzés helyes.

$$ACC(\underline{y}, \underline{\hat{p}}) = \begin{cases} 1, & \text{ha } \arg \max(\underline{y}) = \arg \max(\underline{\hat{p}}) \\ 0, & \text{különben} \end{cases} \quad (61)$$

Döntenünk kellett, hogy hány opciós modellben értékeljük ki a mérkőzések eredményeit. Mivel döntetlen lehetséges eredmény a futballban, így a 3- és 5-opciós modellek közül választhattunk. Mivel a mérkőzések száma elég

magas, és az 5-opciós modell több információt tartalmaz, ezért ezen modell mellett döntöttünk. Így meg kellett határozni, mekkora gólkülönbségtől számít a győzelem nagy győzelemnek. 2 gólnál húztam meg ezt a határt, mivel onnan már nagyon nehéz változtatni az eredményen, azaz azon, hogy győz vagy veszít a csapat; ugyanis ha egy gólt szerez is a hátrányban lévő csapat, még akkor is vesztesre áll.

Az 5 legjobb európai elsőosztályú labdarúgó bajnokság eredményeit értékeltem ki és jeleztem előre az egyes modellek esetén. Míg a 2021/22, 2022/23-as évet tanításra, addig a 2023/24-es évet előrejelzésre használtam. Csak azon mérkőzéseket vettem figyelembe, ahol mindkettő csapat mind a három évben benne volt az elsőosztályú ligában. Összesen 3066 mérkőzést használtam fel tanításra és előrejelzésre: 546 angol, 630 francia, 630 német, 630 olasz, 630 spanyol mérkőzést.

Átlagosan 25% körül alakult a döntetlenek, 37% volt a kis győzelmek (1 gólkülönbséges mérkőzések), 38% a nagy győzelmek (2 vagy több gólkülönbséggel végződő mérkőzések) aránya.

A labdarúgás esetén a hazai csapatnak általában előnye van, mivel a lelátókról neki szurkolnak. Előnyt alkalmazó modelleket is alkalmaztam (BT_5_K és $BT_5_K_V$) a kiértékelésekhez. Minden egyes forduló eredményeit az azt megelőző fordulók mérkőzéseiből jeleztem előre. Mivel a csapatok erősségei változnak idővel, ezért időbeli súlyozást is alkalmazó modellt is alkalmaztam ($BT_5_K_V$). A korábban bemutatott időbeli súlyozáshoz egy olyan q értéket kerestem, amellyel a legkorábbi meccsek még nem veszítik el teljesen értéküket, de idővel ténylegesen csökken a mérkőzések súlya, így esett a választásom a $q = 0,999$ -re. Nagyjából 3 évet ölel fel az időszak amikor mérkőzések eredményeit értékeljük ki, ami nagyjából 1000 napnak felel meg. Az összehasonlítások (mérkőzések) súlyai az idő múlásával fokozatosan csökkennek; például egy 1000 nappal ezelőtti mérkőzés súlya $0,999^{1000} = 0,3677$ lesz.

A négy módszer (P_I_1 , BT_5_N , BT_5_K és $BT_5_K_V$) LogLoss értékeit a 27. táblázat mutatja be. Az eredmények számszerű értékei összhangban vannak a szakirodalomban korábban közölt alacsony értékekkel. Az átlagos teljesítmény tekintetében egyértelműen megállapítható, hogy a BT_5_N modell szerepel a leggyengébben, jelentős a lemaradása a BT_5_K mögött, amelyet a P_I_1 követ. A legalacsonyabb átlagos LogLoss értéket a $BT_5_K_V$ modell produkálja, ami arra utal, hogy az átlagos eredmények (lásd az utolsó sort) alapján ez a modell tekinthető a legjobban teljesítőnek.

Az átlagértékek alkalmazása javasolt, mivel a modellek teljesítménye ligánként eltérő lehet, és egyes modellek bizonyos ligáknál (több győzelem vagy döntetlen esetén) pontosabban képesek előrejelezni, mint mások. Például a BT_5_K modell felülmúlja a LogLoss mérőszám tekintetében a többit az angol, a francia és a spanyol bajnokságok esetében, míg a német és olasz bajnokságok esetén a P_I_1 modell mutatkozik a legeredményesebbnek.

27. táblázat. A LogLoss értékei az egyes modellek alapján

	P_I_1	BT_5_N	BT_5_K	$BT_5_K_V$
Angol	0,9858	1,0076	0,9814	0,9853
Francia	1,0276	1,0281	1,0217	1,0228
Német	0,9665	1,0099	0,9851	0,9769
Olasz	0,9951	1,0199	1,0124	1,0079
Spanyol	0,9951	0,9894	0,9720	0,9752
Átlag	0,9940	1,0110	0,9945	0,9936

Az RPS által kapott eredmények a 28. táblázatban találhatóak. Az eredmények azt mutatják, hogy az egyes esetekben mennyire térnek el az előrejelzések a tényleges eredményektől. Az RPS eredmények közepesen jónak mondhatóak. A módszerek a LogLoss-hoz viszonyítva azonos rangsort adnak RPS esetén is. Az angol, francia és spanyol csapatok esetén BT_5_K felülmúlja mind a P_I_1 -t mind a $BT_5_K_V$ -t. Átlagos értékeket tekintve továbbra is a $BT_5_K_V$ a legjobb, míg a BT_5_N majdnem minden esetben az utolsó helyen szerepel a négy modell közül az eredmények tekintetében.

28. táblázat. Az RPS értékei az egyes modellek alapján

	P_I_1	BT_5_N	BT_5_K	$BT_5_K_V$
Angol	0,2071	0,2157	0,2063	0,2070
Francia	0,2107	0,2103	0,2079	0,2083
Német	0,1946	0,2080	0,2007	0,1979
Olasz	0,2006	0,2065	0,2038	0,2021
Spanyol	0,1897	0,1894	0,1849	0,1862
Átlag	0,2005	0,2060	0,2007	0,2003

A BS szerint számolt eredmények a 29. táblázatban láthatóak. Itt kicsit változik a rangsor az átlagos értékek tekintetében a modellek között. Továbbra

is utolsó helyen van jelentős lemaradással a BT_5_N , előtte a P_I_1 , majd a BT_5_K következik. A legjobban teljesítő modell a $BT_5_K_V$ az átlagos értékek tekintetében.

Az eredmények alapján az angol, a francia és a spanyol csapatok esetén mind a BT_5_K , mind a $BT_5_K_V$ jobb teljesítményt nyújt, mint a P_I_1 . A német és olasz csapatok esetén azonban a P_I_1 kissé jobban teljesít, mint a $BT_5_K_V$ és a BT_5_K . Érdekes megjegyezni, hogy a 0,2 körüli Brier-pontszám jó eredménynek számít a három kimenetelű modellek esetében, mint például a labdarúgó-mérkőzések kimenetelének előrejelzésében használt modelleknél.

29. táblázat. A BS értékei az egyes modellek alapján

	P_I_1	BT_5_N	BT_5_K	$BT_5_K_V$
Angol	0,1951	0,2008	0,1940	0,1947
Francia	0,2061	0,2065	0,2046	0,2049
Német	0,1915	0,2006	0,1955	0,1938
Olasz	0,1998	0,2033	0,2016	0,2005
Spanyol	0,1979	0,1976	0,1935	0,1940
Átlag	0,1981	0,2018	0,1978	0,1976

A pontosság (ACC) a helyes eredmény-előrejelzések arányát mutatja. Az ACC által kapott eredmények a 30. táblázatban találhatók. A szakirodalomban a 0,52-es pontosság már jónak számít, ha nem túlillesztett gépi tanulási modell eredménye. Ebben az esetben is a BT_5_N szerepel leggyengébben, míg továbbra is a legjobban a $BT_5_K_V$ szerepel átlagosan. A BT_5_K második helyen van a mérőszám szerint jelentősen megelőzve a harmadik P_I_1 modellt. Nem szignifikáns a különbség az első kettő modell között.

