

DIPLOMADOLGOZAT

Tulit Imre

2024



Magyar Agrár- és Élettudományi Egyetem

Kaposvári Campus

Vidékfejlesztés és Fenntartható Gazdaság Intézet

Pénzügy mesterképzési szak

**KOVARIANCIAMÁTRIX BECSLÉSI MÓDSZEREK VIZSGÁLATA
BEFEKTETÉSI STRATÉGIÁKBAN**

Belső konzulens:	Dr. Sipiczki Zoltán
	Egyetemi docens
Belső konzulens intézete/tanszéke:	Pénzügyi és Számviteli Tanszék
Külső konzulens:	Dr. Pál László
	Egyetemi docens
Készítette:	Tulit Imre

Kaposvár

2024

Kovarianciamátrix becslési módszerek vizsgálata befektetési stratégiákban

*A dolgozat áttekintő jellege miatt a szerkezete eltér a javasolt felépítéstől a MATE egységes szakdolgozat / diplomadolgozat / záródolgozat / portfólió készítési útmutatója 3.1 pontja alapján.

A portfólióoptimalizálás lehetővé teszi a befektetők számára, hogy maximalizálják hozamaikat, miközben minimalizálják a kockázatot. Ez a befektetések diverzifikálásával érhető el, különböző eszközosztályokba és szektorokba történő befektetéssel, ami csökkenti az egyedi kockázatokat. Az optimalizált portfólió segít a pénzügyi célok elérésében, függetlenül a piaci körülményektől.

Az optimalizálás során kiemelten fontos a kovarianciamátrix kiszámítása. A nagyméretű kovarianciamátrixok esetében a hagyományos módszerek nem mindig adnak jól kondicionált, invertálható mátrixot, ezért új, stabilabb becslési módszerekre van szükség. Ha sikerül olyan kovarianciamátrixot becsülni, amely pontosan tükrözi az eszközök közötti kapcsolatokat és stabil is, akkor jobb eredmények érhetők el.

Dolgozatomban különböző kovarianciamátrix-becslési módszerek teljesítményét fogom vizsgálni portfólióoptimalizálási technikákban.

Examining covariance matrix estimation methods in investment strategies

Portfolio optimization allows investors to maximize their returns while minimizing risk. This is achieved by diversifying investments across asset classes and sectors, which reduces individual risks. An optimized portfolio helps you achieve your financial goals, regardless of market conditions.

The calculation of the covariance matrix is of paramount importance in the optimization process. For large covariance matrices, traditional methods do not always provide a well-conditioned, invertible matrix, meaning that new, more robust estimation methods are needed. If a covariance matrix can be estimated that accurately reflects the relationships between the assets and is stable, better results can be obtained.

In my thesis, I will investigate the performance of different covariance matrix estimation methods in portfolio optimization techniques.

Tartalomjegyzék

1.	Bevezető	6
2.	Kovarianciamátrix becselő módszerek	7
2.1	Kovarianciamátrix értelmezése	7
2.2	Kovarianciamátrix becselő módszerek	8
2.2.1	Historikus	8
2.2.2	Ledoit-Wolf módszer	9
2.2.3	jLoGo módszer	10
2.2.4	Korrelációs mátrix zajcsökkentő (denoising) technikák	10
3.	Kovarianciamátrix becselő módszerek vizsgálata	13
3.2	Vizsgálat véletlenszerűen generált adatokon	14
3.2.1	Módszertan	14
3.2.2	Eredmények	16
3.3	Vizsgálat az S&P 500 részvényein	20
3.3.1	Módszertan	21
3.3.2	Eredmények	21
4.	Portfólió optimalizálási modellek	27
4.1	Markowitz optimalizálási modell	27
4.2	Legkisebb szórású modell	28
5.	Legkisebb szórású modell vizsgálata különböző kovarianciamátrix becselő módszerek mellett	28
5.1	Módszertan	29
5.2	Eredmények	33
6.	Következtetések	41
	Irodalomjegyzék	42

1. Bevezető

A portfólió optimalizálás az egyik legfontosabb pénzügyi stratégia, amely a befektetők számára lehetővé teszi, hogy maximalizálják a hozamokat, miközben minimalizálják a kockázatot. Ennek lényege a befektetések megfelelő diverzifikálása, vagyis különböző eszközosztályokba és szektorokba való befektetés, hogy csökkentsék az egyedi kockázatok hatását. Az optimalizált portfólió segít abban, hogy a befektetők elérjék pénzügyi céljaikat, függetlenül a piaci körülményektől.

Az optimalizálás során különösen fontos a kovarianciamátrix kiszámítása. Számos probléma olyan kovariancia-mátrix becslőt igényel, amely esetében a mátrix nemcsak invertálható, hanem jól kondicionált is vagyis invertálása nem növeli a becslési hibát. Nagyméretű kovarianciamátrixok esetén a szokásos mintakovarianciamátrix jellemzően nem jól kondicionált, és még az is előfordulhat, hogy nem is invertálható (pld. a részvények száma nagyobb, mint a megfigyelések száma). Ebből kifolyólag új módszerek bevezetésére volt szükség, amelyek képesek stabilabb kovarianciamátrix előállítására. Ha olyan kovarianciamátrixot tudunk becsülni, amely még mindig megragadja az eszközök közötti kapcsolatokat, és egyúttal stabilabb is, akkor optimálisabb eredményekre számíthatunk.

Dolgozatom során különböző kovarianciamátrix becslő módszerek teljesítményét fogom vizsgálni portfólió optimalizálási módszerekben.

2. Kovarianciamátrix becselő módszerek

Dolgozatom ezen részében először a kovarianciamátrix megértéséhez szükséges alapfogalmakról, majd a kovarianciamátrix értelmezéséről lesz szó, hogyan is néz ki egy kovarianciamátrix, illetve, hogy milyen tulajdonságai vannak, továbbá, hogy milyen módszerek vannak. Ezzel kapcsolatban tudni kell, hogy több módszer is rendelkezésre áll kovarianciamátrix becslésre, dolgozatom elkészítése során én négy módszert vizsgáltam, ezeket fogom bemutatni.

2.1 Kovarianciamátrix értelmezése

Még mielőtt rátérnék a kovarianciamátrixra, néhány alapfogalmat (Simon & Pál, 2019) ismertetnék, először szeretném a kovarianciával kezdeni. A statisztikában és a valószínűségszámításban a kovariancia két véletlen változó együttes változékonyságának mértékegysége. Ha a véletlen változók hasonló viselkedést mutatnak, akkor általában nagy a kovariancia közöttük. Matematikailag az X kovarianciáját az Y -hoz viszonyítva a (2.1) -ben látható módon fejezzük ki:

$$\text{COV}(X, Y) = E[(X - E[X])(Y - E[Y])] \quad (2.1)$$

Amit ezzel kapcsolatban még érdemes tudni, hogy ahogy azt az alábbi (2.2) képlet is mutatja, ha X kovarianciáját önmagával vesszük, akkor eredményként X varianciáját kapjuk, ami megmutatja, hogy egy valószínűségi változó milyen mértékben szóródik a várható érték (középérték) körül.

$$\text{COV}(X, X) = E[(X - E[X])(X - E[X])] = E[(X - E[X])^2] = \text{VAR}(X) \quad (2.2)$$

Továbbá a kovariancia másik tulajdonsága, hogy szimmetrikus, ami annyit jelent, hogy X és Y kovarianciájának azonos Y és X kovarianciájával, ahogy az az alábbi, (2.3) képletben is látható.

$$\text{COV}(X, Y) = \text{COV}(Y, X) \quad (2.3)$$

A kovarianciát használhatjuk a különböző eszközök közötti hasonlóságok számszerűsítésére. Ha két eszköznek magas a kovarianciája, akkor általában ugyanúgy fognak viselkedni. A különösen magas kovarianciájú eszközök lényegében helyettesíthetik egymást.

Most, hogy az alapvető dolgok tisztázásra kerültek, a továbbiakban rátérek a kovarianciamátrixra és annak fontos tulajdonságaira. A kovarianciamátrix képlete a (2.4) képletben látható módon írható fel:

$$\Sigma = \begin{bmatrix} \text{VAR}(X_1) & \cdots & \text{COV}(X_1, X_N) \\ \vdots & \ddots & \vdots \\ \text{COV}(X_N, X_1) & \cdots & \text{VAR}(X_N) \end{bmatrix} \quad (2.4)$$

Ahogy nő az eszközök száma, úgy nő a kapcsolatukat leíró kovariancia-mátrix dimenziója is. Ha N eszköz közötti kovarianciát vesszük, akkor egy $N \times N$ kovarianciamátrixot kapunk. Illetve ahogy az a (2.4)-es képletben is látható az i . átló az i . eszköz varianciája, a (i, j) és (j, i) értékek pedig a i eszköz és az j eszköz közötti kovarianciára vonatkoznak. Továbbá fontos megemlíteni, hogy a kovarianciamátrix pozitív szemidefinit, ami annyit jelent, hogy minden sajátértéke pozitív.

2.2 Kovarianciamátrix becselő módszerek

Ahogy azt korábban már említettem, most bemutatom az általam vizsgált négy darab kovarianciamátrix becselő módszert. Ezek pontosabban a *Historikus*, *Ledoit-Wolf*, *jLoGo* és *Fixed* módszer.

2.2.1 Historikus

Sorrendbe haladva előbb *Historikus* (Simon & Pál, 2019) módszerrel kezdem. Amit tudni kell, hogy a kovarianciamátrix becslése kritikus jelentőségűvé válik az erre épülő módszerek használatakor. A becslések stabilitása és pontossága alapvető fontosságú ahhoz, hogy stabil súlyokat kapjunk. Sajnos a kovariancia-mátrix becslésének legkézenfekvőbb módja, a mintakovariancia kiszámítása közismerten instabil. Ha kevesebb időbeli megfigyeléssel rendelkezünk az eszközeinkről, mint ahány eszközünk van a becslés különösen megbízhatatlanná válik. Ha olyan kovarianciamátrixot tudunk becsülni, amely még mindig megragadja az eszközök közötti kapcsolatokat, és egyúttal stabilabb is, akkor optimálisabb eredményekre számíthatunk. Ennek egyik fő kezelési módja, hogy valamilyen zsugorítási becslőt használunk.

2.2.2 Ledoit-Wolf módszer

A *Ledoit-Wolf* (Ledoit & Wolf, 2003; Ledoit & Wolf, 2004; Simon & Pál, 2019) a zsugorítás egy sajátos formája, ahol a zsugorítási együtthatót O. Ledoit és M. Wolf képlete alapján számítják ki. Számos probléma olyan kovariancia-mátrix becslőt igényel, amely nemcsak invertálható, hanem jól kondicionált is vagyis invertálása nem növeli a becslési hibát. Nagyméretű kovarianciamátrixok esetén a szokásos becslő - a mintakovarianciamátrix - jellemzően nem jól kondicionált, és még az is előfordulhat, hogy nem is invertálható. A lineáris zsugorodási becslő Ledoit és Wolf (2004) által javasolt és a mintakovariancia-mátrix és az azonossági mátrix aszimptotikusan optimális konvex lineáris kombinációja.

A zsugorítás a minta kovarianciamátrixának átalakítási eljárása, amelyet azért alkalmaznak, hogy egy robusztusabb kovarianciamátrixot kapjanak. A mátrix zsugorításának alapvető módja a minta kovarianciamátrix szélső értékeinek csökkentése, hogy közelebb húzzuk azokat a középponthez.

Adott egy kovarianciamátrix, $\mathbf{S}$, az átlagos variancia, $\boldsymbol{\mu}$ és a δ zsugorodási konstans, a zsugorított becslt kovariancia matematikailag a (2.5) képletben látható módon írható fel:

$$(1 - \delta)\mathbf{S} + \delta\boldsymbol{\mu}\boldsymbol{\mu}^T \quad (2.5)$$

A zsugorodási konstans (δ) értékét $0 \leq \delta \leq 1$ intervallumra van korlátozva, így egy súlyozott átlagot ad a minta kovariancia és az átlagos variancia mátrix között. *Ledoit Wolf* esetében az optimálisnak meghatározott zsugorodási konstans (δ) értékét a (2.6) képlettel határozták meg:

$$\hat{\delta}^* = \max \left\{ 0, \min \left\{ \frac{\hat{k}}{T'}, 1 \right\} \right\} \quad (2.6)$$

A képletben T' a múltbeli megfigyelések száma, a $\hat{k}$ pedig egy olyan konstans, ami nem ismert Ledoit és Wolf pedig ennek a becslését oldotta meg. Ennek részleteit a dolgozatban nem tárgyaljuk.

2.2.3 jLoGo módszer

A *jLoGo* (Barfuss et al., 2016) módszer információsűrő hálózatokat használ fel, hogy a hálózat kis alrészein kiszámított lokális inverz kovarianciák összegéből robusztus becslést készítsen. Mivel ez a módszer helyi és alacsony dimenziós inverziókon alapul, számítási szempontból nagyon hatékony és statisztikailag robusztus, még nagy dimenziójú, zajos és rövid idősorok inverz kovarianciájának becslésére is. A megoldás analitikus és ezért rendkívül gyors. Lokális jellege lehetővé teszi, hogy a számításokat párhuzamosan végezzük, és eszközt biztosít a dinamikus adaptációhoz részleges frissítéssel, amikor egyes változók tulajdonságai megváltoznak anélkül, hogy a teljes modell újraszámítására lenne szükség. Ez teszi különösen alkalmassá a nagyszámú változót tartalmazó nagy adathalmazok kezelésére.

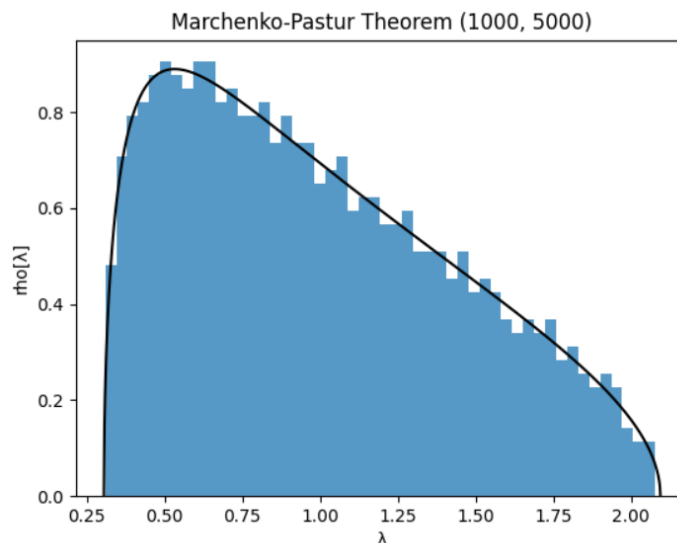
Pénzügyi adatokra alkalmazva a módszer eredményei számítási szempontból hatékonyabbak, mint a legmodernebb módszerek. Felhasználható például egy portfólió nyereség/veszteség eloszlásának kialakítására.

2.2.4 Korrelációs mátrix zajcsökkentő (denoising) technikák

Végül pedig a *Fixed* módszer következik. Először is amit tudni kell, hogy a kovarianciamátrix zajmentesítésének (Bun et al., 2017; López de Prado, 2016) fő gondolata a kovarianciamátrix azon sajátértékeinek eltávolítása, amelyek zajt képviselnek, és nem hasznos információt jelentenek. A pénzügyekben az empirikus korrelációs mátrixok becslését ismert módon befolyásolja a mérési hiba formájában jelentkező zaj, részben a számításukhoz jellemzően használt eszközhozamok idősorainak rövid hossza miatt. Ami még rosszabb, a nagy empirikus korrelációs mátrixokról kimutatták, hogy annyira zajosak, hogy a legnagyobb sajátértékek és a hozzájuk tartozó sajátvektorok kivételével lényegében véletlenszerűnek tekinthetők. Ezért az empirikus korrelációs mátrix használata előtt általában javasolt a zajcsökkentés.

