

Vysoká škola ekonomická v Praze

Diplomová práce

2017

David Scholle

Vysoká škola ekonomická v Praze

Fakulta financí a účetnictví

Studijní obor: Bankovníctví a pojišťovnictví



Název diplomové práce:

**Modelování kovariancí v časových řadách metodou DCC GARCH
a její aplikace na analytický výpočet Value at Risk**

Autor diplomové práce: David Scholle

Vedoucí diplomové práce: prof. RNDr. Jiří Witzany, Ph.D.

P r o h l á š e n í

Prohlašuji, že jsem tuto diplomovou práci
vypracoval samostatně s využitím literatury a informací,
na něž odkazuji.

V Praze dne 30.10.2017

Podpis

Poděkování

Rád bych poděkoval vedoucímu práce, jímž byl pan prof. RNDr. Jiří Witzany, Ph.D., za důležité rady a směřování v ohledu práce, za pomoc s tématem a mnohá doporučení, a stejně tak za jeho trpělivost.

Název diplomové práce:

Modelování kovariancí v časových řadách metodou DCC GARCH a její aplikace na analytický výpočet Value at Risk

Abstrakt:

Cílem této práce je otestovat možnost zapojení modelu s nekonstantním rozptylem při modelování volatility tržních indexů. Využito navíc bude složitějších metod DCC GARCH, které navíc zohledňují proměnlivost korelace mezi světovými trhy v průběhu času. Protože tyto korelace se ukazují jako nekonstantní, v zájmu splnění předpokladů je potřeba zapojení těchto složitých modelů. Jejich užitečnost bude potom ozkoušena porovnáním s benchmarkovým modelem, jímž bude EWMA. Odhadnutá hodnota volatility potom bude využita k výpočtu Value at Risk, kde by měla nalézt uplatnění díky zpřesnění odhadu podstoupeného rizika pro investory.

Klíčová slova:

DCC GARCH, EWMA, Value at Risk, heteroskedastické modely

Title of the Master's Thesis:

Modelling of covariance in time series using DCC GARCH and its application in analytical computation of Value at Risk

Abstract:

The target of this thesis is to test the option of utilizing models without constant variance when modelling the volatility of market indices. We will use DCC GARCH, an even more robust method enabling us to adapt to time variability of correlation between world markets. As these correlations were shown to be non-constant, we should use these complex models. Their usefulness will be examined by the comparison to EWMA, a benchmark method. The estimated volatility will then be used to calculate Value at Risk. That should lead to its utilization by creating better estimates of investors' risk exposure.

Key words:

DCC GARCH, EWMA, Value at Risk, heteroscedastic models

Obsah

Úvod.....	8
Value at Risk.....	10
Úvod.....	10
Definice, Výhody, nevýhody	10
Formální definice	11
Využití.....	11
Postupy výpočtu.....	13
Alternativy VaR.....	16
Modely s podmíněnou heteroskedasticitou.....	18
Definice modelu EWMA	21
Definice modelu GARCH(p,q)	24
Odhad pomocí maximální věrohodnosti	28
Bayesovský přístup	31
Bayesovský odhad GARCH(1,1)	33
Empirická studie	34
Data	36
Odhad modelu pomocí EWMA	43
Odhad GARCH(1,1) pomocí ML	51
Odhad Value at Risk pomocí ML	55
Odhad pomocí bayesovského přístupu.....	60
Portfoliový přístup	65
Závěr	69

1. Úvod

Při modelování hodnoty VaR se nabízí využívat nějakou z klasických modelů časových řad. Mohli bychom například pro střední hodnotu využít model AR a potom za využití nějakého předpokladu o rozdělení výnosů odhadnout hodnotu VaR. K tomu je potřeba odhad volatility, kterým by v takovém případě mohla být například výběrová směrodatná odchylka. Ve finančním světě ovšem při tomto řešení podobných problémů vždy narážíme na stejný problém. Pohyby na trzích nevykazují homoskedastické chování (Martin & Klemkovsky, 1975). Volatilita není konstantní a tím pádem není možno ji odhadovat pomocí konstanty.