Ezen mérőszám nagyban eltér a korábbi három mérőszám eredményeitől. Míg a LogLoss, RPS, BS eredményeken alapuló rangsorok egybecsengenek, addig az ACC-n alapuló rangsor jelentősen eltér a többitől.

A spanyol bajnokság esetén különösen szembetűnő az, hogy a BT_5_N modell sokkal jobban teljesít, mint a többi modell. Ebben a bajnokságban is elmondható átlagosan, hogy a hazai pálya előnyt nyújt a csapatok számára. Amikor erősebb vagy gyengébb csapat játszik másik csapattal, akkor ez meg is mutatkozik, de amikor középerős csapatok mérkőznek meg egymással, akkor sokszor idegenbeli győzelem született. Ez magyarázhatja a BT_5_N kiváló teljesítését ezen bajnokság esetén és a többi modell gyengébb szereplését.

A négy mérőszámot figyelembe véve az általunk alkalmazott három

30. táblázat. Az ACC értékei az egyes modellek alapján

	P_I_1	BT_5_N	BT_5_K	$BT_5_K_V$
Angol	0,5495	0,4945	0,5385	0,5385
Francia	0,4905	0,4714	0,5238	0,5190
Német	0,5238	0,4905	0,5286	0,5476
Olasz	0,5333	0,5000	0,5286	0,5286
Spanyol	0,4857	0,5143	0,4952	0,4857
Átlag	0,5166	0,4941	0,5229	0,5239

módszer igen jól teljesített. A BT_5_N modell hátránya bár jelentős volt a többi modellel szemben, de a modell nem használ sem előnykezelést, sem időbeli súlyozást, így ezt figyelembe véve mondható, hogy jól teljesített. A P_I_1 modell és a BT_5_K modell fej-fej mellett teljesítettek a különféle mérőszámok tekintetében az átlagos eredményeket tekintve. Míg összesítve a legjobb eredményeket a $BT_5_K_V$, az előnykezelést és időbeli súlyozást is alkalmazó modellünk érte el. Nagyon jól látszik az eredményekből, hogy érdemes olyan modelleket alkalmazni előrejelzésre, amelyek a mérkőzések pusztá eredményein kívül többletinformációt is feldolgoznak.

7. Összefoglalás

Páros összehasonlítási modellekkel alapszakos hallgatóként kezdtem el foglalkozni, majd mesterszakon is folytattam ezen kutatásaimat. A (fél)évek során 5 TDK dolgozatot írtam, s ezek közül négyet az OTDK-n is bemutattam. A munkát PhD tanulmányaim során is folytattam, s így jutottam el a dolgozatban ismertetett eredményekhez.

Az elmúlt 7 év alatt részletesen megismertem a sztochasztikus háttérrel rendelkező Thurstone-motivált modelleket, és jelentős eredményeket sikerült elérnem a különböző modellekkel kapcsolatosan témavezetőim segítségével. Az eredményeim egy része elméleti eredmény, másik része pedig a modellek innovatív alkalmazása. Ezen alkalmazások megvilágítják és egyben kidomborítják a modellek előnyös tulajdonságait.

Az eredményeimet 3 fő csoportba soroltam.

Az első terület elméleti eredményeket tartalmaz. Kiemelkedően jelentős eredmény a 3-opciós modellben az adatok kiértékelhetőségének szükséges és elégséges feltételrendszerének megadása, amit egy D1-es újságban publikáltunk. A feltételrendszert nagyszámú számítógépes szimuláció eredményei elemzése alapján sejtettem meg, s ezután sikerült matematikai eszközökkel bizonyítani. A bizonyításhoz analízisbeli és gráfelméletben alkalmazott eszközök voltak szükségesek. Ennek köszönhetően most pontosan eldönthető, hogy az adatok kiértékelhetők-e 3-opciós modellben, vagy sem. Megemlíteném, hogy a 2-opciós modell esetén speciális esetben ismert szükséges és elégséges feltételt is sikerült általánosítani az eloszlásfüggvények széles osztályára.

Az adatok kiértékelhetőségéhez kapcsolódnak még a modell opciószám-választásával kapcsolatos eredményeim. Ezek megerősítik azon logikus emberi sejtést, hogy a komplexebb (azaz a nagyobb opciószámú) modellek jobban teljesítenek az információ kinyerés szempontjából, amennyiben elégséges páros összehasonlítási adat áll rendelkezésre.

Kutatási eredményeim tekintetében egy másik jelentős területnek tartom a sztochasztikus páros összehasonlítási modellek közül a 2-opciós Bradley-Terry modell és a PCM alapú modellek összekapcsolását a konzisztencia fogalmán keresztül. A PCM alapú modelleknél általánosan használt konzisztenciafogalmon keresztül sikerült a Bradley-Terry modellben is definíciót megfogalmazni az adatok konzisztenciájának értelmezésére. Ilyen

fogalommal korábban nem találkoztunk a szakirodalomban. Továbbá sikerült egy olyan ekvivalencia állítást bizonyítani, amely alapot adhat más Thurstone-motivált modell esetén is az adatok konzisztenciájának definiálására. A PCM és a Thurstone-motivált modellek közti kapcsolatokat erősíti, hogy számítógépes szimulációkkal megmutattam, hogy mindkét esetben ugyanazokat az összehasonlítási struktúrákat érdemes az adatgyűjtéskor alkalmazni, amennyiben nem tudunk minden objektumot minden más objektummal összehasonlítani. Ez újabb kérdéseket vet fel a két modelles család közti kapcsolatok tekintetében, ami a kutatások folytatását indukálja. A bebizonyított elméleti és szimulációs eredményeket szintén egy D1-es folyóiratban jelentettük meg.

A harmadik terület a modellek újszerű alkalmazása. A modelleket már korábban is alkalmazták különféle területeken, például sporteredmények kiértékelésének területén is. Én kibővítettem az alkalmazásokat, és ezeken a példákon keresztül illusztráltam a modellek előnyös tulajdonságait. Bemutattam, hogy a modellek alkalmasak előrejelzések készítésére akkor is, ha kevés információ áll rendelkezésre, vagy például elmaradt mérkőzések nehezítik az előrejelzést. Megmutattam, hogy kevés összefűző információ esetén is képesek hatékony együttes sorrend készítésére. Összehasonlítva ezen képességét a PCM alapú módszerek és a neurális hálók által biztosított módszerekkel, megmutattam, hogy jobb eredményeket kaphatunk, ha Thurstone-motivált modelleket alkalmazunk. Használtam a modelleket sportág-kiválasztásra, és megmutattam, hogy segítségével hatékonyan és megbízhatóan adhatunk javaslatokat e tekintetben. Munkámban újszerű többkritériumos döntéstámogatáshoz alkalmaztam a 7-opciós modellt, mellyel kiváló eredményeket értem el. Ezzel bemutattam, hogy akár sportolók teljesítményének méréséhez, szakágakba sorolásához nagy segítség lehet egy Thurstone-motivált modell alapú szoftver. Alkalmaztam a modellt valós összehasonlítási eredményeken fényforrások csoportosítására nemzetközi kutatói együttműködés keretében. A modellek előrejelző képességének javítására a szokásos 3-opciós modell helyett 5-opciós modellt alkalmaztam a korábban is használt időbeli súlyozás és hazai pálya előnyének beépítése mellett. Megmutattam, hogy az így megalkotott modell előrejelző képessége vetekszik az irodalomban publikált, legjobbnak talált modell előrejelző képességével, sok esetben meghaladja azt. Eredményeimet Q1/Q2/Q3 besorolású nemzetközi folyóiratokban publikáltam.

Kutatásaim során számos kérdésre választ találtam, de természetesen újabb kérdések is felvetődtek, amelyek megválaszolása további kiterjedt kutatásokat igényel a tématerületen. Többek között érdemes lenne tudni az adatok

kiértékelhetőségének szükséges és elégséges feltételrendszerét általános s -opciós modellben, hasznos lenne inkonzisztencia mérőszámot definiálni a sztochasztikus háttérű modellek esetén is, illetve újabb területeken alkalmazni a modelleket.