A véletlen mátrixelmélet legfontosabb eredménye a Marchenko-Pastur-tétel (Portfoliooptimizer.io, 2022), amely leírja a nagy véletlen kovarianciamátrixok sajátérték-eloszlását. A Marchenko-Pastur-tétel megállapítja, hogy egy nagy véletlenszerű kovarianciamátrix sajátértékeinek eloszlása valójában univerzális, mivel az eloszlás független az alapul szolgáló megfigyelési mátrixtól. Ez az eloszlás segít azonosítani azokat a sajátértékeket, amelyek jelentősen eltérnek az elméleti eloszlástól, jelezve, hogy ezek a sajátértékek valószínűleg zajt képviselnek.

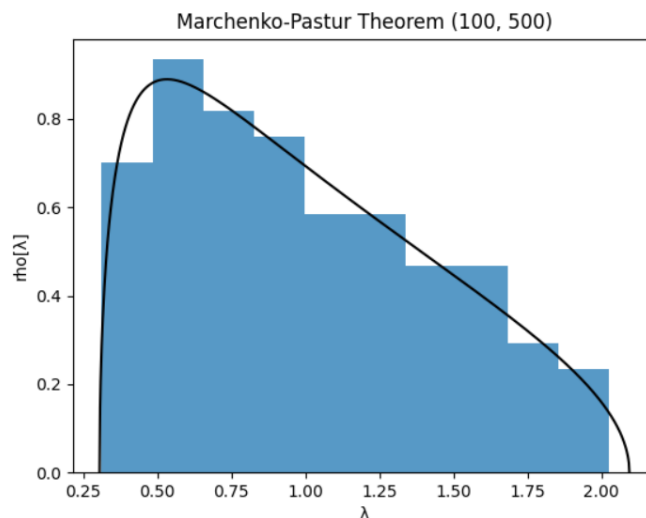
A Marchenko-Pastur-tétel szemléltetése az alábbi, 2.1. ábrán látható, az empirikus korrelációs mátrix sajátérték-sűrűsége $n = 1000$ és $T = 5000$ standard gauss-változókól álló véletlen megfigyelési mátrix.



2.1. ábra. Marchenko-Pastur $n = 1000$, $T = 5000$

(Forrás: <https://portfoliooptimizer.io/blog/correlation-matrices-denoising-results-from-random-matrix-theory/>)

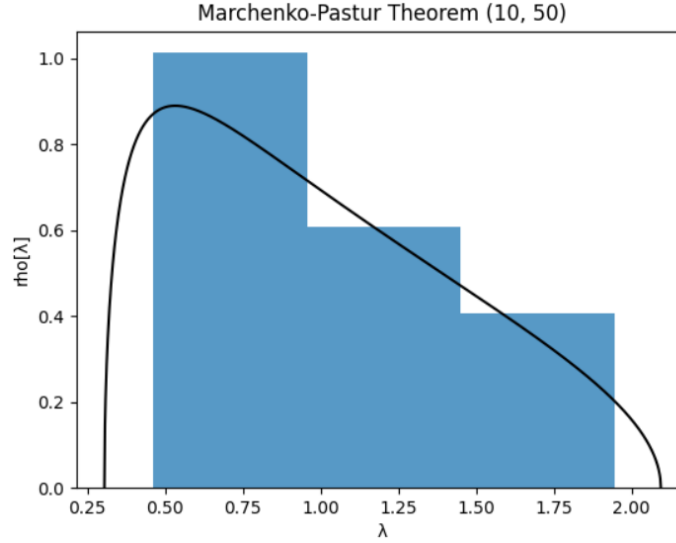
A következő 2.2. ábrán megvizsgálom, hogy mi történik akkor, ha n és T értékek kisebbek? Az n értékét 100-ra csökkentettem és T értéket pedig 500-ra.



2.2. ábra. Marchenko-Pastur $n = 100$, $T = 500$

(Forrás: <https://portfoliooptimizer.io/blog/correlation-matrices-denoising-results-from-random-matrix-theory/>)

A következő, 2.3. ábrán pedig egy még kisebb esetet vizsgálunk meg, ahol $n = 10$ és $T = 50$



2.3. ábra. Marchenko-Pastur $n = 10$, $T = 50$

(Forrás: <https://portfoliooptimizer.io/blog/correlation-matrices-denoising-results-from-random-matrix-theory/>)

A fenti ábrák alapján (2.1, 2.2, 2.3. ábra) megállapítható, hogy a Marchenko-Pastur-tétel alkalmazható, ha n és T értéke megközelítőleg 100, viszont nem árt az óvatosság abban az esetben, ha az értékük 10 körül van.

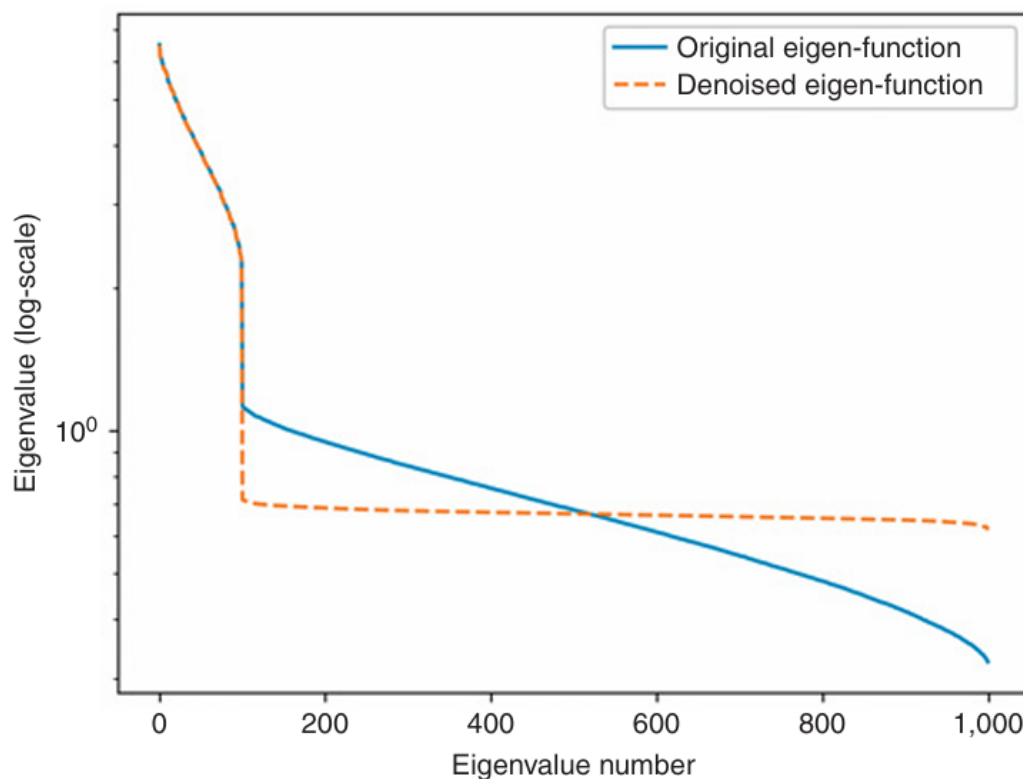
A zajmentesített mátrixok megbízhatóbb bemeneteket biztosítanak az optimalizáló algoritmusok számára, ami jobb eszközallokációt és kockázatkezelést eredményez.

A *Fixed* (López de Prado, 2020) módszer lényege, hogy minden véletlenszerű sajátvektorhoz állandó sajátértéket állítunk be. $\{\lambda_n\}_{n=1, \dots, N}$ az összes sajátérték halmaza, csökkenő sorrendben, és i a sajátérték pozíciója, úgy, hogy $\lambda_i > \lambda_{i+1}$ és $\lambda_{i+1} \leq \lambda_{i+2}$. Ezután beállítjuk $\lambda_i = 1/(N - i) \sum_{k=i+1}^N \lambda_k$, $j = i + 1, \dots, N$, így megmarad a korrelációs mátrix nyomvonala. A $VW = W\Lambda$ sajátvektor-dekompozíciót tekintve a C_1 zajcsökkentett korrelációs mátrixot a következőképpen alakítjuk ki ((2.7). és (2.8) képlet):

$$\tilde{C}_1 = W\tilde{\Lambda}W' \quad (2.7)$$

$$C_1 = \tilde{C}_1 \left[(\text{diag}[\tilde{C}_1])^{\frac{1}{2}} (\text{diag}[\tilde{C}_1])^{\frac{1}{2}} \right]^{-1} \quad (2.8)$$

A képletben $\tilde{\Lambda}$ az a diagonális mátrix, amely a korrigált sajátértékeket tartalmazza, illetve az aposztróf (') transzponálja a mátrixot és a `diag[.]` nullázza a négyzetes mátrix összes nem diagonális elemét. A második transzformáció oka a C1 mátrix átskálázása, hogy a C1 mátrix fő átlója egyeseket tartalmazzon. Az alábbi, 2.4. ábra összehasonlítja a sajátértékek logaritmusát a történő zajmentesítés előtt és után.



2.4. ábra. Sajátértékek zajcsökkentés előtt és után

(Forrás: Marcos M. López de Prado, 2020, Machine Learning for Asset Managers, 30)

3. Kovarianciamátrix becslő módszerek vizsgálata

Most, hogy a kovarianciamátrixal és a módszerekkel kapcsolatos elméleti részek tisztázva lettek, a gyakorlati rész következik. Dolgozatom ezen részében először véletlenszerűen generált adatokon vizsgálom a módszereket, majd ezután pedig valós adatokon, pontosabban az S&P 500 részvényeinek adatait fogom felhasználni.

3.2 Vizsgálat véletlenszerűen generált adatokon

Ahogy azt korábban említettem a módszerek vizsgálata következik véletlenszerűen generált adatokon, először bemutatom a módszertant, azon belül, hogy hogyan határoztam meg az eredeti kovarianciamátrixot, milyen mátrix és adathalmaz méreteket vizsgáltam, illetve milyen metrikákat vettem figyelembe.

3.2.1 Módszertan

A vizsgálat során két fontos paramétert használok, az egyik a mátrix mérete, a másik a generált adathalmaz mérete. Összesen öt különböző mátrix méret és generált adat paramétert használtam fel, minden esetben az első érték a mátrix méretét a másik az adathalmaz méretét jelenti, vagyis a következőképpen néz ki: 30 és 35, 50 és 60, 100 és 120, 200 és 240, 400 és 480. A két paraméter közötti eltérés alacsony, viszont minden esetben a mátrix mérete kisebb, mint a generált adathalmaz mérete.

A mátrix mérete az eredeti mátrix meghatározásához szükséges, ami, egy véletlenszerűen generált ortogonális mátrix és egy meghatározott sajátértékkel (1, 3, és 10 arányban) rendelkező véletlenszerűen generált diagonális mátrix segítségével került meghatározásra. Ezzel kapcsolatban érdemes tudni, hogy egy valós számokkal vagy elemekkel rendelkező négyzetmátrixot ortogonális mátrixnak nevezünk, ha a transzponálása egyenlő az inverz mátrixával. Vagy azt is mondhatjuk, hogy ha egy négyzetmátrix és transzponálásának szorzata identitásmátrixot ad, akkor a négyzetmátrixot ortogonális mátrixnak nevezzük. A diagonális mátrixszal kapcsolatban pedig érdemes tudni, hogy a fő átlós elemek kivételével minden elem nulla (Chen, 1997; Horn & Johnson, 2013).

A meghatározott eredeti mátrix és a megadott adathalmaz méret felhasználva, generálok egy véletlenszerű normális eloszlású adathalmazt, amelyet a kovarianciamátrix számításához használok fel a továbbiakban.

Az eredeti és a becsült kovarianciamátrix alapján pedig kiszámítom a hibákat, ebben az esetben olyan metrikákat számolok, mint az átlagos abszolút hiba ("Mean absolute error" - továbbiakban MAE), átlagos négyzetes hiba ("Mean squared error" – továbbiakban MSE), Minimum variance loss (továbbiakban min_var_loss) és Stein's Loss.

A továbbiakban kicsit bővebben kifejteném ezt a négy metrikát, kezdve az átlagos abszolút hibával. Az **átlagos abszolút hiba (MAE)** (Willmott & Matsuura, 2005) egy egyszerű, de hatékony mérőszám, amelyet a regressziós modellek pontosságának értékelésére használnak. Az előrejelzett értékek és a tényleges célértékek közötti átlagos abszolút különbséget méri. Más metrikáktól eltérően a MAE nem emeli négyzetre a hibákat, ami azt jelenti, hogy minden hibának azonos súlyt ad. Képlete alább, a (3.1) képletnél látható:

$$MAE = \frac{\sum_{i=1}^n |y_i - x_i|}{n} \quad (3.1)$$

A képletben n az adatok számát, y_i az adatok tényleges értékét, x_i pedig az adatok becsült értékét jelenti. A MAE számos előnnyel rendelkezik, a legfontosabb, amit kiemelnék, hogy kevésbé érzékeny az adatok szélsőséges értékeire (kiugró értékekre), illetve, hogy egyszerű kiszámítani és megérteni.

Az **átlagos négyzetes hiba (MSE)** (Willmott & Matsuura, 2005) esetében a megfigyelések és az előre jelzett értékek összehasonlításakor a különbségeket négyzetre kell emelni, mivel egyes adatértékek nagyobbak lesznek, mint az előrejelzés (és így a különbségük pozitív lesz), mások pedig kisebbek (és így a különbségük negatív lesz). Mivel a megfigyelések ugyanolyan valószínűséggel lesznek nagyobbak, mint az előre jelzett értékek, mint amilyen valószínűséggel kisebbek, a különbségek összege nulla lesz. A különbségek négyzetre emelése kiküszöböli ezt a helyzetet. Az átlagos négyzetes hiba képlete (3.2) a következő.

$$MSE = \frac{\sum_{i=1}^n (y_i - x_i)^2}{n} \quad (3.2)$$

A képletben n az adatok számát, y_i az adatok tényleges értékét, x_i pedig az adatok becsült értékét jelenti. Ha az előrejelzés minden adatponton áthalad, akkor az átlagos négyzetes hiba nulla.

A **minimum variance loss** (Ledoit & Wolf, 2020; Engle et al., 2019) egy olyan metrika, amelyet a kovarianciamátrix becslők hasznosságának mérésére használnak. Az ilyen típusú becslőket gyakran használják az eredeti változók olyan kombinációinak megtalálására, amelyek minimális szórással rendelkeznek egy lineáris megkötés mellett. A kovarianciamátrix-becslő minőségét ezután az eredeti változók lineáris kombinációjának valódi varianciájával mérjük: az alacsonyabb variancia jobb. Ezen az alapon Engle, Ledoit és Wolf ((2019), 1. definíció) javasolta az alábbi (3.3) hibafüggvényt:

$$\mathcal{L}_n^{MV}(\hat{\Sigma}_n, \Sigma_n) := \frac{\frac{Tr(\hat{\Sigma}_n^{-1} \Sigma_n \hat{\Sigma}_n^{-1})}{p}}{\left[\frac{Tr(\hat{\Sigma}_n^{-1})}{p} \right]^2} - \frac{1}{\frac{Tr(\Sigma_n^{-1})}{p}}, \quad (3.3)$$

ahol $Tr(\cdot)$ egy négyzetes mátrix nyoma, n minta mérete, p a kovarianciamátrix mérete. Tehát kiszámítja a minimális variancia hibafüggvényt a becsült kovarianciamátrix és a minta-kovarianciamátrix között.