Při pozorování dat se dá ovšem zjistit, že volatilita přece jen vykazuje nějaké predikovatelné chování. Tím je především to, že je vyšší v dobách krize a tato změna nějakou dobu přetrvává a postupem času se po šoku snižuje (Sandoval, 2012). Pokud totiž například na trh s akciemi přijde nová informace, nebo jiný podnět, mnoho investorů tento signál využije jako důvod k investicím, což začne pohybovat s cenami. Tento pohyb cen má potom za následek další periodu zvýšeného objemu obchodování, který vede k dalšímu zvýšení volatility. Postupem času se však situace na trhu uklidňuje a pokud se neobjeví další informace, volatilita klesne na svůj dlouhodobý průměr. Popsané chování zřejmě vykazuje prvky, které by vedly k AR modelu odhadování časové řady volatility. Jev, při kterém se objevují periody se zvýšenou či sníženou volatilitou, se v angličtině nazývá „volatility clustering“.

Odtud vznikl nápad pro vytvoření tzv. ARCH (autoregressive conditional heteroskedasticity) modelu. Ten využívá toho, že ačkoliv nepodmíněná volatilita (tedy volatilita bez jakýchkoliv známých informací o jejích předchozích hodnotách) je v předpokladu konstantní, podmíněná volatilita je v některých obdobích vyšší. V modelu ARCH je toto zavedeno pomocí historických chyb predikce (takzvané inovace). Pokud v posledním období došlo při odhadu k chybě, kdy skutečná volatilita byla vyšší, než jsme předpokládali, potom pro další období predikci navýšíme o nějakou část této chyby. Tento model, ačkoliv je užitečný a rychle se na trzích rozšířil, měl stále nedostatky. Jedním z nich bylo, že v modelu jako významné vycházely i chyby ze vzdálenější minulosti. Délka zpoždění tedy při odhadu byla příliš vysoká na to, aby byly parametry modelu jednoduše spočteny.

Proto vznikl návrh nahradit model odhadu volatility ARMA modelem. Kromě chyby odhadu se v modelu vyskytuje také předchozí hodnota volatility. Tím vznikl takzvaný GARCH model. Jeho nejjednodušší verze GARCH(1,1) v sobě, díky vlastnostem ARMA modelu, zahrnuje celou minulost chyb predikce. Tyto mají ale klesající efekt, nejvyšší vliv tedy mají nedávné hodnoty. Započtením celé historie tak zanikl problém, modelu ARCH, kdy bylo nutné odhadovat příliš velké množství parametrů. Ty jsou nyní zahrnuty v jediném parametru a u starších hodnot se vyskytuje jeho mocnina. Alternativou je využití modelu EWMA (Exponentially weighted moving average), který podobně jako GARCH přiřazuje starším hodnotám nižší váhy, ty se snižují exponenciálně. Ten je využíván v rámci velmi rozšířeného modelu RiskMetrics (RiskMetrics Technical Document, 1996). Tento model využijeme v práci jako benchmarkový. Výsledky metody GARCH tedy budeme porovnávat s modelem EWMA, abychom určili užitečnost. EWMA je přitom specifický typ modelu GARCH, takže v podstatě zjišťujeme, zda se vyplatí zvyšovat složitost v zájmu zpřesnění modelu.

Při odhadu Value at Risk se tedy nejdříve musí vygenerovat střední hodnota budoucího výnosu (ta může být generována jako samostatná časová řada, použitelná je například regrese, případně může být také použita ARMA), potom se za pomoci modelu GARCH simuluje budoucí volatilita. Ta se potom spolu s předpoklady o rozdělení výnosů použije pro nalezení 95-kvantilu, který odpovídá Value at Risk. Při generování můžeme rovnou vytvořit 95% interval spolehlivosti, jehož dolní hranice bude právě odhadovaná VaR.