Hivatkozások

- [1] Saaty, T. L. (1977). A scaling method for priorities in hierarchical structures, *Journal of Mathematical Psychology*, 15(3), 234–281, [https://doi.org/10.1016/0022-2496\(77\)90033-5](https://doi.org/10.1016/0022-2496(77)90033-5).
- [2] Saaty, T. L. (1990). How to make a decision: the analytic hierarchy process. *European Journal of Operational Research*, 48(1), 9–26
- [3] Bozóki, S.; Fülöp, J.; Rónyai, L. (2010). On optimal completion of incomplete pairwise comparison matrices. *Mathematical and Computer Modelling*, 52(1), 318–333.
- [4] Brunelli, M. (2018). A survey of inconsistency indices for pairwise comparisons. *International Journal of General Systems*, 47(8), 751–771. <https://doi.org/10.1080/03081079.2018.1523156>
- [5] Thurstone, L. (1927). A law of comparative judgment. *Psychological Review*, 34, 273–286. <https://doi.org/10.1037/h0070288>.
- [6] Mosteller, F. (1951). Remarks on the method of paired comparisons: I. The least squares solution assuming equal standard deviations and equal correlations. *Psychometrika*, 16, 3–9. <https://doi.org/10.1007/BF02313422>.
- [7] Bradley, R. A.; Terry, M. E. (1952). Rank analysis of incomplete block designs: I. The method of paired comparisons. *Biometrika*, 39, 324–345.
- [8] Ford Jr, L. R. (1957). Solution of a ranking problem from binary comparisons. *American Mathematical Monthly*, 64, 28–33.
- [9] Glenn, W. A.; (1960). David, H. A. Ties in paired-comparison experiments using a modified Thurstone-Mosteller model. *Biometrics*, 16, 86–109.
- [10] Rao, P. V.; Kupper, L. L. (1967). Ties in paired-comparison experiments: A generalization of the Bradley-Terry model. *Journal of the American Statistical Association*, 62, 194–204. <https://doi.org/10.1080/01621459.1967.10482901>
- [11] Davidson, R. R. (1970). On extending the Bradley-Terry model to accommodate ties in paired comparison experiments. *Journal of the American Statistical Association*, 65, 317–328.

- [12] Luce, R. D. (1959). Individual choice behavior. New York: Wiley.
- [13] Luce, R. D. (1977). The choice axiom after twenty years. *Journal of Mathematical Psychology*, 15(3), 215–233.
- [14] Agresti A. (1992). Analysis of ordinal paired comparison data. *Applied Statistics*, 41(2), 287–297 <https://doi.org/10.2307/2347562>
- [15] Conner, G. R.; Grant, C. P. (2000). An extension of Zermelo’s model for ranking by paired comparisons. *European Journal of Applied Mathematics*, 11(3), 225–247. <https://doi.org/10.1017/S0956792500004150>
- [16] Zermelo, E. (1929). Die Berechnung der Turnier-Ergebnisse als ein Maximumproblem der Wahrscheinlichkeitsrechnung. *Mathematische Zeitschrift*, 29(1), 436–460. <https://doi.org/10.1007/BF01180541>
- [17] Yan, T. (2016). Ranking in the generalized Bradley–Terry models when the strong connection condition fails. *Communications in Statistics-Theory and Methods*, 45(2), 340–353.
- [18] Kovács, B.; Orbán-Mihálykó, É.; & Mihálykó, C. (2025). Analysis of Thurstone-Motivated Models in the Case of Non-Evaluable Data: Methods for Extracting Information. *Mathematics*, 13(16), 2578.
- [19] Hunter, D. R. (2004). MM algorithms for generalized Bradley-Terry models. *The Annals of Statistics*, 32(1), 384–406.
- [20] Orbán-Mihálykó, É.; Mihálykó, Cs.; Koltay, L. (2019). A generalization of the Thurstone method for multiple choice and incomplete paired comparisons. *Central European Journal of Operations Research*, 27, 133–159. <https://doi.org/10.1007/s10100-017-0495-6>.
- [21] Orbán-Mihálykó, É.; Mihálykó, Cs.; Koltay, L. (2019). Incomplete paired comparisons in case of multiple choice and general log-concave probability density functions. *Central European Journal of Operations Research*, 27, 515–532.
- [22] Orbán-Mihálykó, É.; Mihálykó, Cs.; Kajtár, P.; (2019). Általánosított Thurstone-módszer alkalmazásokkal, *Alkalmazott Matematikai Lapok*, 36(2), 255–262 <https://doi.org/10.37070/AML.2019.36.2.09>
- [23] Sahroni, T. R.; Ariff, H. (2016). Design of analytical hierarchy process (AHP) for teaching and learning. In *Proceedings of the*

2016 11th International Conference on Knowledge, Information and Creativity Support Systems (KICSS), IEEE, Yogyakarta, Indonesia, 10–12 November 2016; pp. 1–4.

- [24] Kosztyán, Z. T.; Orbán-Mihálykó, É.; Mihálykó, Cs.;Csányi, V. V.; Telcs, A. (2020). Analyzing and clustering students' application preferences in higher education. *Journal of Applied Statistics*, 47, 2961–2983.
- [25] Cattelan, M.; Varin, C.; Firth, D. (2013). Dynamic Bradley–Terry modelling of sports tournaments. *Journal of the Royal Statistical Society: Series C (Applied Statistics)*, 62, 135–150.
- [26] Macri Demartino, R.; Egidi, L. & Torelli, N. (2025). Alternative ranking measures to predict international football results. *Computational Statistics*, 40, 1899–1917 <https://doi.org/10.1007/s00180-024-01585-z>
- [27] Schauburger, G.; Groll, A.; & Tutz, G. (2018). Analysis of the importance of on-field covariates in the German Bundesliga. *Journal of Applied Statistics*, 45(9), 1561–1578. <https://doi.org/10.1080/02664763.2017.1383370>
- [28] Tutz, G.; Schauburger, G. (2015). Extended ordered paired comparison models with application to football data from German Bundesliga. *AStA – Advances in Statistical Analysis*, 99, 209–227 <https://doi.org/10.1007/s10182-014-0237-1>
- [29] Jeon, J. J.; Kim, Y. (2013). Revisiting the Bradley-Terry model and its application to information retrieval. *Journal of the Korean Data and Information Science Society*, 24, 1089–1099.
- [30] Trojanowski, T. W.; Kazibudzki, P. T. (2021). Prospects and Constraints of Sustainable Marketing Mix Development for Poland's High-Energy Consumer Goods. *Energies*, 14, 8437.
- [31] Montequín, V. R.; Balsera, J. M. V.; Piloñeta, M. D.; Pérez, C.Á. (2020). A Bradley-Terry model-based approach to prioritize the balance scorecard driving factors: The case study of a financial software factory. *Mathematics*, 8, 107.
- [32] Canco, I.; Kruja, D.; Iancu, T. (2021). AHP, a Reliable Method for Quality Decision Making: A Case Study in Business. *Sustainability*, 13, 13932.