A **Stein's Loss** (Ledoit & Wolf, 2017) a kovariancia-mátrix becslésével összefüggésben használt mérték, különösen nagy dimenziós adatok esetén. Ez egy skála-invariáns vesztességfüggvény, amely a becsült kovarianciamátrix teljesítményét értékeli a valódi kovarianciamátrixhoz képest. A vesztesség függvény célja, hogy mérje a becsült kovarianciamátrix és a valódi kovarianciamátrix közötti különbséget. Képlete (3.4) a következőképpen írható fel:

$$\mathcal{L}_n^S(\Sigma_n, \tilde{\Sigma}_n) := \frac{1}{p} Tr(\Sigma_n^{-1} \tilde{\Sigma}_n) - \frac{1}{p} \log \det(\Sigma_n^{-1} \tilde{\Sigma}_n) - 1 \quad (3.4)$$

A képletben a Σ az eredeti mátrixot S pedig a becsült mátrixot jelöli, $Tr(-)$ pedig a mátrix nyomvonalát.

Összességében ezt a négy metrikát használtam fel a módszerek vizsgálatához, minden mátrix méret, generált adat méret páros esetében 50 alkalommal futtattam és minden metrikának ebből az ötven futtatásból számoltam átlagot. Továbbá vizsgáltam a kovarianciamátrix becselő módszerek futási idejét is.

3.2.2 Eredmények

A Dolgozat ezen részében az eredményeket fogom vizsgálni. Az előbb említett mátrix méret és generált adathalmaz kombinációk esetében számolt metrikák eredményei az alábbi 3.1-es táblázaton láthatók.

Cov Size	Data Size	Method	MAE	MSE	Min_var_loss	Stein_loss
30	35	Ledoit	0.453	0.324	1.801	0.372
		JLogo	0.705	0.802	3.097	0.595
		Fixed	0.663	0.702	2.845	0.575
		Hist	0.749	0.914	16.931	0.749
50	60	Ledoit	0.345	0.187	1.769	0.372
		JLogo	0.511	0.42	2.868	0.594
		Fixed	0.512	0.415	2.721	0.562
		Hist	0.563	0.508	15.565	0.683
100	120	Ledoit	0.242	0.092	1.764	0.38
		JLogo	0.336	0.182	2.811	0.627
		Fixed	0.368	0.214	2.715	0.571
		Hist	0.395	0.248	15.338	0.658
200	240	Ledoit	0.171	0.046	1.763	0.382
		JLogo	0.227	0.082	2.786	0.649
		Fixed	0.286	0.129	2.636	0.463
		Hist	0.279	0.123	14.446	0.653
400	480	Ledoit	0.121	0.023	1.76	0.384
		JLogo	0.156	0.039	2.771	0.666
		Fixed	0.203	0.065	2.626	0.457
		Hist	0.197	0.061	13.805	0.646

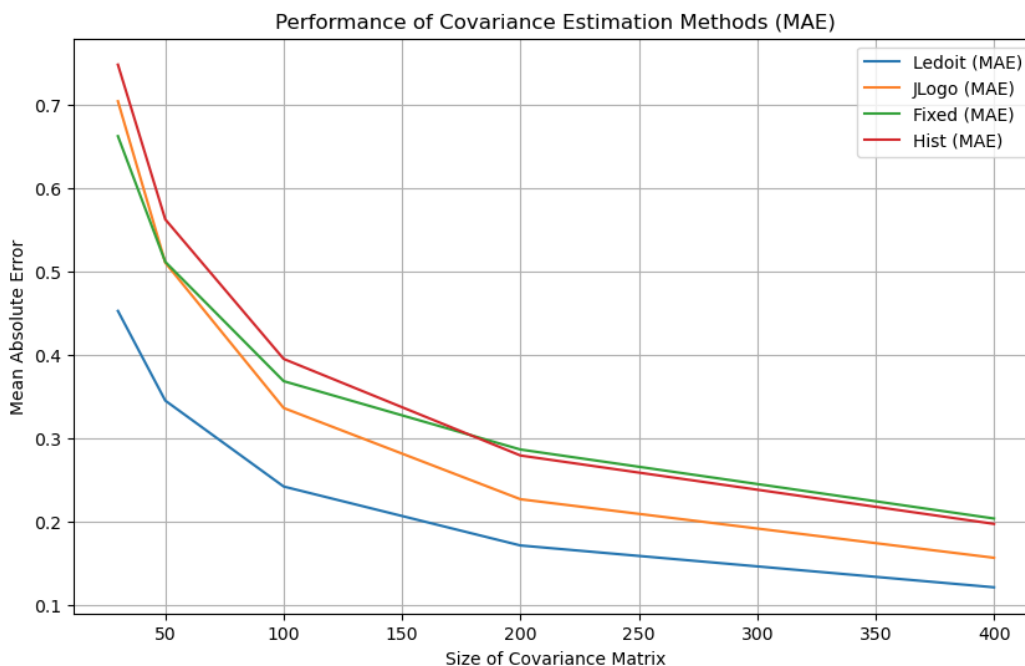
3.1. táblázat. Mátrix méret és generált adathalmaz kombinációk eredményei (MAE, MSE, Minimum variance loss, Stein loss)

Az 3.1. táblázat alapján, amit elsőként megállapíthatunk, hogy a *Ledoit* módszer minden esetben minden metrikát tekintve jobb teljesített a többinél. MAE és MSE tekintetében megfigyelhető, hogy ahogy a kovarianciamátrix mérete és az adatok mérete növekszik úgy nagy mértékben csökken a MAE és az MSE értéke is minden esetben. Min_var_loss esetében is csökkenés látható viszont már kisebb mértékű, a Stein loss esetében pedig változó a *Ledoit* és a *JLoGo* esetében növekedik, a *Fixed* és a *Hist* pedig csökken. Továbbá érdemes megfigyelni, hogy *Ledoit* alacsonyabb kovarianciamátrix és adat méret mellett is jól teljesít, és persze minden esetben a legjobb azonban kovarianciamátrix és adat méret növekedésével a módszerek közötti eltérés csökken.

A táblázat egészét tekintve ahogy már említettem a *Ledoit* módszer minden metrika szerint a legjobb teljesítményt nyújtja, különösen nagyobb méretű adatok és kovarianciamátrixok esetében. A *JLoGo* és *Fixed* módszerek közepes teljesítményt mutatnak, de általában rosszabbak, mint a *Ledoit*. A *Hist* módszer, ahogy az várható is volt a legrosszabb teljesítményt nyújtja minden

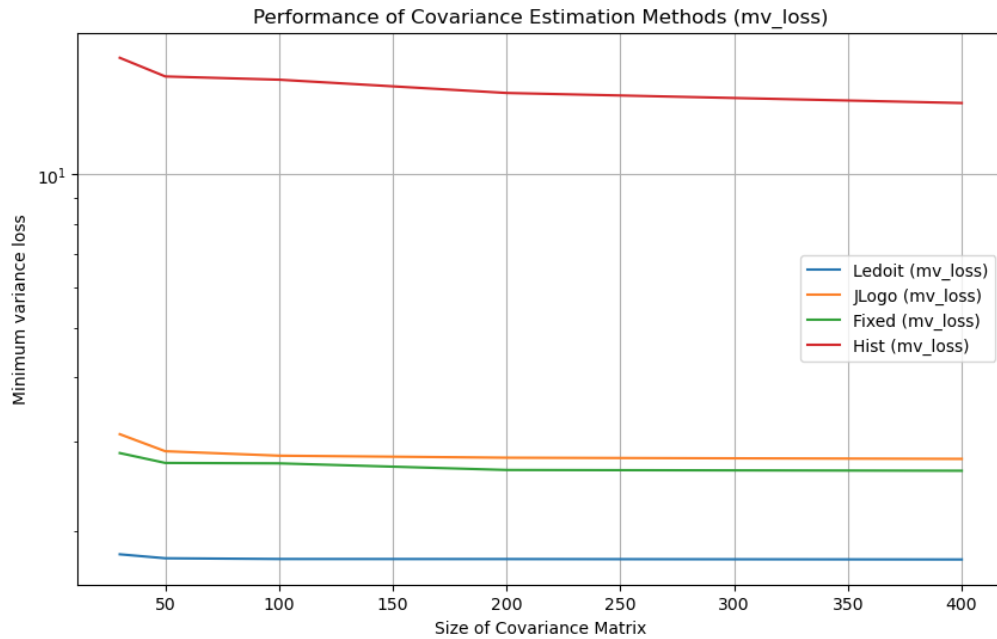
esetben, különösen a `min_var_loss` tekintetében. Tehát ez alapján az eredmények alapján a *Ledoit* módszer a legmegfelelőbb választás a kovarianciamátrix becslésére.

A továbbiakban a MAE és a `min_var_loss`-ról láthatók ábrák, illetve a módszertan ismertetése során említett átlagos futási idő vizsgálatáról is itt láthatunk majd egy diagramot.



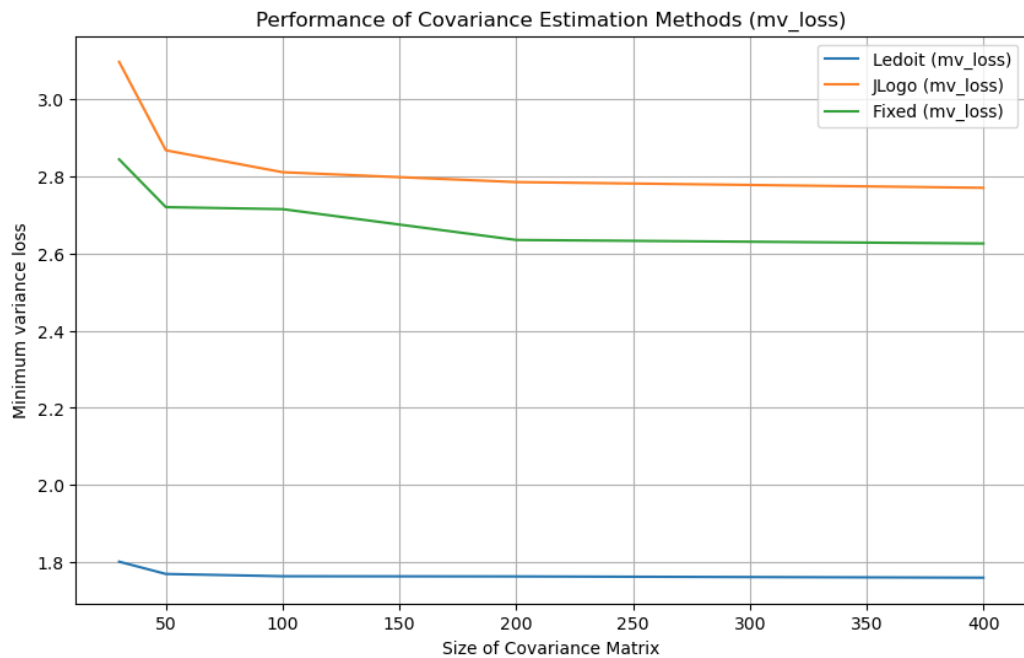
3.1. ábra. Vizsgálat véletlenszerűen generált adatokon – MAE eredmény

A MAE diagramot megtekintve (3.1. ábra) is láthatjuk, hogy a *Ledoit* módszer a legjobb, illetve az elején még a *Fixed* a második legjobb majd, ahogy növekszik a kovarianciamátrix mérete a *Fixed* egyre rosszabban teljesít és a második helyet a *JLoGo* módszer veszi át. Illetve a diagramon még jobban látható, hogy a kovarianciamátrix és az adatok méretének növekedésével a hiba minden módszer esetében csökken és a módszerek közti különbségek is csökkennek.



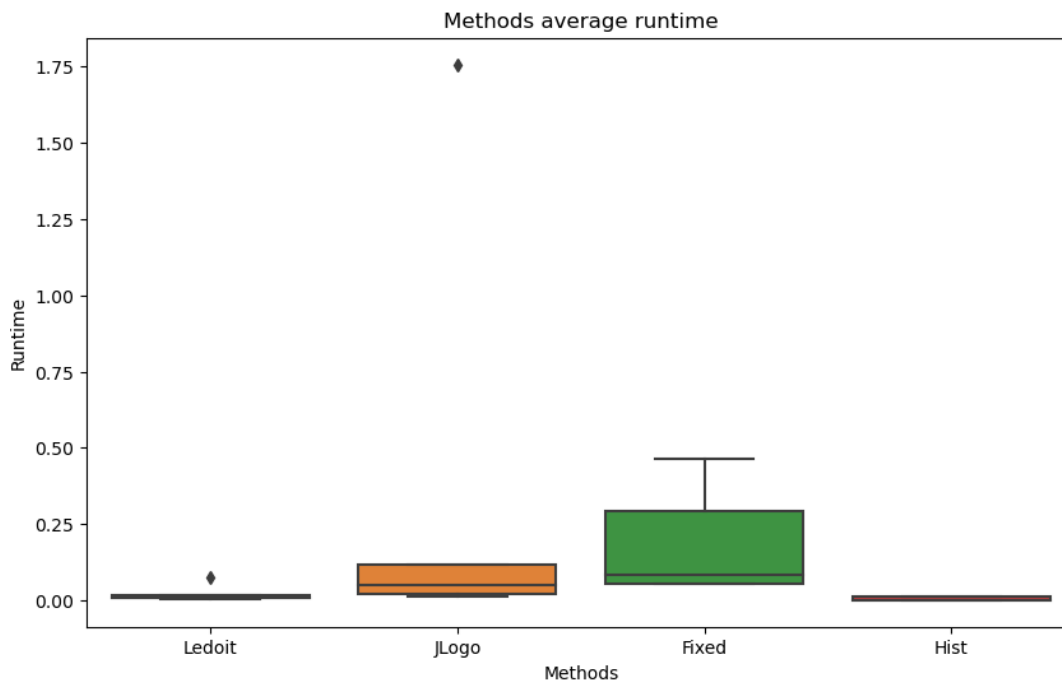
3.2. ábra. Vizsgálat véletlenszerűen generált adatokon – Minimum variance loss eredmény

A fenti 3.2. ábrán `min_var_loss` tekintetében vannak hasonlítva a módszerek, itt érdemes megfigyelni, illetve nagyon jól látható, hogy ha bár a *Ledoit* a legjobb, a másik két módszer is mennyivel jobban teljesít, mint a *Hist* módszer. Tehát jól látható a *Hist* gyenge teljesítménye és egyértelműen látható a kovarianciamátrix becslő módszerek szükségessége.



3.3. ábra. Vizsgálat véletlenszerűen generált adatokon – Minimum variance loss eredmény *Hist* nélkül

Készítettem egy külön ábrát is ahol csak a három módszer látható (3.3. ábra), hogy egyértelműek legyenek a különbségek, a *Ledoit* módszert követve szorosan egymás után a *jLoGo* és a *Fixed* módszer következik.



3.4. ábra. Vizsgálat véletlenszerűen generált adatokon – Módszerek átlagos futási ideje

Végül pedig ahogy a fenti 3.4. ábrán is látható, futási idő tekintetében hasonlítottam össze a módszereket, az ábrán az értékek másodpercben vannak kifejezve. Megfigyelhető, hogy ha bár a *Hist* gyenge eredményeket ad viszont nem mondható lassúnak. A *Ledoit* pedig nem csak, hogy jól teljesít, viszont gyors is, őt követi a *jLoGo* és végül futási idő szempontjából a leglassabb a *fixed*. Tehát futási idő szempontjából is a *Ledoit* képes kiemelkedő teljesítményt nyújtani.

Összefoglalva, ahogy már azt többször is említettem, minden általam figyelembe vett metrikát tekintve a *Hist*, *Ledoit*, *jLoGo* és *Fixed* módszerek közül a legjobb teljesítményt a *Ledoit* módszer képes nyújtani.