Protože při výpočtu VaR po diverzifikaci portfolia potřebujeme znát především korelace mezi výnosy použitých aktiv, budeme v této práci využívat alternativu metody, zvanou DCC (dynamic conditional correlation) GARCH. Ta předpokládá především nekonstantní korelace mezi jednotlivými aktivy. Jde o zobecnění metody CCC (constant conditional correlation), která sice sdílí předpoklad heteroskedasticity, ale korelační matice zůstává konstantní. Tím je potřeba odhadnout nižší množství parametrů, ale předpoklad konstantní korelace mezi aktivy může být velmi silný a je potřeba ověřit, zda tato podmínka skutečně platí.

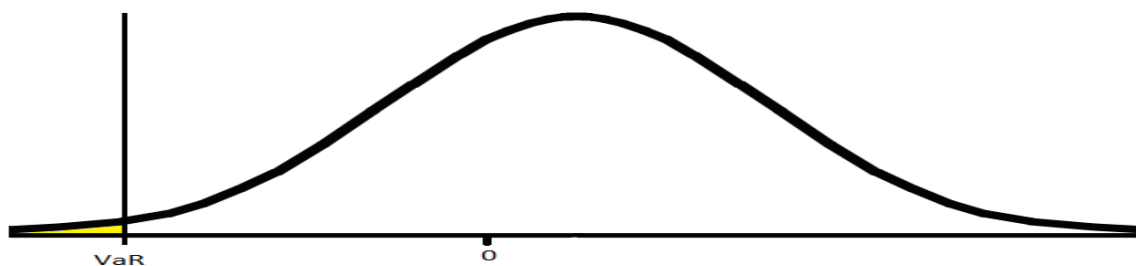
2. Value at Risk

2.1. Úvod

Problematika odhadu podstoupeného rizika je jedním z nejdůležitějších problémů v oblasti řízení rizik. Mezi nejpoužívanější míry rizika patří Value at Risk (název se do češtiny většinou nepřekládá, v literatuře se občas dá nalézt jako „hodnota v riziku“). Díky relativní jednoduchosti výpočtu a snadné interpretaci je VaR velmi rozšířený a často používaný jako jedna z hlavních metod řízení rizika. Podle Basel II je aktuálně nejdůležitější způsob měření tržního rizika. Ten vyžaduje od finančních institucí řízení tržního rizika na základě 10denního VaR s hladinou významnosti 1 %. A to i přesto, že se v poslední době objevuje čím dál více metod pro měření rizika, které mají napravit její nedostatky.

2.2. Definice, Výhody, nevýhody

Hodnota Value at Risk na hladině významnosti α je matematicky definována jako hodnota kvantilové funkce rozdělení výnosů v α . To zjednodušeně znamená, že s pravděpodobností $1-\alpha$ nedosáhne portfolio ztráty vyšší než $\text{VaR}(\alpha)$. Někdy se používá zápis $\text{VaR}(1-\alpha)$ – to nevede k žádným problémům, neboť VaR se používá pouze pro nadměrné ztráty, na druhé straně by tato hodnota postrádala smysl.



Obrázek 2-1: Ilustrace Value at Risk

Na obrázku 2-1 máme normální rozdělení, žlutá oblast zabírá 5 % celkové plochy pod křivkou. V praxi ale nikdy nemáme spojitá data, pouze diskrétní naměřené hodnoty. Navíc jedním z hlavních problémů je to, že neznáme parametry rozdělení – často dokonce i předpoklad normality je příliš silný. VaR by se jinak dal zjistit i z distribuční funkce rozdělení. V hodnotě VaR je distribuční funkce rovna hladině významnosti.