- [33] Courcoux, P.; Semenou, M. (1997). Preference data analysis using a paired comparison model. *Food Quality and Preference*, 8, 353–358.
- [34] Rotter, P. (2012). Multimedia information retrieval based on pairwise comparison and its application to visual search. *Multimed Tools Appl*, 60, 573–587. <https://doi.org/10.1007/s11042-011-0828-8>
- [35] Dean, M. L. (1980). Presentation order effects in product taste tests. *The Journal of Psychology*, 105(2), 107–10. <https://doi.org/10.1080/00223980.1980.9915138>.
- [36] Eliason, S. R. (1993). *Maximum likelihood estimation: Logic and practice*. Sage Publications.
- [37] Hvattum, L. M. (2020). Offensive and Defensive Plus–Minus Player Ratings for Soccer. *Applied Sciences*, 10, 7345. <https://doi.org/10.3390/app10207345>
- [38] Held, L.; Vollnhals, R. (2005). Dynamic rating of European football teams. *IMA Journal of Management Mathematics*, 16(2), 121–130. <https://doi.org/10.1093/imaman/dpi004>
- [39] Gyarmati, L.; Mihálykó, C.; Orbán-Mihálykó, É.; & Mihálykó, A. (2026). Evaluability of paired comparison data in stochastic paired comparison models: Necessary and sufficient condition. *European Journal of Operations Research*, 333(3), 852–867. <https://doi.org/10.1016/j.ejor.2026.01.029>.
- [40] Saaty, T. L. (2003). Decision-making with the AHP: Why is the principal eigenvector necessary, *European Journal of Operational Research*, 145(1), 85–91, [https://doi.org/10.1016/S0377-2217\(02\)00227-8](https://doi.org/10.1016/S0377-2217(02)00227-8)
- [41] S. Bozóki S.; Tsyganok. V. (2019). The (logarithmic) least squares optimality of the arithmetic (geometric) mean of weight vectors calculated from all spanning trees for incomplete additive (multiplicative) pairwise comparison matrices. *International Journal of General Systems*, 48(4), 362–381.
- [42] Hartung, F. (2004). *Bevezetés a numerikus analízisbe*, Veszprémi Egyetemi Kiadó, Veszprém
- [43] Orbán-Mihálykó, É.; Mihálykó, C.; & Gyarmati, L. (2024). Evaluating the capacity of paired comparison methods to aggregate rankings of separate groups. *Central European Journal of Operations Research*, 32(1), 109–129. <http://doi.org/10.1007/s10100-023-00839-3>

- [44] Orbán-Mihálykó, É.; Mihálykó, C.; & Gyarmati, L. (2022). Application of the Generalized Thurstone Method for Evaluations of Sports Tournaments' Results. *Knowledge*, 2(1), 157–166. <http://doi.org/10.3390/knowledge2010009>
- [45] Gyarmati, L.; Orbán-Mihálykó, É.; & Mihálykó, C. (2023). Comparative Analysis of the Existence and Uniqueness Conditions of Parameter Estimation in Paired Comparison Models. *Axioms*, 12(6), 575. <http://doi.org/10.3390/axioms12060575>
- [46] Prékopa, A. (1973). On logarithmic concave measures and functions. *Acta Scientiarum Mathematicarum*, 34, 335.
- [47] Stern, H. (1990). A continuum of paired comparisons models. *Biometrika*, 77, 265–273.
- [48] Spearman, C. (1904). The proof and measurement of association between two things. *The American Journal of Psychology*, 15(1), 72–101. <https://doi.org/10.2307/1412159>
- [49] Gyarmati, L.; Mihálykó, C.; & Orbán-Mihálykó, É. (2024). Application of Different Option Numbers in Thurstone Motivated Models. In 2024 IEEE 3rd Conference on Information Technology and Data Science (CITDS) Proceedings, Debrecen, Hungary, 26-28 August 2024 ; pp. 1–6. <http://doi.org/10.1109/CITDS62610.2024.10791388>
- [50] Gyarmati, L.; Mihálykó, C.; & Orbán-Mihálykó, É. (2024). Algorithm for Option Number Selection in Stochastic Paired Comparison Models. *Algorithms*, 17(9). <http://doi.org/10.3390/a17090410>
- [51] Mihálykóné Orbán, É.; Mihálykó, C.; & Kajtár, P. (2019). Általánosított Thurstone-módszer alkalmazásokkal. *Alkalmazott Matematikai Lapok*, 36(2), 255–262. <http://doi.org/10.37070/AML.2019.36.2.09>
- [52] Wikipédia. (2025). 2024-es labdarúgó-Európa-bajnokság. Wikipédia https://hu.wikipedia.org/wiki/2024-es_labdar%C3%BAG%C3%B3-Eur%C3%B3pa-bajnoks%C3%A1g; Hozzáférés dátuma: 2025-10-28
- [53] Bozóki, S.; Fülöp, J.; Poesz, A. (2015). On reducing inconsistency of pairwise comparison matrices below an acceptance threshold. *Central European Journal of Operations Research*, 23, 849–866. <https://doi.org/10.1007/s10100-014-0346-7>

- [54] Koyun Y. S.; Ozkir, V. (2018). Extended Consistency analysis for pairwise comparison method. *International Journal of the Analytic Hierarchy Process*, 10(1). <https://doi.org/10.13033/ijahp.v10i1.506>
- [55] Pant, S.; Kumar, A.; Ram, M.; Klochkov, Y.; Sharma, H. K. (2022). Consistency Indices in Analytic Hierarchy Process: A Review. *Mathematics*, 10(8), 1206. <https://doi.org/10.3390/math10081206>
- [56] Bozóki, S.; Rapcsák, T. (2008). On Saaty's and Koczkodaj's inconsistencies of pairwise comparison matrices. *Journal of Global Optimization*, 42, 157–175 <https://doi.org/10.1007/s10898-007-9236-z>
- [57] Starczewski, T. (2023). Simulation Research on the Relationship between Selected Inconsistency Indices Used in AHP. *Entropy*, 25(10), 1464. <https://doi.org/10.3390/e25101464>
- [58] Rácz, A. (2022). Mixed-integer linear-fractional programming model and it's linear analogue for reducing inconsistency of pairwise comparison matrices, *Information Sciences*, 592, 192–205, <https://doi.org/10.1016/j.ins.2022.01.077>.
- [59] Rácz, A. (2025). Optimization model for reducing inconsistency of pairwise comparison matrices with minimal change. *Opsearch*, 62, 2272–2288 <https://doi.org/10.1007/s12597-024-00898-3>
- [60] Szybowski, J.; Kułakowski, K.; Prusak, A. (2020). New inconsistency indicators for incomplete pairwise comparisons matrices. *Mathematical Social Sciences*, 108, 138–145 <https://doi.org/10.1016/j.mathsocsci.2020.05.002>
- [61] Gyarmati, L.; Orbán-Mihálykó, É.; Mihálykó, C.; Szádoczki, Z.; & Bozóki, S. (2023). The incomplete analytic hierarchy process and Bradley–Terry model: (In)consistency and information retrieval. *Expert Systems with Applications*, 229(Part B). <http://doi.org/10.1016/j.eswa.2023.120522>
- [62] Szádoczki, Z.; Bozóki, S. (2025). Optimal sequences for pairwise comparisons: the graph of graphs approach. *Annals of Operations Research*, 353, 1099–1122. <https://doi.org/10.1007/s10479-025-06752-z>
- [63] (2008). Pearson's Correlation Coefficient. In: Kirch, W. (Ed.), *Encyclopedia of Public Health*. Springer, Dordrecht. https://doi.org/10.1007/978-1-4020-5614-7_2569

- [64] Kendall, M. G. (1938). A new measure of rank correlation. *Biometrika*, 30(1–2), 81–93.
- [65] Mihálykó-Orbán É. (2022). Simulation data for the paper entitled "The incomplete Analytic Hierarchy Process and Bradley-Terry Model: (In)consistency and information retrieval." <https://math.uni-pannon.hu/~orban/Results.zip>; Hozzáférés dátuma: 2025-11-11
- [66] Tóth-Merényi, A.; Mihálykó, C.; Orbán-Mihálykó, É.; & Gyarmati, L. (2025). Exploring Consistency in the Three-Option Davidson Model: Bridging Pairwise Comparison Matrices and Stochastic Methods. *Mathematics*, 13(9), 1374. <http://doi.org/10.3390/math13091374>
- [67] Qin, Y.; Hashim, S. R. M.; Sulaiman, J. (2022). An Interval AHP Technique for Classroom Teaching Quality Evaluation. *Education Sciences*, 12, 736.
- [68] Xu, W.; Zhao, S. (2020). The influence of entrepreneurs' psychological capital on their deviant innovation behavior. *Frontiers in Psychology*, 11, 1606.
- [69] Hwang, J. W.; Kim, J. Y. (2020). A study on industrial security psychology. In *Proceedings of the 2020 IEEE International Conference on Big Data and Smart Computing (BigComp)*, Busan, Republic of Korea, 19–22 February 2020; pp. 500–506.
- [70] Improta, G.; Converso, G.; Murino, T.; Gallo, M.; Perrone, A.; Romano, M. (2019). Analytic hierarchy process (AHP) in dynamic configuration as a tool for health technology assessment (HTA): The case of biosensing optoelectronics in oncology. *International Journal of Information Technology & Decision Making*, 18, 1533–1550.
- [71] Szabó, F.; Gazdag-Keri, R. Schanda, J.; Csuti, P.; Orbán-Mihálykó, É. (2014). A study of preferred colour rendering of light sources: Home lighting. *Lighting Research and Technology*, 48. [10.1177/1477153514555536](https://doi.org/10.1177/1477153514555536).
- [72] Katarina, D.; Nurrohman, A.; Putra, A. S.; et al. (2021). Decision Support System For The Best Student Selection Recommendation Using Ahp (Analytic Hierarchy Process) Method. *International Journal of Educational Research and Social Sciences*, 2, 1210–1217.