3.3 Vizsgálat az S&P 500 részvényein

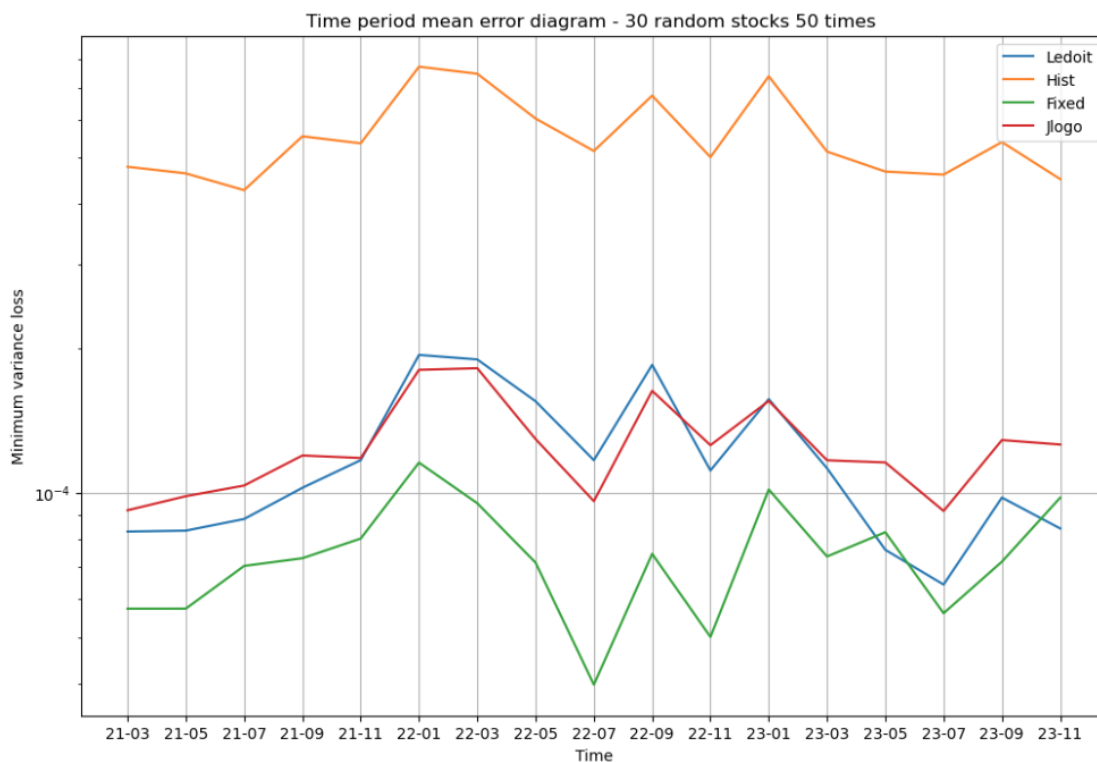
Dolgozatom ezen részében a módszerek vizsgálata következik az S&P 500 részvényein, a felépítés ugyanaz, mint az előbb, először ismertettem a módszertant, majd bemutatom az eredményeket.

3.3.1 Módszertan

A vizsgálat során három év adatait vettem figyelembe, tehát a vizsgált periódus 2021-2023. A részvények kiválasztása az S&P 500 részvényei közül véletlenszerűen történt. Három esetet vizsgáltam, 30 részvényt 2 hónapnyi adattal, 50 részvényt 3 hónapnyi adattal és végül pedig 100 részvényt 6 hónapnyi adattal. A vizsgálat úgy történt, hogy párokra bontottam a számolt kovarianciamátrixokat, vagyis az aktuális hónap eredményét vettem a becsült mátrixnak és az aktuális hónap előtti hónapot tekintettem az eredeti mátrixnak és ezekkel számoltam ki a metrikákat. Ez a módszer során a MAE és a min_var_loss metrikákat vettem figyelembe, ezeket a dolgozatom korábbi részében már ismertettem.

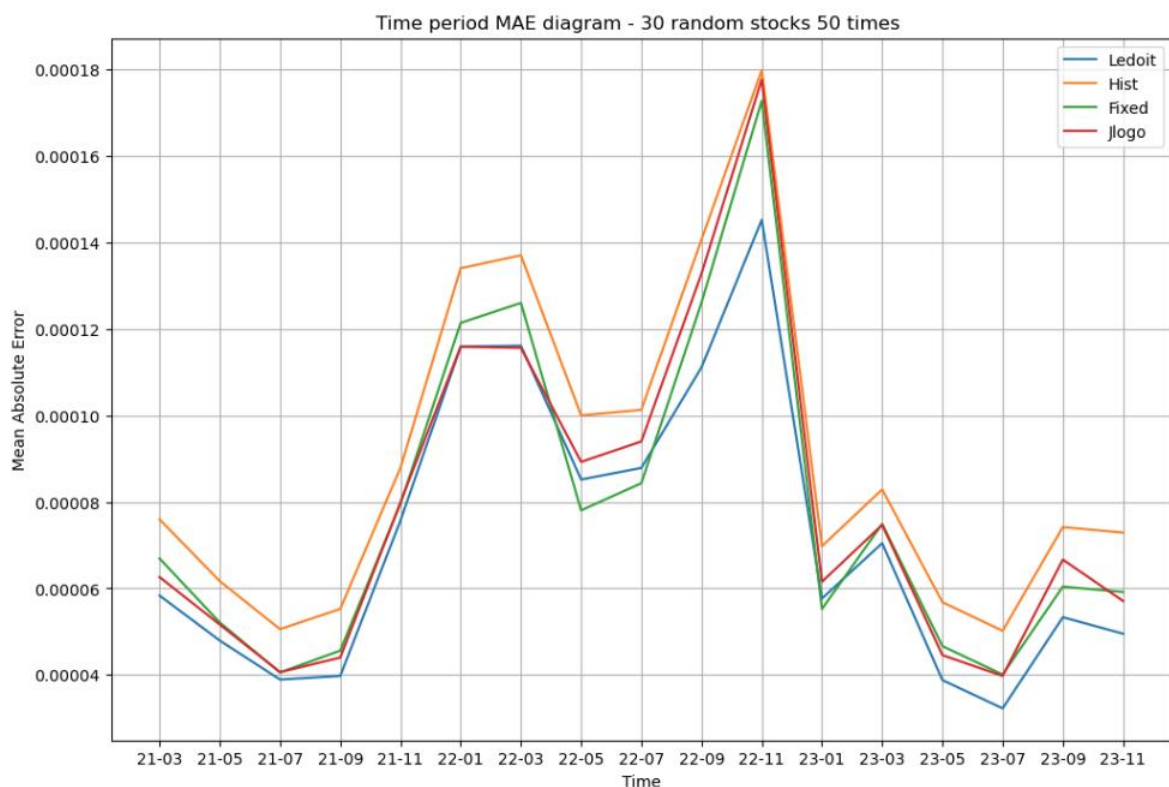
3.3.2 Eredmények

Most, hogy a módszertant ismertettem az eredmények ismertetése következik, mindhárom esetben a minimum variance loss-t és a MAE-t vizsgáltam meg. Elsőkén sorba haladva a 30 részvény, 2 hónapnyi adattal eset eredményeit mutatom be.



3.5. ábra. Minimum variance loss - 30 részvény, 2 hónapnyi adattal, 50 alkalommal

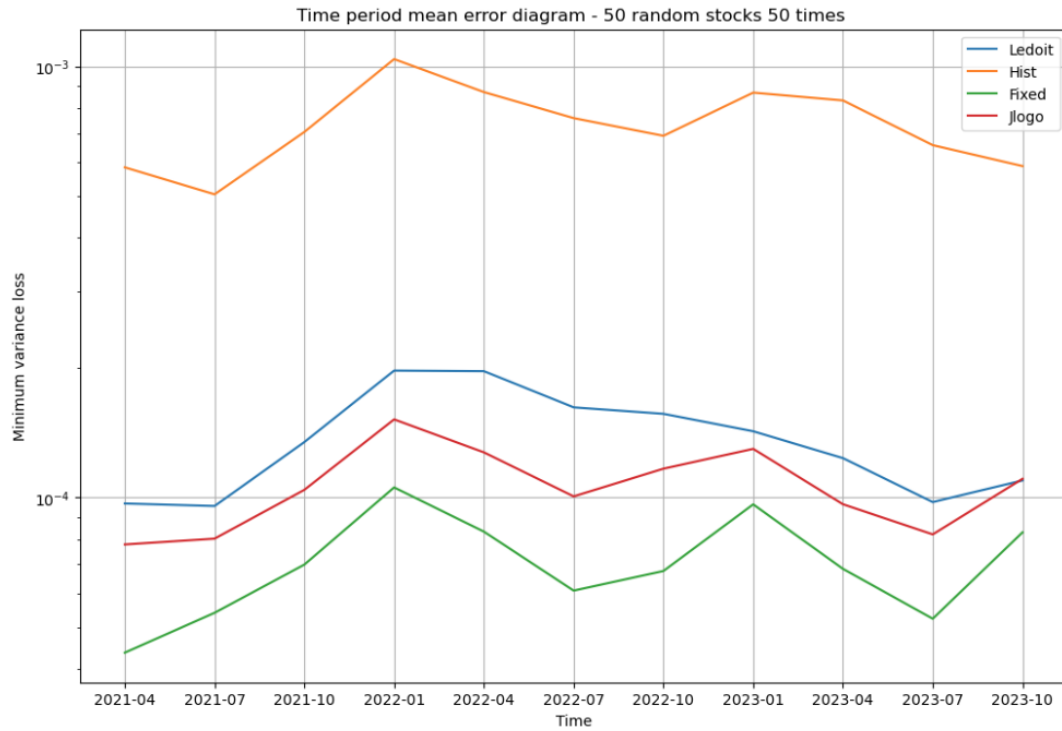
Ahogy az a fenti 3.5. ábrán is látható *min_var_loss*-t tekintve, a *Hist* módszer továbbra is, ahogyan arra számítani is lehetett a legrosszabb, a korábbi random adatokkal ellentétben, ahol a *Ledoit* módszer volt a legjobb, ebben az esetben az látható, hogy 1-2 alkalommal jobb az összesnél viszont ennél a vizsgálatnál a *Fixed* módszer érte el a legjobb eredményt. Ez az eltérés abból is adódhat, hogy a vizsgálat során eredeti mátrixnak tekintett érték mindig egy előző periódusból származó becsült kovarianciamátrix, így nem tud annyira pontos eredményt adni, mint egy olyan vizsgálat, ahol ismert az eredeti mátrix.



3.6. ábra. MAE - 30 részvény, 2 hónapnyi adattal, 50 alkalommal

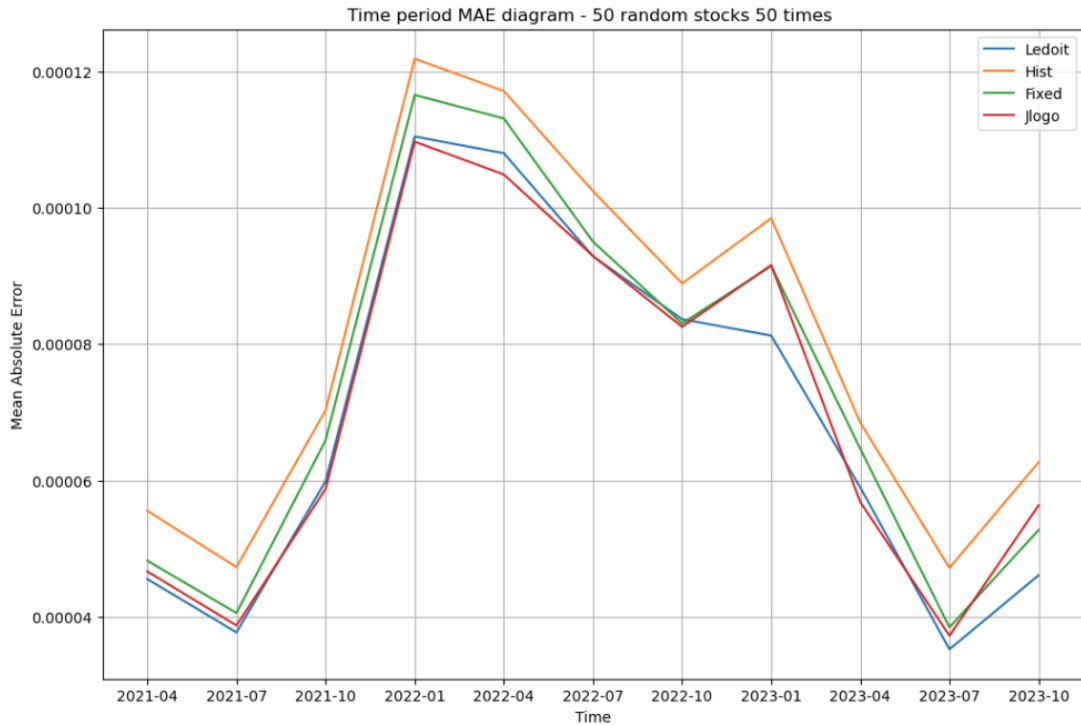
Ahogy említettem, ebben az esetben megnéztem a MAE-t is (3.6. ábra) ahol hasonlóan a random generált adatoknál való vizsgálat során itt is elmondható, hogy a *Ledoit* módszer teljesített a legjobban ugyanis a vizsgált periódusban ennek az esetében volt a leggyakoribb, hogy a legkisebb hibával rendelkezett. Továbbra megfigyelhető, hogy itt is a *Hist* a leggyengébb, illetve, hogy a három másik módszer közül bár a *Ledoit* a legjobb az eltérések nagyon minimálisak.

Most, hogy kielemeztük a 30 részvény, 2 hónapnyi adatra vonatkozó eset eredményeit, következhet az 50 részvény, 3 hónapnyi adattal történt vizsgálat eredménye.



3.7. ábra. Minimum variance loss - 50 részvény, 3 hónapnyi adattal, 50 alkalommal

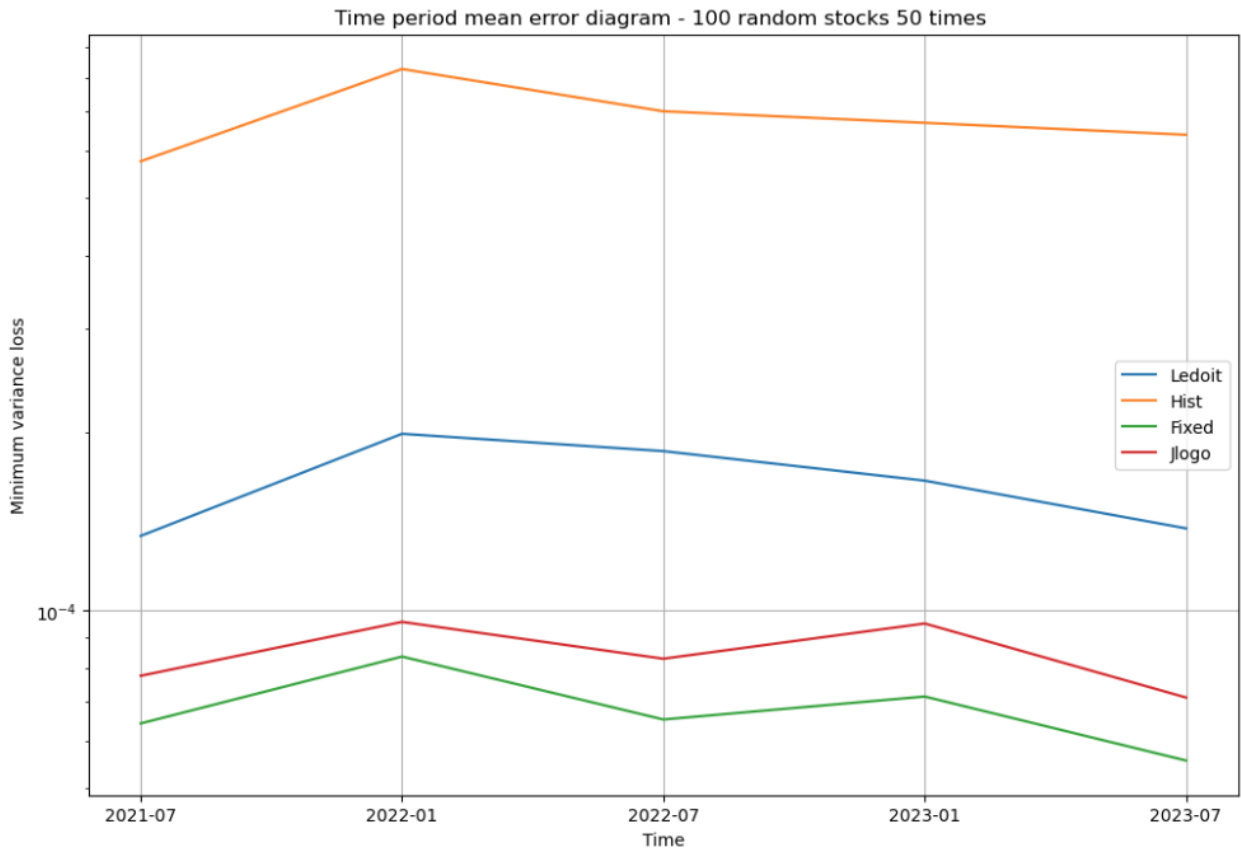
Ahogy az a fenti 3.7. ábrán is látható, a részvények számának növekedésével, illetve a periódus növelésével min_var_loss szempontjából a sorrend meg inkább egyértelművé vált, vagyis ez a vizsgálat során a legjobb eredményt a *fixed* adja, második a *JLoGo*, illetve harmadik a *Ledoit*.