- [73] Anderson, A. (2015). A Monte Carlo comparison of alternative methods of maximum likelihood ranking in racing sports. *Journal of Applied Statistics*, 42, 1740–1756.
- [74] Bozóki, S.; Csató, L.; Temesi, J. (2016). An application of incomplete pairwise comparison matrices for ranking top tennis players. *European Journal of Operations Research*, 248, 211–218. <https://doi.org/10.1016/j.ejor.2015.06.069>.
- [75] Csató, L. (2013). Ranking by pairwise comparisons for Swiss-system tournaments. *Central European Journal of Operations Research*, 21, 783–803.
- [76] Csató, L. (2017). On the ranking of a Swiss system chess team tournament. *Annals of Operations Research*, 254, 17–36.
- [77] Urbaniak, K.; Wątróbski, J.; Sałabun, W. (2020). Identification of players ranking in e-sport. *Applied Sciences*, 10, 6768.
- [78] Kiani Mavi, R.; Kiani Mavi, N.; Kiani, L. (2012.) Ranking football teams with AHP and TOPSIS methods. *International Journal of Decision Sciences, Risk and Management*, 4, 108–126.
- [79] Gyarmati, L.; Orbán-Mihálykó, É.; Mihálykó, C.; & Vathy-Fogarassy, Á. (2023). Aggregated Rankings of Top Leagues' Football Teams: Application and Comparison of Different Ranking Methods. *Applied Sciences-Basel*, 13(7). <http://doi.org/10.3390/app13074556>
- [80] Gyarmati, L.; Edvy, L.; Mihálykó, C.; & Orbán-Mihálykó, É. (2023). Decision support for sports choices based on performance measurement using the generalized Thurstone method: the model and a case study. In *10th MathSport International Conference Proceedings 2023*, Budapest, Hungary, 26-28 June 2023; pp. 35–40.
- [81] Gyarmati, L.; Edvy, L.; Mihálykó, C.; & Orbán-Mihálykó, É. (2024). Decision support for sports selection based on performance measurement applying the generalized Thurstone method. *International Journal of Sports Science and Coaching*. <http://doi.org/10.1177/17479541241240609>
- [82] Gyarmati, L.; Mihálykó, C.; Orbán-Mihálykó, É. (2025). Forecasting outcomes using multi-option, advantage-sensitive Thurstone-motivated models. *Forecasting*, 7(3), 34. <https://doi.org/10.3390/forecast7030034>

- [83] Burges, C.; Shaked, T.; Renshaw, E.; Lazier, A.; Deeds, M.; Hamilton, N.; Hullender, G. (2005). Learning to rank using gradient descent. In Proceedings of the 22nd International Conference on Machine Learning, Bonn, Germany, 7–11 August 2005; pp. 89–96.
- [84] FIFA.(2021). Transfermarkt. <https://www.transfermarkt.com>; Hozzáféres dátuma: 2021-01-31.
- [85] Kou, G.; Ergu, D.; Chen, Y.; & Lin, C. (2016). Pairwise comparison matrix in multiple criteria decision making. Technological and Economic Development of Economy, 22(5), 738-765. <https://doi.org/10.3846/20294913.2016.1210694>
- [86] Azhar, N.; Mohamed Radzi, N. A.; & Wan Ahmad, W. S. H. M. (2021). Multi-criteria Decision Making: A Systematic Review. Recent Advances in Electrical & Electronic Engineering, 14(8), e291021197468. <https://doi.org/10.2174/2352096514666211029112443>
- [87] Soltanifar, M.; Tavana, M. (2024) A novel pairwise comparison method with linear programming for multi-attribute decision-making. EURO Journal on Decision Processes, 12,100051. <https://doi.org/10.1016/j.ejdp.2024.100051>.
- [88] FitBack. (2023). NETFIT - Hungarian best practice. <https://www.fitbackeurope.eu/en-us/monitoring-fitness/best-practice/hungary-netfit>; Hozzáféres dátuma: 2023-05-23
- [89] WEBEXPERTS. (2022). <https://www.eurohandball.com/en/news/en/gyor-top-power-ranking-going-into-play-offs/>; Hozzáféres dátuma: 2022-01-22
- [90] Lasek, J.; Gagolewski, M. (2021). Interpretable sports team rating models based on the gradient descent algorithm. International Journal of Forecasting, 37(3), 1061–1071. <https://doi.org/10.1016/j.ijforecast.2020.11.008>.
- [91] Koopman, S. J.; Lit, R. (2019). Forecasting football match results in national league competitions using score-driven time series models. International Journal of Forecasting, 35(2), 797–809. <https://doi.org/10.1016/j.ijforecast.2018.10.011>
- [92] Hvattum, L. M.; Arntzen, H. (2010). Using ELO ratings for match result prediction in association football. International Journal of Forecasting, 26(3), 460–470. <https://doi.org/10.1016/j.ijforecast.2009.10.002>.

- [93] Ley C, Wiele T. V. de; Eetvelde, H. V. (2019). Ranking soccer teams on the basis of their current strength: A comparison of maximum likelihood approaches. *Statistical Modelling*, 19(1), 55–73. <https://doi:10.1177/1471082X18817650>
- [94] Constantinou, A. C.; Fenton, N. E. (2012). Solving the Problem of Inadequate Scoring Rules for Assessing Probabilistic Football Forecast Models. *Journal of Quantitative Analysis in Sports*, 8(1). <https://doi.org/10.1515/1559-0410.1418>