3.8. ábra. MAE - 50 részvény, 3 hónapnyi adattal, 50 alkalommal

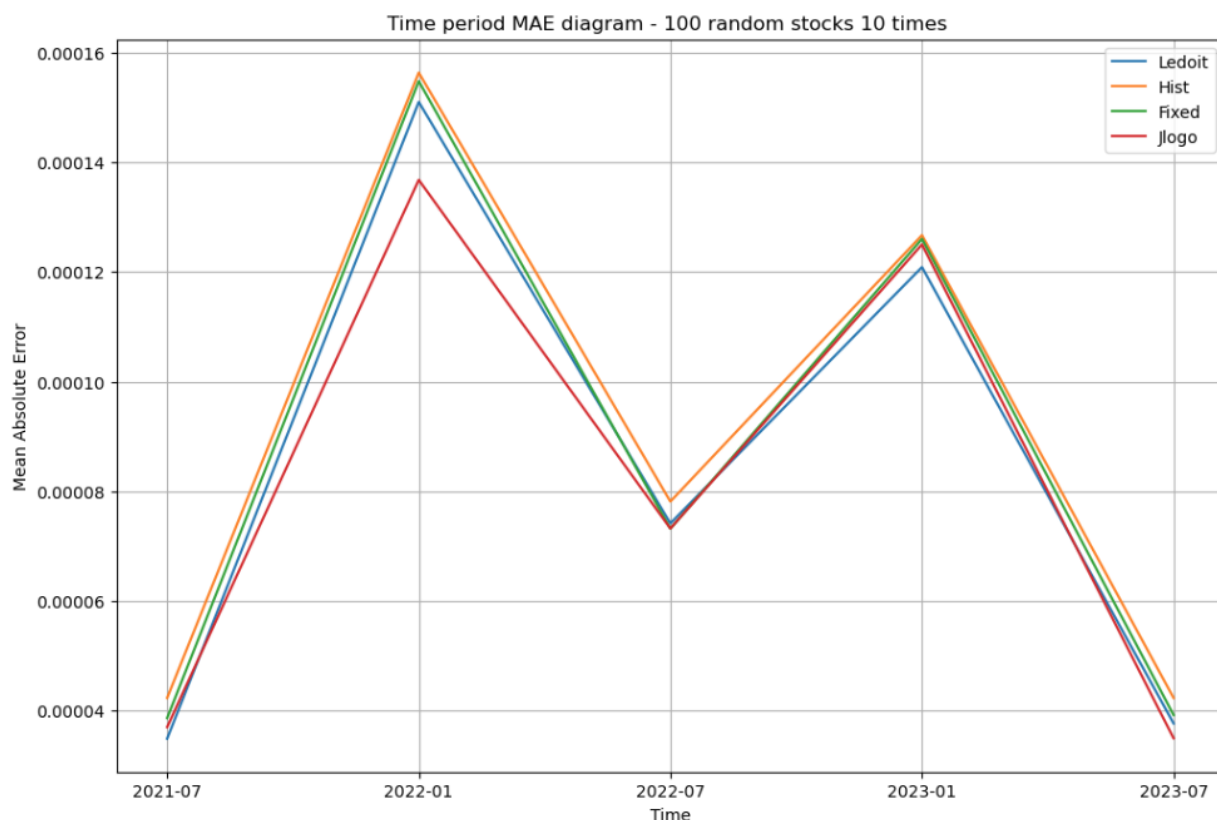
Továbbá ebben az esetben is megnéztem a MAE-t (3.8. ábra), ebben a tekintetben a *Ledoit* a jobb, ugyanis a periódus egészét tekintve, habár szorosan együtt mozog a *JLoGo* módszerrel, nagyobb részben tud alacsonyabb hiba értéket elérni, vagyis több időtartamon keresztül jobban teljesít.

Ahogy azt már korábban említettem van még egy eset, amit vizsgáltam, ennek során a részvények számát 100-ra növeltem és 6 hónapnyi adattal dolgoztam, a továbbiakban az itt kapott eredményeket mutatom be.



3.9. ábra. Minimum variance loss - 100 részvény, 6 hónapnyi adattal, 50 alkalommal

Ahogy a fenti 3.9. ábrán is látható, a részvények és a hónapok számának nevelkedésével az eltérések meg inkább láthatók lettek, a *Hist* továbbra is a leggyengébb, ahogy az várható is volt, viszont amit érdemes megfigyelni, hogy a *JLoGo* és a *Fixed* továbbra is együtt mozog, viszont a *Ledoit* esetében nagyobb lett az eltávolodás az előző esetekhez képest, tehát továbbra is ebben a kísérletben a *Fixed* teljesít a legjobban.



3.10. ábra. MAE - 100 részvény, 6 hónapnyi adattal, 50 alkalommal

Végül pedig ebben az esetben is megvizsgáltam a MAE-t (3.10. ábra), ebben az esetben már megoszlik majdnem 50-50% arányban, hogy a *Ledoit* vagy a *JLoGo* módszer ér el jobb eredményt, látható az ábrán, hogy a vizsgált időszak feléig a *JLoGo* teljesít jobban, majd onnantól kezdve a *Ledoit* vagy ha nem is a két módszer nagyon azonos hiba értéket ér el.

Összességében a három esetet megvizsgálva arra a következtetésre jutottam, hogy a *Fixed* módszer a *min_var_loss* szempontjából jobban teljesített minden esetben, különösen a nagyobb minták és hosszabb időszakok esetében volt feltűnően látható, míg a MAE szempontjából a *Ledoit* módszer bizonyult a legjobbnak. A *Ledoit* módszer kisebb hibával dolgozott a vizsgált időszakok többségében. A *Hist* módszer pedig minden esetben az elvárt eredményt hozta, vagyis a legrosszabb eredményeket érte el minden esetben. A *min_var_loss*-t tekintve az előző kísérlethez képest meglepő a *Fixed* teljesítménye.

4. Portfólió optimalizálási modellek

A portfólióoptimalizálás (Pál, 2022) az egyes eszközök (pl. egyes részvények, eszközosztályok, kötvények, készpénz) közötti lehető legjobb allokáció kiválasztásának folyamata valamilyen konkrét cél, például a kockázat minimalizálása, a hozam, a kockázattal korrigált hozam, a diverzifikáció maximalizálása stb. alapján. Az eszközallokáció általában kihívást jelentő feladat, mivel számos tényező befolyásolhatja az eredményeket. Az eredményt befolyásolhatja a befektetési időszak és az eszközuniverzum, a befektető kockázattűrő képessége és így tovább.

Dolgozatom ezen részében néhány portfólió optimalizálási módszer bemutatása következik, pontosabban két módszer, a Markowitz optimalizálási modell, és a legkisebb szórású modell.

4.1 Markowitz optimalizálási modell

Elsőként a Markowitz optimalizálási modellel (Pál, 2022; Markowitz 1952) kezdem, amit tudni kell, hogy Henry Markowitz (1952) forradalmasította a portfólióoptimalizálást azzal, hogy a várható hozamot és a szórást tekintette az eszköz hozamának és kockázatának számszerűsítéséhez szükséges kulcsszervezőknek. Ez az elmélet a portfólióépítést kvadratikus optimalizálási problémaként fogalmazza meg, ahol a cél a kockázathoz tartozó hozam maximalizálása, vagy ennek megfelelően a kockázat minimalizálása a várható hozam adott szintje mellett. Bár a Markowitz-modell elméletileg megalapozott, és nagy hatással van a portfóliókutatásra, a valós életben való alkalmazása kihívást jelent. A gyakorlatban a kutatónak vagy a portfóliókezelőnek meg kell becsülnie az értékpapírok hozamának ismeretlen várható értékét és varianciáját, hogy alkalmazni tudja azokat a modellben. A kockázatot és a hozamot általában pontatlanul számítják ki a *Historikus* minta felhasználásával, ami elfogadhatatlan, a mintán kívüli rosszabb teljesítményű megoldásokhoz vezet. Ezért a portfóliók kisszámú, szélsőséges súlyú részvényt tartanak, így ezek a portfóliók kevésbé diverzifikáltak. Továbbá a bemeneti adatokban bekövetkező kis változások gyakran jelentősen módosított súlyokat okoznak az optimális portfólióban. Az évek során a kutatók óriási erőfeszítéseket tettek a becslési hibák kezelésére, hogy javítsák a Markowitz-modell teljesítményét.

4.2 Legkisebb szórású modell

A legkisebb szórású portfólió (Clarke et al., 2011) a Markowitz modell egy sajátos változata, ahol a kockázatot minimalizáljuk és a hozamra semmilyen megkötésünk nincs. A befektető úgy kombinálja a részvényállományokat, hogy a teljes portfólió árfolyam-ingadozása csökkenjen.

A legkisebb szórású portfólió minimalizálja a befektetési kockázatokat. Ebben a módszerben a befektetők a portfólió volatilitásának csökkentése érdekében határozzák meg a varianciát. Ezért a minimális variancia biztosítása érdekében a befektetők diverzifikálják befektetéseiket.

A diverzifikált portfólió különböző ágazatokból származó részvényeket és különböző méretű vállalatokat tartalmaz. A részvényt piac kétféle kockázatot jelent - rendszerszintű és egyedi (nem rendszerszintű) kockázatokat. A rendszerszintű kockázatok minden ágazatot érintenek, ezért elkerülhetetlenek. A diverzifikációs stratégia azonban jelentősen csökkentheti a nem rendszeres kockázatokat.

Ha egy portfóliónak magas a szórása, az magasabb kockázatot jelent, ugyanakkor magasabb hozamot is. Ezért a varianciát a portfóliók megítélésére szolgáló mérőszámként használják. A kockázatkerülő hajlandóságtól függ, hogy ki mennyire kockázatos eszközbe fektet be. A legkisebb kockázatú portfóliót azért is szokták vizsgálni, mert csak a kockázat becslésével kell foglalkozni és az átlagos hozam számítását lehet mellőzni. A hozam becslése a kockázat becslésénél is kritikusabb probléma.

5. Legkisebb szórású modell vizsgálata különböző kovarianciamátrix becslő módszerek mellett

Dolgozatom ezen részében a legkisebb szórású modellt vizsgáltam, különböző, a dolgozat elején említett kovarianciamátrix becslő módszerek mellett. Ahogy a dolgozat korábbi részeiben tettem, most is előbb a módszertan ismertetésével kezdem, illetve azután fog következni az eredmények bemutatása.

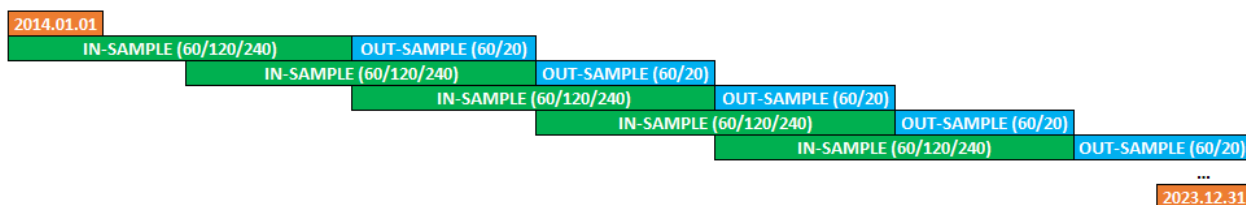
5.1 Módszertan

A vizsgálat során a figyelembe vett periódus az elmúlt 10 év volt, vagyis 2014 elejétől 2023 végéig rendelkezésre álló adatokat használtam fel. Adatok tekintetében továbbra is az S&P500 részvényeit használtam fel, annyi eltéréssel, hogy a korábbi random kiválasztástól eltérően most piaci kapitalizáció szerint rendeztem őket és mindig annyit választottam ki belőle amennyi az adott elemzéshez szükséges volt. A kísérletek során Walk-forward optimalizálási módszert használtam.

A Walk-forward optimalizálás (Kirkpatrick & Dahlquist, 2010) egy pénzügyi kereskedési koncepció, amely egy kereskedési stratégia tesztelési módszerét foglalja magában. Ez magában foglalja a historikus adatok számos periódusra való bontását, a stratégia paramétereinek optimalizálását minden egyes szegmensre, majd az optimalizált stratégia értékelését a következő szegmensre. Ebben a megközelítésben a kereskedőknek lehetőségük van arra, hogy megfigyeljék, hogyan reagál a technikájuk a különböző piaci fázisokra, és mennyire robusztus a változó piaci környezetekkel szemben. A Walk-forward optimalizációs módszer használatához először szükség van egy múltbeli adathalmazra és egy tesztelni kívánt kereskedési stratégia kiválasztására. Ezután az információhalmazt több kisebb szegmensre kell bontani, például évekre, hónapokra vagy negyedévekre. Minden szegmensnek két különböző részt kell tartalmaznia, amelyek egy mintavételen belüli (in-sample) és egy mintavételen kívüli (out-of-sample) részből állnak. Ezután a kereskedési stratégia paramétereit kell optimalizálni a mintán belüli (in-sample) részen. Az optimalizálási folyamatot minden egyes szegmensre vonatkozóan el kell végeznie. Majd használjuk az optimalizált stratégia paramétereit az out-of-sample részen, és dokumentáljuk a teljesítményt, majd folytassuk a folyamatot a következő szegmensre, és a korábbi out-of-sample részt használjuk új in-sample részként. Végül pedig értékeljük az eredményeket.

Ahogy említettem az általam vizsgált periódus tíz évet foglal magába, vizsgálataimat ismét különböző mennyiségű részvénytávval végeztem, illetve az in sample és out of sample értékeket is folyamatosan változtattam. Amit érdemes tudni, hogy az in-sample és out-of-sample értékek esetében napokat jelentenek, vagyis például egy 20 érték 20 kereskedelmi napot jelent, ami tulajdonképpen egy hónapra vonatkozó adat. A kísérleteim itt a következőképpen alakultak, először a piaci kapitalizáció szerinti legjobb első 50 részvényt vettem 60 in-sample értékkel és 20, valamint 60 out-of-sample értékekkel. Ezután növeltem a részvénytávot 100-ra az in-sample értéket 120-ra és az out-of-sample esetében 20 majd 60-ra növeltem. Végül az utolsó kísérletet pedig 200 részvényt 240 in-sample értékkel és ismételtén 20 majd 60 out-of-sample értékkel. Az

alábbi 5.1. ábrán ábrázolom a Walk-forward módszert, ábrázoltam az általam használt in-sample és out-of-sample értékeket a megadott 10 évnyi periódusban.



5.1. ábra. Walk-forward optimalizálási módszer

Ezt követően elemeztem a kapott portfólió hozamokat, a Sharpe ráta, CAGR (Compound annual growth rate), Éves kockázat és Max drawdown metrikákat figyelembe véve, illetve ábrázoltam a “growth of wealth”-et, vagyis portfólió növekedését az időszak során. Továbbiakban ezeknek a metrikáknak a bemutatása következik.

Sharpe-arány (vagy Sharpe-index vagy módosított Sharpe-arány) (Pav, 2021; Sharpe, 1966; Sharpe, 1994) nevét William Sharpe amerikai közgazdászról kapta. A Sharpe-arány a befektetés hozama (R_p) és a kockázatmentes hozam (R_f) közötti különbség osztva a befektetés szórásával (σ_p), ahogy az az alábbi (5.1) -es képletben is látható:

$$\text{Sharpe Arány} = \frac{R_p - R_f}{\sigma_p} \quad (5.1)$$

Egyszerűbben fogalmazva, a Sharpe-arány korrigálja a teljesítményt a befektető által vállalt többletkockázathoz képest. Arra használják, hogy egy befektetés teljesítményét a kockázat figyelembevételével mérték. Minél magasabb az arány, annál nagyobb a befektetés hozama a vállalt kockázathoz képest, és így annál jobb a befektetés. Az arány használható egyetlen részvény vagy befektetés, illetve egy teljes portfólió értékelésére.

A **CAGR** (Compound Annual Growth Rate, összetett éves növekedési ráta) (Chan, 2012) a befektetések átlagos éves növekedését méri egy adott időszak alatt. Megmutatja, hogy egy év alatt átlagosan mekkora volt a befektetések hozama. A CAGR nem egy valódi megtérülési ráta, hanem inkább egy reprezentatív számadat. Lényegében egy olyan szám, amely azt írja le, hogy egy befektetés milyen ütemben növekedett volna, ha minden évben ugyanolyan ütemben növekedett volna, és a nyereséget minden év végén újra befektetik. A valóságban ez a fajta teljesítmény

valószínűtlen. A CAGR azonban használható a hozamok simítására, így azok az alternatív módszerekhez képest könnyebben érthetők. Kiszámításához az időszak végén lévő befektetés értékét el kell osztani az időszak elején lévő értékével, az így kapott eredményt 1/évek száma hatványra kell emelni, ebből kivonni egyet és beszorozni százszal, hogy eredményként egy százalékos értéket kapjunk, ahogy ez az alábbi (5.2) képlet alapján is látható:

$$CAGR = \left(\left(\frac{EV}{BV} \right)^{\frac{1}{n}} - 1 \right) * 100 \quad (5.2)$$

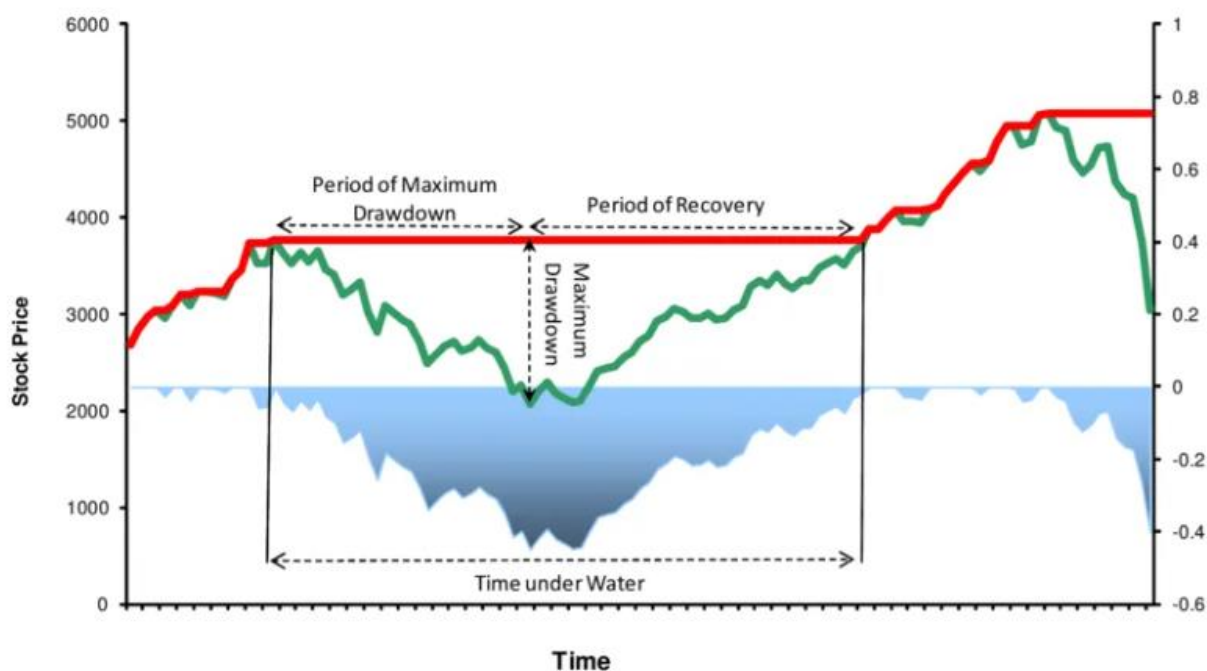
A képletben EV az időszak végén lévő befektetés értéke, BV az időszak elején lévő befektetés értéke, n pedig az évek számát jelenti.

Éves kockázat (Steiner, 2012) vagy annual volatility egy befektetés vagy portfólió kockázatának számszerűsítésére szolgál, azt jelezve, hogy a befektetés értéke egy adott időszakban valószínűleg mennyire ingadozik. Minél magasabb az évesített volatilitás, annál nagyobb a befektetés kockázata. Az évesített volatilitást jellemzően a múltbeli hozam adatok alapján számítják ki, és százalékban fejezik ki. Az éves volatilitási képlet (5.3) a következő:

$$\text{Éves kockázat} = \sigma * \sqrt{\text{időszakok száma egy évben}} \quad (5.3)$$

A képletben a szórás a pénzügyi eszköz hozamainak szórását jelenti egy adott időszakban. Az "időszakok száma egy évben" az évenkénti időszakok számát jelenti, ami 252. Azért használjuk a 252 négyzetgyökét, mert évente körülbelül 252 kereskedési nap van. A képletben azért szerepel az évenkénti időszakok számának négyzetgyöke, hogy a hozamok szórását évesített értéként fejezzük ki.

A maximum drawdown (MDD) (Steiner, 2010) a befektetés értékének maximális csökkenését méri, amelyet a legalacsonyabb mélypont és a mélypontot megelőző legmagasabb csúcs értéke közötti különbség ad. Az MDD-t hosszú időre számítják ki, amikor egy eszköz vagy befektetés értéke több fellendülési és visszaesési cikluson ment keresztül. Az alábbi 5.2. ábra jól szemlélteti, hogy ez az érték pontosan hogyan is határozható meg.



5.2. ábra. Maximum drawdown

(Forrás: corporatefinanceinstitute.com 2024)

Vagyis lényegében kiszámításához a portfólió csúcsértékéből kivonjuk a portfólió mélyponti értékét és az eredményt elosztjuk a portfólió csúcsértékével, ahogy az az alábbi (5.4) képlet is mutatja.

$$\text{Maximum drawdown} = \frac{\text{Portfólió csúcsértéke} - \text{Portfólió mélyponti értéke}}{\text{Portfólió csúcsértéke}} \quad (5.4)$$

A befektetők végső soron a legjobb eszközportfólió birtoklására törekednek, amely a legmagasabb várható hozamot és a legalacsonyabb kapcsolódó kockázatot kínálja. A maximum drawdown alkalmazása egy ilyen portfólió felépítéséhez hatékony módja annak, hogy biztosítsa a nyereséges befektetés lehetőségét.

Továbbá ahogy említettem “growt of wealth” (Choi et al., 2019) vagyis portfólió növekedési diagramokat készítettem. A portfólió növekedését úgy követjük nyomon, hogy feltételezzük, hogy pénzünket a felépített portfóliókba fektetjük. Ha egy egységnyi pénzt

befektetünk a j eszközbe ($j=1, 2, \dots, p$) a $(k-1)$ időszak elején, akkor az az időszak végére az alábbi (5.5) képlet alapján látható módon növekszik:

$$d_j^{(k-1)} := \prod_{t=t_{k-1}}^{t_{k-1}+L-1} (1 + r_{tj}) \quad (5.5)$$

Tehát amíg a CAGR egy évesített hozamot mutat százalékosan, addig a „growth of wealth” folyamatosan mutassa a teljes periódus alatt a portfólió növekedést egy kezdeti befektetett összegre. Pld. 1 befektetett dollár a periódus végén a több mint 3 dollárt ér a *Ledoit* esetén, ahogy majd azt a dolgozatom későbbi részében fogjuk látni.

Tehát összességében a vizsgálat során különböző esetekben a Walk-forward optimalizációt alkalmaztam, majd kiszámoltam a Sharpe arány, CAGR (Compound annual growth rate), éves kockázat és Max drawdown metrikákat, ábrázolom a “growth of wealth”-et vagyis portfólió növekedését és ezek alapján levontam a következtetéseket. Továbbá amit megemlítenék, hogy az említett metrikák közül a legfontosabb az éves volatilitás vagyis a kockázat ugyanis ezt próbáljuk minimalizálni.

5.2 Eredmények

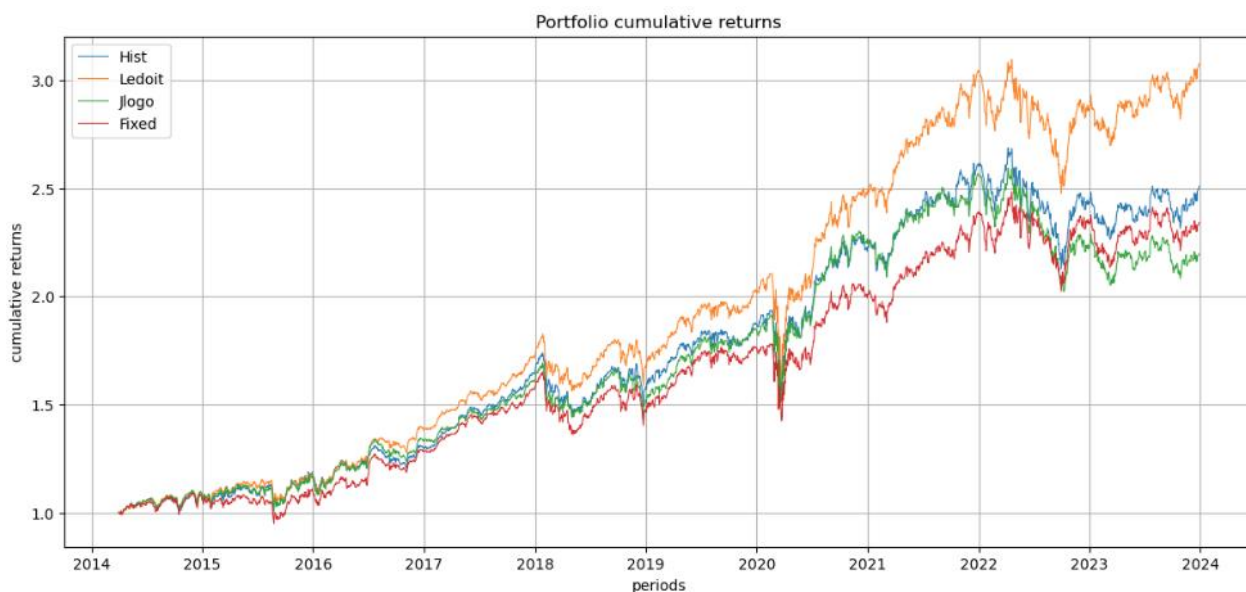
Most, hogy ismertettem a módszertant és érthető, hogy pontosan mit is vizsgáltam, rátérnék az eredmények bemutatására és értelmezésére. Sorrendben haladva először az 50 részvény, 60 in-sample, 20 out-of-sample vizsgálat eredményeivel kezdem.

Ahogy az az alábbi 5.1-es táblázaton is látható a *jLoGo* módszer rendelkezik a legalacsonyabb éves volatilitással, ami azt jelenti, hogy ez a módszer a legkevésbé ingadozó. A *Hist* módszer mutatja a legmagasabb volatilitást. A *Ledoit* módszer rendelkezik a legmagasabb Sharpe aránnyal, ami azt jelzi, hogy ez a módszer nyújtja a legjobb kockázat-korrigált hozamot a négy közül. A *jLoGo* módszer Sharpe aránya a legalacsonyabb. A *Fixed* módszer eredményezi a legkisebb maximális visszaesést, ami azt jelenti, hogy ez a portfólió a legkevésbé kockázatos a nagyobb veszteségek szempontjából. A *Ledoit* módszer viszont a legnagyobb visszaesést mutatja. CAGR-t tekintve a *Ledoit* módszernek a legmagasabb éves növekedési rátát, míg a *jLoGo* a legalacsonyabbat érte el.

50stock_60in_20out	Hist	Ledoit	jLoGo	Fixed
Sharpe ratio	0.728	0.879	0.645	0.685
Max. drawdown	-20.825	-23.327	-21.965	-19.949
CAGR	9.902	12.223	8.418	9.130
Annual volatility	14.385	14.276	14.072	14.232

5.1. táblázat 50 részvény, 60 in-sample, 20 out-of-sample

Az alábbi 5.3. ábrán a portfólió növekedését ábrázoltam, jól látható, hogy a legjobb eredményt a *Ledoit* módszerrel sikerült elérni, ami meglepő, hogy a másik két módszer a *Hist* módszerhez viszonyítva is gyengébben teljesít



5.3. ábra Portfólió növekedés - 50 részvény, 60 in-sample, 20 out-of-sample

Összesítve 50 részvény 60 in-sample, 20 out-of sample kísérlet során az éves volatilitást tekintve a *JLoGo* esetében a legalacsonyabb az érték. *Ledoit* esetében volt a legjobb a Sharpe arány, a drawdown a *Fixed* esetében volt a legalacsonyabb, CAGR-t tekintve pedig a *Ledoit* módszernek a legmagasabb éves növekedési rátája.