A. Appendix

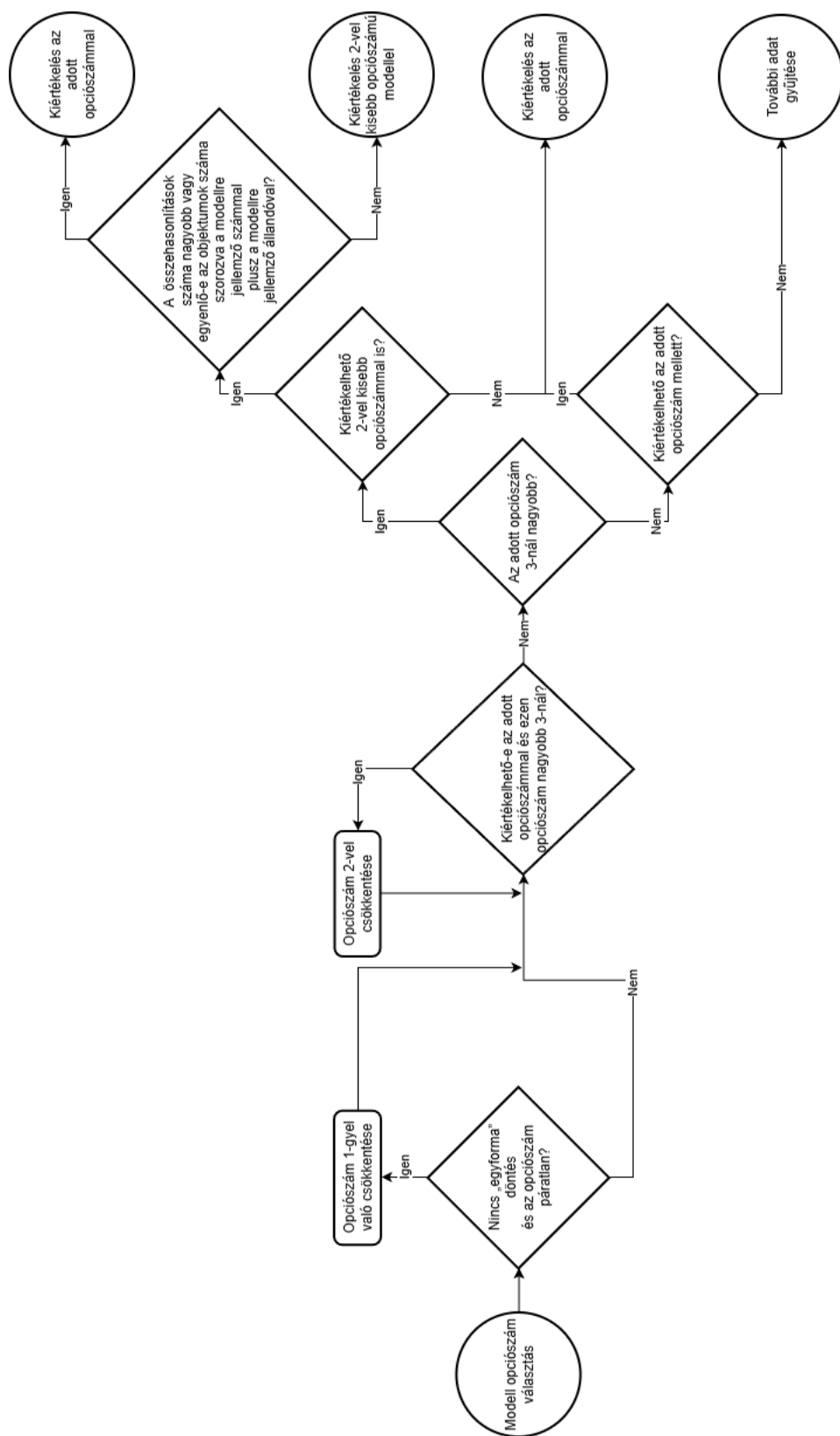
A.1. Optimális opciószám választása

31. táblázat. A 2- és 4-opciós modell összehasonlítása szimuláció segítségével ($m_i : 0 - 2, d_1 = -1, 35$)

$ A $	2 és 4	csak 4	egyik se	$\overline{e^{(2)}}$	$\overline{e^{(4)}}$	$R^{(2),(4)}$
10	23	4006	95971	3,887	6,059	-0,359
11	85	8611	91304	4,566	6,405	-0,287
12	279	14819	84902	4,880	6,823	-0,285
13	708	21872	77420	5,168	7,231	-0,285
14	1473	29086	69441	5,313	6,917	-0,232
15	2646	35962	61392	5,421	6,736	-0,195
16	4195	42121	53684	5,576	6,655	-0,162
17	6401	47080	46519	5,709	6,465	-0,117
18	8981	50959	40060	5,839	6,332	-0,078
19	11999	53683	34318	5,924	6,179	-0,041
20	15499	55244	29257	5,946	5,970	-0,004
21	19284	55878	24838	5,971	5,811	0,027
22	23220	55720	21060	5,939	5,643	0,052
23	27247	54962	17791	5,921	5,456	0,085
24	31453	53502	15045	5,882	5,284	0,113
25	35552	51612	12836	5,808	5,090	0,141
30	54845	39613	5542	5,424	4,296	0,263
40	79977	18949	1074	4,441	3,151	0,410
50	91265	8514	221	3,569	2,417	0,477
60	96142	3797	61	2,898	1,929	0,503
70	98224	1763	13	2,388	1,590	0,502
80	99167	831	2	2,011	1,343	0,497
90	99587	412	1	1,729	1,163	0,487
100	99781	219	0	1,510	1,023	0,476
200	99999	1	0	0,650	0,460	0,412
300	100000	0	0	0,414	0,297	0,394
400	100000	0	0	0,304	0,219	0,387
500	100000	0	0	0,240	0,174	0,382
600	100000	0	0	0,199	0,144	0,381
700	100000	0	0	0,169	0,123	0,378
800	100000	0	0	0,148	0,107	0,376
900	100000	0	0	0,131	0,095	0,374
1000	100000	0	0	0,117	0,085	0,373

32. táblázat. A 3- és 5-opciós modell összehasonlítása szimuláció segítségével
 $(m_i : 0 - 2, d_1 = -1, 35, d_2 = -0, 1)$

$ A $	3 és 5	csak 5	egyik se	$\overline{e^{(3)}}$	$\overline{e^{(5)}}$	$R^{(3),(5)}$
10	39	2283	97678	5,715	7,765	-0,264
11	168	4899	94933	6,203	9,178	-0,324
12	454	8313	91233	6,508	8,666	-0,249
13	936	12437	86627	6,838	8,495	-0,195
14	1788	16617	81595	6,870	8,342	-0,177
15	3041	20513	76446	6,959	8,112	-0,142
16	4738	23837	71425	7,091	7,820	-0,093
17	6766	26470	66764	7,074	7,563	-0,065
18	9171	28704	62125	7,045	7,260	-0,030
19	11901	30146	57953	6,913	6,949	-0,005
20	14931	30912	54157	6,847	6,706	0,021
30	46295	21860	31845	5,583	4,466	0,250
40	66621	10441	22938	4,310	3,181	0,355
50	77156	4594	18250	3,370	2,404	0,402
60	82770	1964	15266	2,688	1,896	0,417
70	86060	887	13053	2,206	1,559	0,415
80	88263	412	11325	1,857	1,316	0,411
90	89784	191	10025	1,593	1,137	0,401
100	90892	92	9016	1,391	1,001	0,390
200	95497	0	4503	0,607	0,452	0,341
300	96976	0	3024	0,387	0,292	0,327
400	97762	0	2238	0,285	0,215	0,323
500	98232	0	1768	0,225	0,171	0,319
600	98530	0	1470	0,186	0,142	0,317
700	98714	0	1286	0,159	0,121	0,315
800	98867	0	1133	0,138	0,105	0,314
900	98991	0	1009	0,122	0,093	0,314
1000	99105	0	895	0,110	0,084	0,314

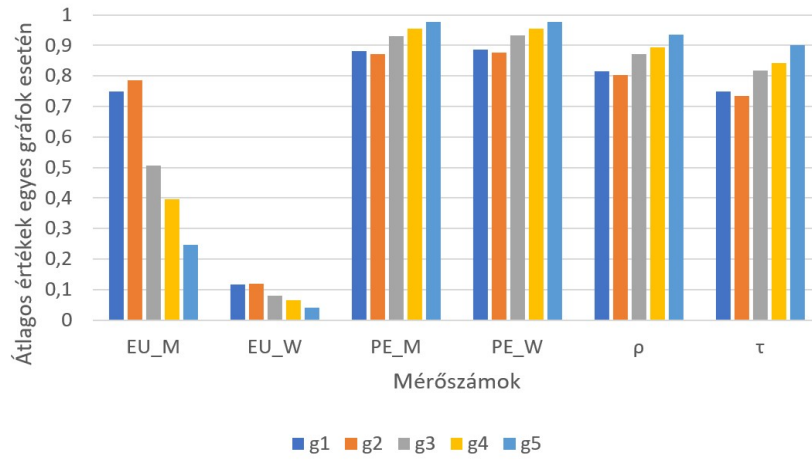


20. ábra. Optimális opciószám kiválasztó algoritmus folyamatábrája

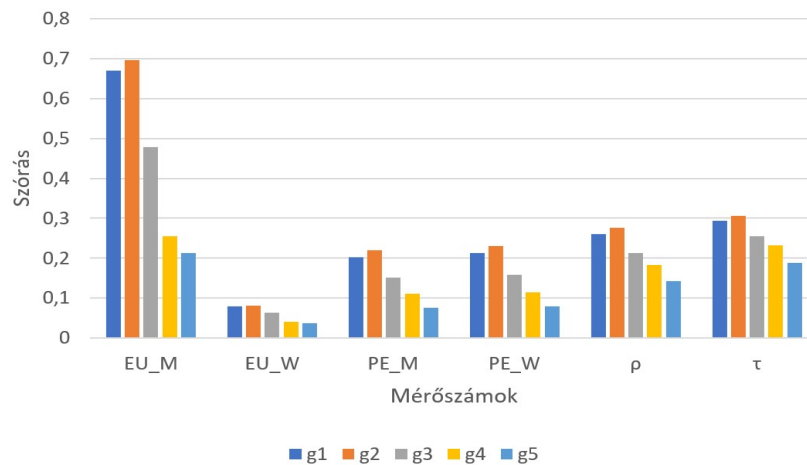
A.2. Optimális összehasonlítási struktúrák a 2-opciós Bradley-Terry modellek esetén inkonzisztens esetben