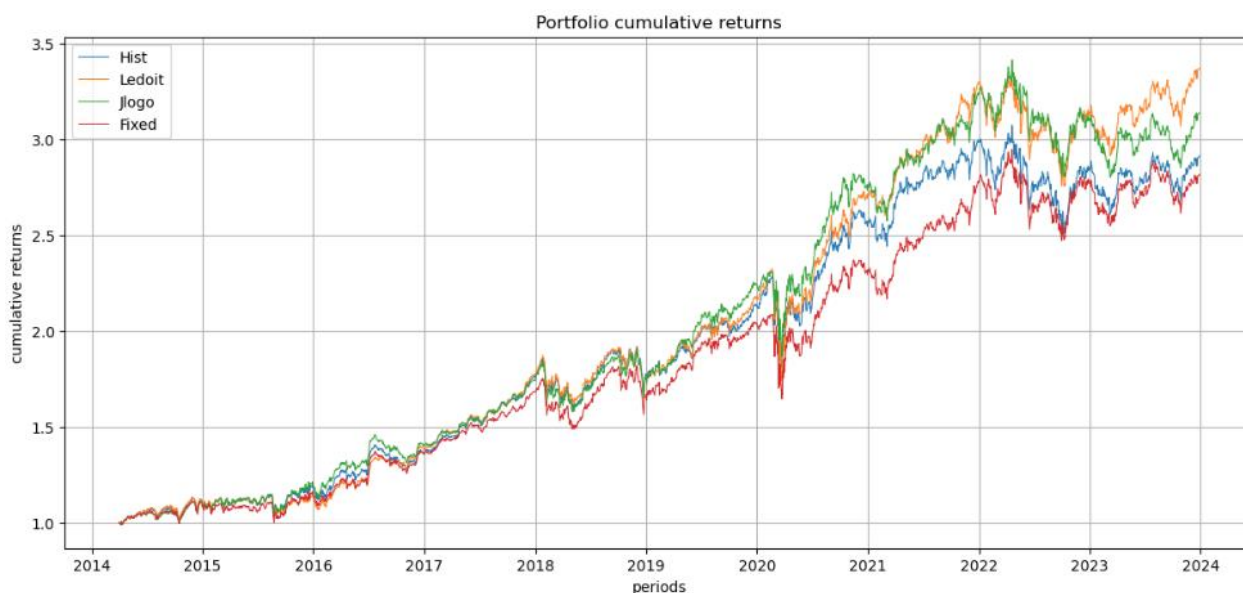
Ezzel tovább is lépnek a következő kísérletre (5.2. táblázat), ahol a korábbiakhoz képest annyi változott, hogy az out-of-sample értéket 60-ra növeltem. Ebben az esetben is a *JLoGo* módszer rendelkezik a legalacsonyabb éves volatilitással. Továbbá ebben az esetben is a *Ledoit* módszer éri el a legmagasabb Sharpe arányt, ami azt jelenti, hogy ez a módszer nyújtja a legjobb kockázattal korrigált hozamot. A *JLoGo* módszer eredményezi a legkisebb maximális visszaesést,

ami azt jelenti, hogy ez a portfólió a legkevésbé kockázatos a nagyobb veszteségek szempontjából. CAGR-t tekintve továbbra is a *Ledoit* esetében a legmagasabb.

50stock_60in_60out	Hist	Ledoit	jLoGo	Fixed
Sharpe ratio	0.824	0.933	0.892	0.806
Max. drawdown	-22.055	-25.061	-19.113	-21.082
CAGR	11.585	13.271	12.443	11.202
Annual volatility	14.590	14.490	14.294	14.473

5.2. táblázat. 50 részvény, 60 in-sample, 60 out-of-sample

Ebben az esetben is ábrázoltam (5.4. ábra) a portfólió növekedését, ahogy az az alábbi ábrán látható. Az out-of sample érték növelésével a *JLoGo* módszer megközelítette, illetve van hol jobban teljesít a *Ledoit*nál, a *Fixed* viszont továbbra is gyengébb a *Hist* módszernél.



5.4. ábra. Portfólió növekedés - 50 részvény, 60 in-sample, 60 out-of-sample

Összesítve 50 részvény 60 in-sample, 60 out-of sample kísérlet során azzal, hogy növeltük az out-of-sample értéket Sharpe és CAGR tekintetében továbbra is a *Ledoit* a legjobb, éves kockázat, illetve maximum drawdown szempontjából pedig a *JLoGo*.

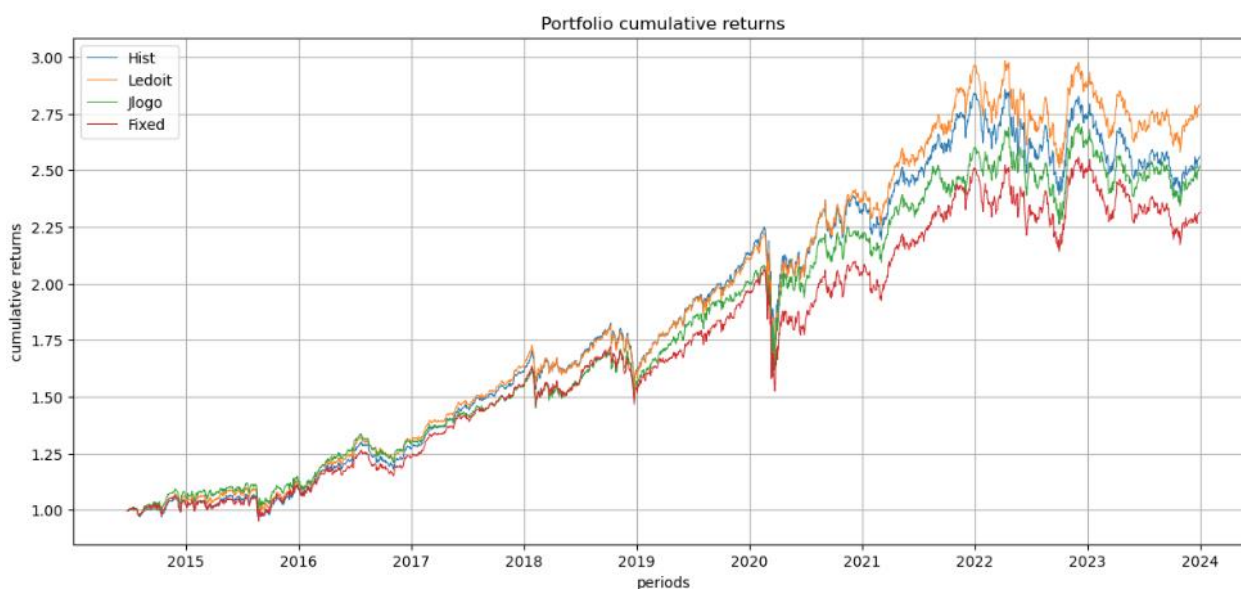
A következőkben a részvények számát 100-ra növeltem, az in-sample értékét 120-ra és az out-of-sample értékét 20-ra csökkentettem (5.3. táblázat). Éves kockázat tekintetében továbbra is a *JLoGo* módszer éri el a legjobb eredményeket. Sharpe arányt tekintve a *Ledoit* módszer a legjobb,

ahogy CAGR tekintetében is. Maximum drawdown esetében pedig a *jLoGo* rendelkezik a legjobb értékkel.

100stock_120in_20out	Hist	Ledoit	jLoGo	Fixed
Sharpe ratio	0.784	0.850	0.779	0.706
Max. drawdown	-24.447	-25.755	-22.206	-26.200
CAGR	10.395	11.403	10.213	9.233
Annual volatility	13.836	13.828	13.679	13.882

5.3. táblázat. 100 részvény, 120 in-sample, 20 out-of-sample

A portfólió növekedését tekintve (5.5. ábra) is jól látható, hogy a *Ledoit* teljesít a legjobban, ami meglepő, hogy ilyen tekintetben a másik két módszer mennyire elmarad a *Hist* módszertől, holott az előző kísérlet során a *JLoGo* majdnem azonosan teljesített, mint a *Ledoit*.



5.5. ábra. Portfólió növekedés - 100 részvény, 120 in-sample, 20 out-of-sample

Végül arra a megállapításra juthatunk, hogy a kísérlet során a *Ledoit* módszer érte el a legjobb Sharpe és CAGR értéket, a max drawdown és az éves kockázat tekintetében pedig a *jLoGo* a legjobb.

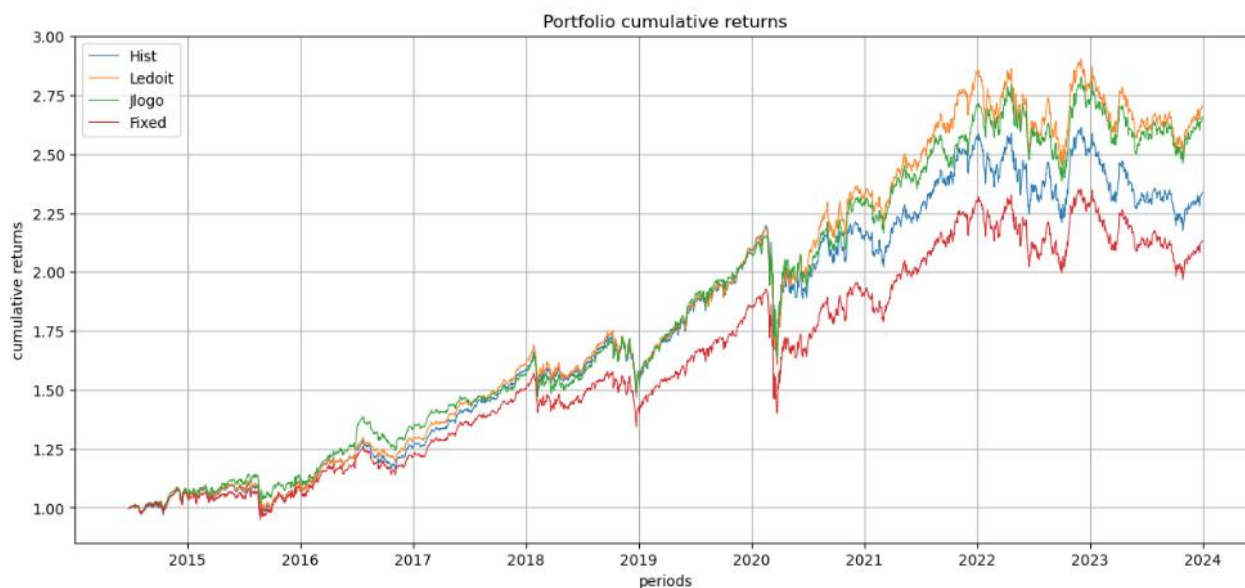
Ezzel tovább is lépnek a következő kísérletre, ahol annyi változott, hogy az out-of-sample értékét 60-ra növeltem. Az alábbi 5.4. táblázat alapján látható, hogy éves kockázat tekintetében továbbra is a *JLoGo* teljesít a legjobban. A *Ledoit* a Sharpe arány és CAGR tekintetében éri el a

legjobb eredményeket. Maximum drawdown esetében pedig a *JLoGo* eredménye számít a legjobbnak.

100stock_120in_60out	Hist	Ledoit	jLoGo	Fixed
Sharpe ratio	0.706	0.816	0.812	0.637
Max. drawdown	-26.716	-26.583	-23.758	-27.299
CAGR	9.344	11.039	10.830	8.292
Annual volatility	14.054	14.053	13.842	14.067

5.4. táblázat. 100 részvény, 120 in-sample, 60 out-of-sample

Az alábbi 5.6. ábrán ismét a portfólió növekedés látható, érdemes megfigyelni, hogy azzal, hogy növeltük a részvények számát és az in sample értéket a *Ledoit* és a *JLoGo* jobban elkülönült és a *Hist* és a *Fixed* lemaradása meg inkább egyértelműbb. Továbbá, hogy az out-of-sample értékének növelésével a *JLoGo* módszer javulást tudott elérni.



5.6. ábra. Portfólió növekedés - 100 részvény, 120 in-sample, 60 out-of-sample

Tehát összegezve, a kísérlet alapján megállapítható, hogy a Sharpe arány és a CAGR-t tekintve a *Ledoit* módszer segítségével kaptuk a legjobb eredményt, maximum drawdown és éves kockázat tekintetében pedig továbbra is a *JLoGo* módszer teljesít a legjobban.

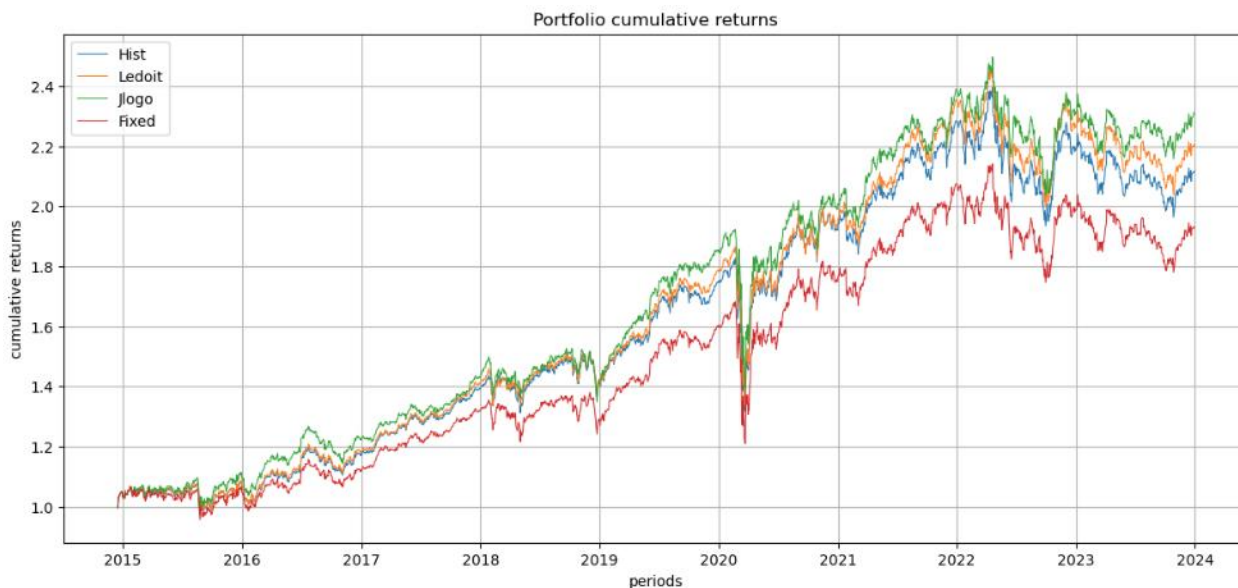
A következő kísérletben megemeltém a részvények számát 200-ra, az in-sample értékét 240-re és az out-of-sample értékének pedig 20-at állítottam be. Megnézve az eredményeket (5.5.

táblázat) látható, hogy az eddigiektől eltérően ebben az esetben nem a *JLoGo* hanem a *Ledoit* érte el a legjobb eredményt. Sharpe arány tekintetében most a *JLoGo* a jobb, ahogy a CAGR és Max drawdown tekintetében is.

200stock_240in_20out	Hist	Ledoit	jLoGo	Fixed
Sharpe ratio	0.665	0.700	0.737	0.587
Max. drawdown	-27.817	-28.246	-27.473	-28.057
CAGR	8.659	9.156	9.730	7.566
Annual volatility	13.972	13.909	13.910	14.146

5.5. táblázat. 200 részvény, 240 in-sample, 20 out-of-sample

A portfólió növekedését megtekintve (5.7. ábra) továbbra is az látható, hogy a *Ledoit* és a *JLoGo* módszer közötti különbség minimális a *Fixed* pedig még a *Hist* módszertől is el van maradva.



5.7. ábra. Portfólió növekedés - 200 részvény, 240 in-sample, 20 out-of-sample

Összefoglalva ebben a kísérletben, a *Ledoit* és a *JLoGo* módszer Sharpe arány és CAGR tekintetében nagyon hasonló eredményeket ér el, a *JLoGo* viszont minimálisan jobb. Max drawdown esetében a *JLoGo* és éves kockázat tekintetében a *Ledoit* ért el jobb eredményt.