33. táblázat. Összehasonlítási struktúrák eredményei 0,05-ös perturbációval $n = 5$ esetén

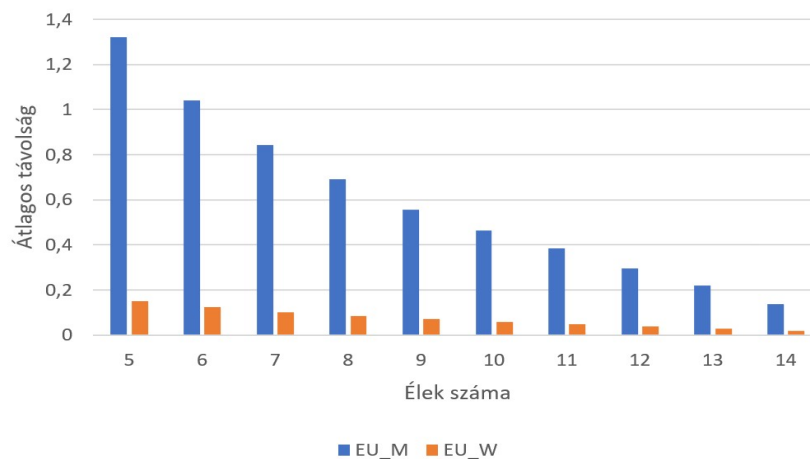
struktúra adatok			Perturbálthoz viszonyítva						Eredetihez viszonyítva					
gráf	sorszám	élszám	EU_M	EU_W	PE_M	PE_W	ρ	τ	EU_M	EU_W	PE_M	PE_W	ρ	τ
	1	4	0,303	0,044	0,981	0,982	0,929	0,878	0,303	0,044	0,981	0,982	0,929	0,878
	2	4	0,321	0,046	0,978	0,979	0,924	0,870	0,321	0,046	0,978	0,979	0,924	0,870
	3	4	0,336	0,048	0,975	0,977	0,919	0,863	0,336	0,048	0,975	0,977	0,919	0,863
	4	5	0,249	0,036	0,986	0,987	0,941	0,897	0,249	0,036	0,986	0,987	0,941	0,897
	5	5	0,229	0,033	0,988	0,989	0,945	0,903	0,229	0,033	0,988	0,989	0,945	0,903
	6	5	0,256	0,037	0,985	0,986	0,939	0,894	0,256	0,037	0,985	0,986	0,939	0,894
	7	5	0,269	0,038	0,983	0,984	0,935	0,888	0,269	0,038	0,983	0,984	0,935	0,888
	8	5	0,207	0,030	0,991	0,991	0,949	0,909	0,207	0,030	0,991	0,991	0,949	0,909
	9	6	0,197	0,028	0,991	0,991	0,953	0,916	0,197	0,028	0,991	0,991	0,953	0,916
	10	6	0,160	0,023	0,994	0,995	0,960	0,927	0,160	0,023	0,994	0,995	0,960	0,927
	11	6	0,186	0,027	0,991	0,992	0,953	0,916	0,186	0,027	0,991	0,992	0,953	0,916
	12	6	0,169	0,024	0,993	0,994	0,958	0,924	0,169	0,024	0,993	0,994	0,958	0,924
	13	6	0,202	0,029	0,990	0,991	0,951	0,914	0,202	0,029	0,990	0,991	0,951	0,914
	14	7	0,137	0,020	0,995	0,995	0,967	0,939	0,137	0,020	0,995	0,995	0,967	0,939
	15	7	0,137	0,020	0,995	0,996	0,966	0,937	0,137	0,020	0,995	0,996	0,966	0,937
	16	7	0,168	0,024	0,993	0,993	0,961	0,929	0,168	0,024	0,993	0,993	0,961	0,929
	17	7	0,128	0,019	0,996	0,996	0,968	0,940	0,128	0,019	0,996	0,996	0,968	0,940
	18	8	0,101	0,015	0,997	0,997	0,975	0,954	0,101	0,015	0,997	0,997	0,975	0,954
	19	8	0,094	0,014	0,998	0,998	0,976	0,954	0,094	0,014	0,998	0,998	0,976	0,954
	20	9	0,059	0,008	0,999	0,999	0,986	0,973	0,059	0,008	0,999	0,999	0,986	0,973
	21	10	0	0	1	1	1	1	0	0	1	1	1	1



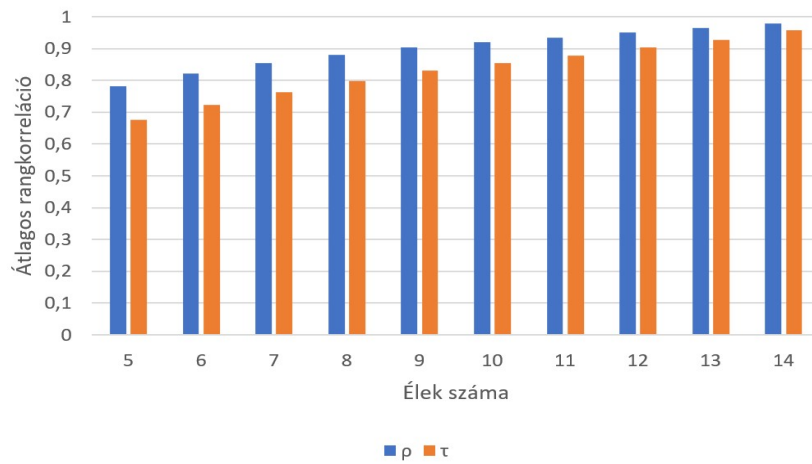
21. ábra. Különböző gráfstruktúrák esetén az információ visszakeresés átlagos mérőszámai: az EU_M, EU_W által meghatározott átlagos távolságok, a PE_M, PE_W, Spearman ρ és Kendall τ által meghatározott átlagos korrelációk, $n = 4$ objektum összehasonlítása esetén $l = 0, 15$ -ös perturbáció alkalmazásával



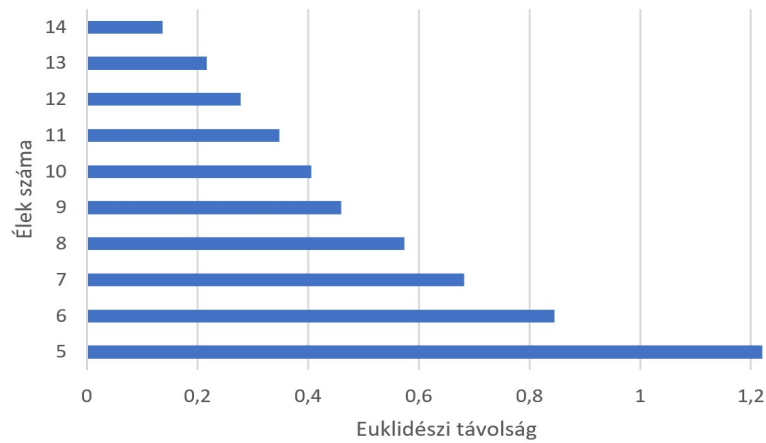
22. ábra. Különböző gráfstruktúrák esetén az információ visszakeresés átlagos szórásai: az EU_M, EU_W által meghatározott szórások, a PE_M, PE_W, Spearman ρ és Kendall τ által meghatározott szórások értékei, $n = 4$ objektum összehasonlítása esetén $l = 0, 15$ -ös perturbáció alkalmazásával



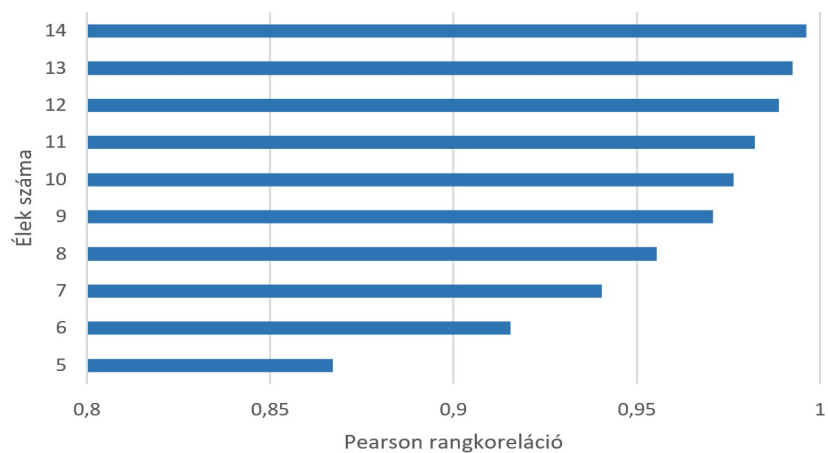
23. ábra. Az átlagos EU_M és EU_W távolságok a gráfok éleinek számának függvényében ($n = 6$)



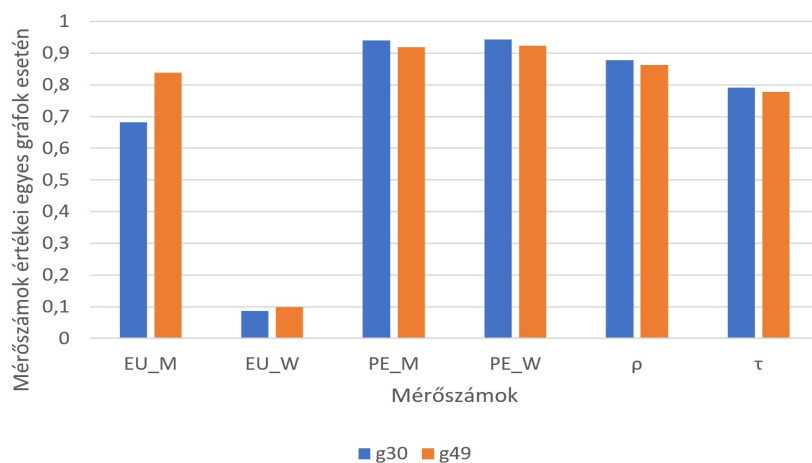
24. ábra. Átlagos rangkorrelációs (Spearman ρ és Kendall τ) értékek a különböző élszámú struktúrák esetén ($n = 6$)



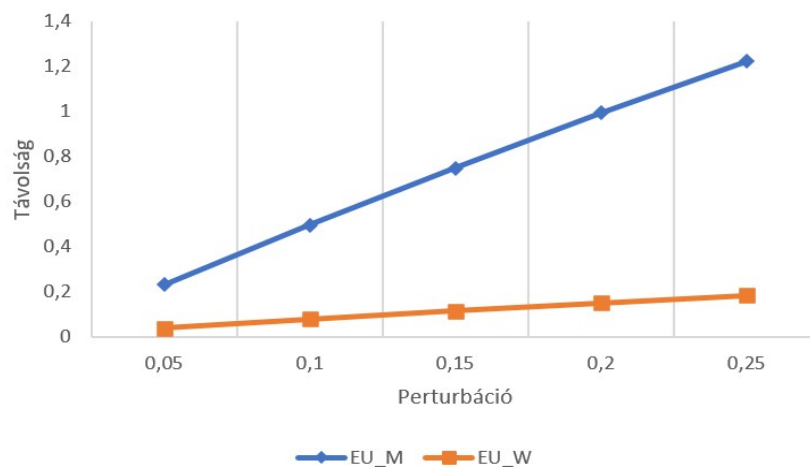
25. ábra. Az adott élszámú optimális struktúrák euklideszi távolsága $n = 6$ objektum esetén



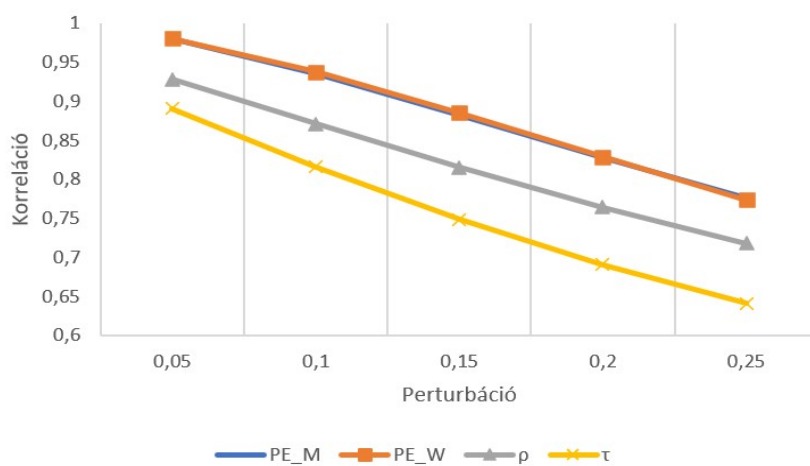
26. ábra. Pearson-korrelációs együtthatók az optimális gráfok és a teljes gráf között a gráfok élszámának függvényében



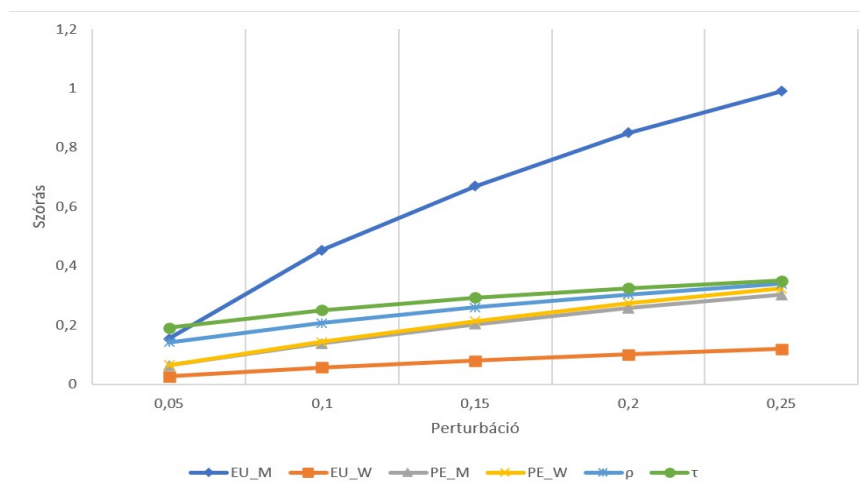
27. ábra. Példa arra, hogy a kevesebb összehasonlítás több információt eredményezhet: a 7 élű optimális struktúra, mint egy 8 élű struktúra ($n=6$ objektum és 0,15 perturbáció esetén).



28. ábra. EU_M és EU_W távolságok a csillaggráfhöz és a teljes gráfhöz tartozó eredmények között a perturbáció értékének függvényében



29. ábra. PE_M és PE_W, Spearman és Kendall rangkorrelációs értékek a csillaggráf és a teljes gráf esetén esetén a különböző perturbáció értékének függvényében (PE_M és PE_W szinte megegyezik)



30. ábra. EU_M, EU_W, PE_M, PE_W, ρ és τ értékek a csillaggráf és a teljes gráf esetén különböző perturbáció értékek mellett