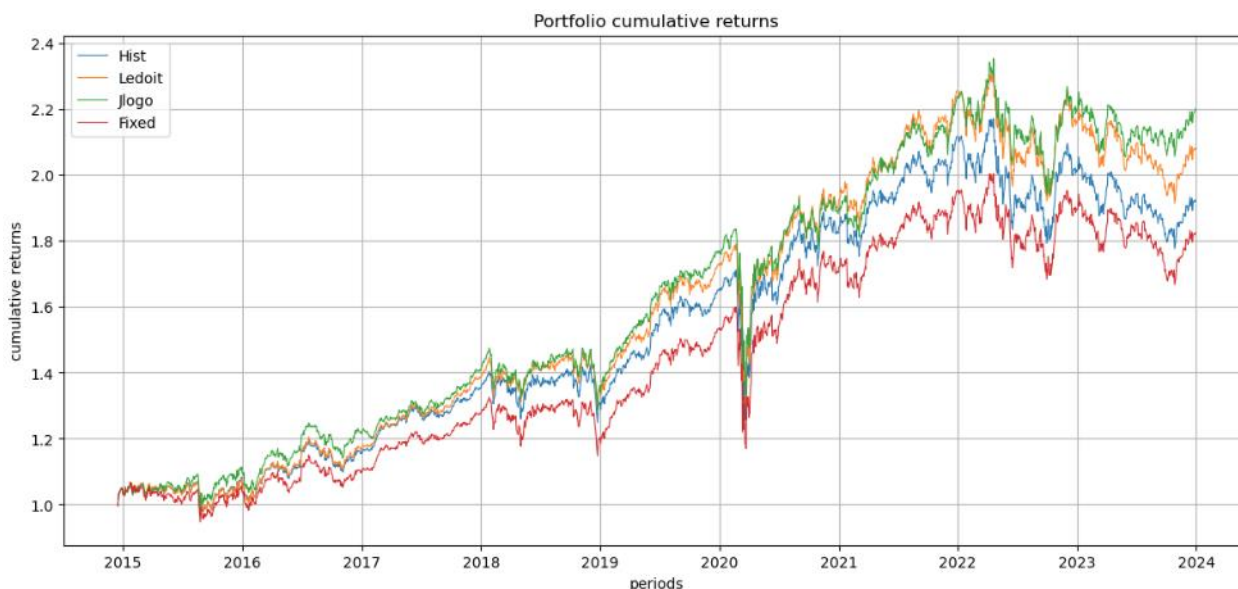
Végül pedig az utolsó kísérletben az out-of-sample értékének 60-at állítottam be, ebben az esetben (5.6. táblázat) éves kockázatot tekintve a *JLoGo* érte el a legjobb eredményt. Sharpe arányt tekintve is a *JLoGo* a legjobb, ami azt jelenti, hogy ez a módszer nyújtja a legjobb kockázattal korrigált hozamot, tehát a *JLoGo* ebben az esetben minimálisan jobb értéket ért el, mint a *Ledoit*.

A CAGR esetében is ugyanez a figyelhető meg, vagyis a *JLoGo* eredménye jobb a *Ledoit* eredményénél. Max drawdown tekintetében meglepetésre a *Hist* módszer lett a legjobb.

200stock_240in_60out	Hist	Ledoit	jLoGo	Fixed
Sharpe ratio	0.584	0.647	0.693	0.536
Max. drawdown	-26.569	-27.018	-27.408	-26.801
CAGR	7.521	8.469	9.133	6.894
Annual volatility	14.130	14.100	14.053	14.364

5.6. táblázat. 200 részvény, 240 in-sample, 60 out-of-sample

A portfólió növekedés diagrammot megtekintve (5.8. ábra) látható, hogy A *Ledoit* és *JLoGo* portfóliók nyújtották a legjobb eredményt, megerősítve a korábbi elemzések eredményeit, amelyek magas Sharpe arányt és CAGR-t mutattak.



5.8. ábra. Portfólió növekedés - 200 részvény, 240 in-sample, 60 out-of-sample

Összefoglalva látható, hogy ebben a kísérletben a *JLoGo* jobb eredményeket tudott nyújtani, mint a *Ledoit*, viszont a különbségek nem annyira jelentősök.

Az alábbi 5.7. táblázatban összesítettem a kísérleteket, minden kísérletnél, minden metrikát vizsgálva az első és a második legjobb módszer tartalmazza a táblázat, ez alapján megállapítható, hogy éves kockázat szempontjából egyértelműen a *JLoGo* módszer éri el a legjobb eredményt, ha

a Sharpe arányt nézzük akkor a *Ledoit* módszer a legjobb, max drawdown szempontból a *JLoGo* módszer tekinthető a legjobbnak, CAGR tekintetében a *Ledoit* módszer teljesít a legjobban.

Stratégia	Sharpe Ratio		Max. Drawdown		CAGR		Annual Volatility	
	1	2	1	2	1	2	1	2
50stock_60in_20out	Ledoit (0.879)	Hist (0.729)	Fixed (-19.949)	Hist (-20.825)	Ledoit (12.223)	Hist (9.902)	JLogo (14.072)	Fixed (14.232)
50stock_60in_60out	Ledoit (0.932)	JLogo (0.892)	JLogo (-19.113)	Fixed (-21.082)	Ledoit (13.271)	JLogo (12.443)	JLogo (14.294)	Fixed (14.473)
100stock_120in_20out	Ledoit (0.850)	Hist (0.784)	JLogo (-22.206)	Hist (-24.447)	Ledoit (11.403)	Hist (10.395)	JLogo (13.629)	Ledoit (13.828)
100stock_120in_60out	Ledoit (0.816)	JLogo (0.812)	JLogo (-23.758)	Ledoit (-26.583)	Ledoit (11.039)	JLogo (10.830)	JLogo (13.842)	Ledoit (14.053)
200stock_240in_20out	JLogo (0.737)	Ledoit (0.700)	JLogo (-27.473)	Hist (-27.817)	JLogo (9.730)	Ledoit (9.156)	Ledoit (13.909)	JLogo (13.910)
200stock_240in_60out	JLogo (0.693)	Ledoit (0.647)	Hist (-26.569)	Fixed (-26.801)	JLogo (9.133)	Ledoit (8.469)	JLogo (14.053)	Ledoit (14.100)

Ledoit
 JLogo
 Hist
 Fixed

5.7.táblázat. Összesített eredmények

6. Következtetések

A dolgozatom elkészítése során látva a *Hist* módszer teljesítményét meg inkább egyértelművé vált számomra, hogy mennyire fontos a megfelelő kovarianciamátrix becselő módszer alkalmazása. Az általam vizsgált módszerek ugye a *Hist*, *Ledoit*, *JLoGo* és a *Fixed* voltak.

Random generált adatokon alkalmazva, ahol ismert volt az eredeti mátrix a *Ledoit* módszer minden tekintetben a legjobb volt. A kísérlet során, ahol az S&P 500 részvényei közül random választottam különböző mennyiségű részvényt és különböző időszakokon becsültem a kovarianciamátrixot a min_var_loss tekintetében a *Fixed* érte el a legjobb eredményt, a MAE metrikát vizsgálva pedig a *Ledoit*, ennél viszont fontos megemlíteni, hogy az eredmények lehet nem a legpontosabbak, mivel ebben az esetben az eredeti mátrix csak feltételezett.

A módszereket a legkisebb szórású modell portfólió optimalizálás során alkalmazva arra jutottam, hogy min_var_loss tekintetében az első kísérletem nyújtott valós eredmény ugyanis a *Ledoit* módszer jónak bizonyult és a *Fixed* sokszor még a *Hist* módszernél is gyengébb eredményeket produkált. Érdeemes megemlíteni a *JLoGo* módszert is ugyanis éves kockázat és max drawdown tekintetében a legjobb és ahogy a kísérlet paraméterei növekedtek a teljesítménye is javult. Sharpe arány és CAGR tekintetében egyértelműen a *Ledoit* a jobb.

Mindezek tudatában az általam végzet kutatás során látható a *Hist* módszer egyértelmű gyengesége és a kovarianciamátrix becselő módszerek szükségessége. Az általam vizsgált módszerek közül az eredmények alapján számomra a *Ledoit* módszer tűnik a legjobbnak, a *JLoGo* is jó eredményeket ért el, viszont a *Ledoit* módszer, ami bármilyen paramétértől függetlenül mindig képes a legjobb eredményt nyújtani, vagy ha nem is a legjobbat, de olyan eredményt, ami teljesen ideális, vagyis jobban illeszkedik a körülményekhez és stabil eredményt nyújt.

Irodalomjegyzék

1. Barfuss, W., Massara, G. P., Di Matteo, T., & Aste, T. (2016). Parsimonious modeling with information filtering networks. *Physical Review E*, 94, 062306.
2. Bun, J., Bouchaud, J.-P., & Potters, M. (2017). Cleaning large correlation matrices: Tools from random matrix theory. *Physics Reports*, 666, 1–109.
3. Chan, E. (2012). *Harvard Business School Confidential: Secrets of Success*. John Wiley & Sons, 185
4. Chen, W. (1997). Linear algebra: Chapter 10, orthogonal matrices. Letöltés dátuma: 2024.05.15, forrás: <https://www.williamchen-mathematics.info/lnlafolder/la10.pdf>.
5. Choi, Y. G., Lim, J., & Choi, S. (2019). High-dimensional Markowitz portfolio optimization problem: Empirical comparison of covariance matrix estimators. *Journal of Statistical Computation and Simulation*, 89(7), 1278–1300.
6. Clarke, R., de Silva, H., & Thorley, S. (2011). Minimum-variance portfolio composition. *The Journal of Portfolio Management*, 37(2), 31-45.
7. corporatefinanceinstitute.com (2024): CFI Team; Maximum Drawdown; <https://corporatefinanceinstitute.com/resources/career-map/sell-side/capital-markets/maximum-drawdown>; Letöltve: 2024.06.19
8. Engle, R. F., Ledoit, O., & Wolf, M. (2019). Large dynamic covariance matrices. *Journal of Business & Economic Statistics*, 37(2), 363–375.
9. Horn, R. A., & Johnson, C. R. (2013). *Matrix Analysis*. Cambridge University Press.
10. Kirkpatrick, C. D., & Dahlquist, J. R. (2010). *Technical Analysis: The Complete Resource for Financial Market Technicians*. FT Press, 548
11. Ledoit, O., & Wolf, M. (2003). Honey, I shrunk the sample covariance matrix. Letöltés dátuma: 2024.05.15, forrás: <http://ledoit.net/honey.pdf>.
12. Ledoit, O., & Wolf, M. (2004). A well-conditioned estimator for large-dimensional covariance matrices. *Journal of Multivariate Analysis*, 88(2), 365-411.
13. Ledoit, O., & Wolf, M. (2017). Optimal estimation of a large-dimensional covariance matrix under Stein's loss. *Journal of Multivariate Analysis*, 1-64.
14. Ledoit, O., & Wolf, M. (2020). Analytical nonlinear shrinkage of large-dimensional covariance matrices. *Annals of Statistics*, 48(5), 3043-3065.

15. López de Prado, M. M. (2020). *Machine Learning for Asset Managers*. Cambridge University Press.
16. López de Prado, M. (2016). Building diversified portfolios that outperform out-of-sample. *Journal of Portfolio Management*, 42, 59–69.
17. Markowitz, H. (1952). Portfolio selection. *Journal of Finance*, 7, 77–91.
18. Pál, L. (2022). Asset Allocation Strategies Using Covariance Matrix Estimators, *Acta Universitatis Sapientiae, Economics and Business*, 10(1), 133-144.
19. Pav, S. E. (2021). *The Sharpe Ratio: Statistics and Applications*. CRC Press.
20. Portfoliooptimizer.io. (2022). Correlation matrices denoising: Results from random matrix theory. Letöltés dátuma: 2024.05.15, forrás: <https://portfoliooptimizer.io/blog/correlation-matrices-denoising-results-from-random-matrix-theory/>.
21. Sharpe, W. F. (1966). Mutual fund performance. *Journal of Business*, 39(S1), 119–138.
22. Sharpe, W. F. (1994). The Sharpe ratio. *The Journal of Portfolio Management*, 21(1), 49–58.
23. Simon, T. L., & Pál, L. (2019). Részvényportfóliók optimalizálása webszolgáltatások segítségével. In L. J. Tóczos & Z. Makó (Eds.), *Stadium sorozat* (pp. 1-25). Sepsiszentgyörgy: T3 Kiadó.
24. Steiner, A. (2010). Ambiguity in calculating and interpreting maximum drawdown. *Electronic Journal*, 1-2.
25. Steiner, A. (2012). Annualized volatility. *Electronic Journal*, 1-8.
26. Willmott, C. J., & Matsuura, K. (2005). Advantages of the mean absolute error (MAE) over the root mean square error (RMSE) in assessing average model performance. *Climate Research*, 30, 79–82.

MATE Szervezeti és Működési Szabályzat

III. Hallgatói Követelményrendszer

III.1. Tanulmányi és Vizsgaszabályzat

6.13. sz. függelék: A MATE egységes szakdolgozat / diplomadolgozat / záródolgozat / portfólió készítési útmutatója

4.2. sz. melléklete: Nyilatkozat a záródolgozat/szakdolgozat/diplomadolgozat/portfólió nyilvános hozzáféréséről és eredetiségéről

NYILATKOZAT

a záródolgozat/szakdolgozat/diplomadolgozat/portfólió¹ nyilvános hozzáféréséről és eredetiségéről

A hallgató neve:

TULIT IMRE

A Hallgató Neptun kódja:

IEWGWS

A dolgozat címe:

KOVARIANCIA-MÁTRIX BECSLÉSI MÓDSZEREK
VIZSGÁLATA BEFEKTETÉSI STRATÉGIÁKBAN
2024

A megjelenés éve:

A konzulens intézetének neve:

VIDEKFEJLESZTÉS ÉS FENNTARTHATÓ
GAZDASÁG INTÉZET

A konzulens tanszékének a neve:

BEFEKTETÉSI, PÉNZÜGYI ÉS SZÁMVITELI
TANSZÉK

Kijelentem, hogy az általam benyújtott záródolgozat/szakdolgozat/diplomadolgozat/portfólió² egyéni, eredeti jellegű, saját szellemi alkotásom. Azon részeket, melyeket más szerzők munkájából vettem át, egyértelműen megjelöltem, és az irodalomjegyzékben szerepeltettem.

Ha a fenti nyilatkozattal valótlan állítottnak tudomásul veszem, hogy a záróvizsga-bizottság a záróvizsgából kizár és a záróvizsgát csak új dolgozat készítése után tehetek.

A leadott dolgozat, mely PDF dokumentum, szerkesztését nem, megtekintését és nyomtatását engedélyezem.

Tudomásul veszem, hogy az általam készített dolgozatra, mint szellemi alkotás felhasználására, hasznosítására a Magyar Agrár- és Élettudományi Egyetem mindenkori szellemi tulajdon-kezelési szabályzatában megfogalmazottak érvényesek.

Tudomásul veszem, hogy dolgozatom elektronikus változata feltöltésre kerül a Magyar Agrár- és Élettudományi Egyetem könyvtári repozitori rendszerébe. Tudomásul veszem, hogy a megvédett és

- nem titkosított dolgozat a védést követően
- titkosításra engedélyezett dolgozat a benyújtásától számított 5 év eltelte után nyilvánosan elérhető és kereshető lesz az Egyetem könyvtári repozitori rendszerében.

Kelt: CSIKSZENTSIMON év 2024 hó október nap 24

Hallgató aláírása

¹ A megfelelő dolgozattípus meghagyása mellett a többi típus törölendő.

² A megfelelő dolgozattípus meghagyása mellett a többi típus törölendő.

MATE Szervezeti és Működési Szabályzat

III. Hallgatói Követelményrendszer

III.1. Tanulmányi és Vizsgaszabályzat

**6.13. sz. függelék: A MATE egységes szakdolgozat /
diplomadolgozat / záródolgozat / portfólió készítési útmutatója**

4.1. sz. melléklete: Konzulensi nyilatkozat

NYILATKOZAT

A hallgató konzulenseként nyilatkozom arról, hogy a záródolgozatot/szakdolgozatot/diplomadolgozatot/portfóliót áttekintettem, a hallgatót az irodalmi források korrekt kezelésének követelményeiről, jogi és etikai szabályairól tájékoztattam.

A záródolgozatot/szakdolgozatot/diplomadolgozatot/portfóliót a záróvizsgán történő védésre javaslom / nem javaslom².

A dolgozat állam- vagy szolgálati titkot tartalmaz: igen nem³

Kelt: 2024.09.14.


belső konzulens

¹ A megfelelő aláhúzendó.

² A megfelelő aláhúzendó.

³ A megfelelő aláhúzendó.