Pokud má tedy portfolio hodnotu $VaR(0,01)=1\,000\,000$, můžeme říct, že s pravděpodobností 99% bude výnos portfolio vyšší než -1 000 000. Tolik ke snadné interpretaci. Již z definice je ale zřejmá jedna z největších nevýhod této metody: co když k extrémní ztrátě dojde? Z hodnoty nemůžeme určit nic o očekávané hodnotě takové ztráty – můžeme mít dvě portfolio se stejnou hodnotou VaR, pokud ale jedno z nich bude mít rozdělení s tzv. těžšími chvosty (může s vyšší pravděpodobností docházet k vyšším ztrátám), logicky bychom pro něj chtěli vyšší míru rizika.

2.3. Formální definice

Pro hladinu významnosti α z otevřeného intervalu $(0,1)$ bude Value at Risk portfolio X definována následovně:

$$VaR_{\alpha}(X) = \inf\{x \in \mathbb{R} : P(\Delta X + x < 0) \leq \alpha\},$$

kde ΔX značí změnu hodnoty portfolio X .

Vzorec zřejmě odpovídá neformální definici, která byla uvedena výše. Hledáme nejmenší x takové, že pravděpodobnost toho, že hodnota portfolio poklesne o více než x , je nižší než hladina významnosti α . Ekvivalentní je matematická definice pomocí kvantilové funkce rozdělení změny hodnoty X :

$$VaR_{\alpha}(X) = \inf\{x \in \mathbb{R} : 1 - F_{\Delta X}(-x) \geq \alpha\},$$

kde F označuje distribuční funkci rozdělení.

Problém výpočtu z těchto definic je, že vyžaduje znalost skutečného rozdělení ztrát, případně distribuční funkce. Případy, kdy rozdělení známe, jsou však velmi vzácné (pravděpodobnosti mohou být známy například u sázkových kanceláří). Ve finančním světě je v podstatě nemožné znát pravděpodobnosti jednotlivých jevů. V postupech se proto dá využít například empirických distribučních funkcí nebo Monte Carlo simulací. Možnými metodami výpočtu se budeme zabývat dále v této práci.

2.4. Využití

Očekávání od VaR je, že ukáže nějakou maximální ztrátu, ke které může s danou pravděpodobností dojít. Pokud například víme, že s 99% pravděpodobností nebude mezidenní ztráta vyšší než 500 000 Kč, můžeme si vyhradit kapitál na pokrytí takových ztrát. Problémem je, že k porušení takové mezidenní ztráty dochází (z principu pravděpodobnosti) průměrně jednou za 100 dní. Pokud bychom se tedy spoléhali na takovéto zajištění, narazili bychom na problémy v průměru téměř čtyřikrát do roka.

Další možností hodnoty VaR je oddělení hranice extrémních ztrát. To znamená, že do hodnoty VaR je možné používat pro odhad hodnoty ztráty klasické modely, neboť pro ztráty v těchto hodnotách máme dostatečné množství dat. K těmto výsledkům dochází s relativně vysokou četností, dá se proto odhadovat jejich rozdělení a pravděpodobnost. U extrémních ztrát z důvodu neočekávaných událostí už takové modely selhávají. Nedochází k nim dostatečně často, aby se dala takovými modely odhadovat jejich velikost. Pro tyto ztráty se potom mohou používat data naměřená právě o neočekávaných situacích na trhu.

Matematickou nevýhodou metody VaR je, že nesplňuje podmínku subaditivity (Daniélsson, 2005). Ta je důležitým základem pro diverzifikaci rizika, neboť vyžaduje splnění podmínky, že portfolio dvou aktiv nemůže mít vyšší riziko, než mají dohromady jednotlivá aktiva. Tento problém je ale spíše teoretický, protože podmínka je splněna pro všechna rozdělení, která jsou eliptická, tj. i pro normální. V praxi proto porušení podmínky subaditivity nedělá problémy.

Praktickou nevýhodou je, že tato hodnota nevypovídá o podstoupeném riziku tolik, kolik bychom mohli požadovat. Odděluje sice hranici extrémních ztrát, neříká ale nic o tom, co se nachází za ní. Důvodem je náročnost odhadu neobvyklých událostí. Nemáme totiž pro tyto jevy dostatečné množství měření a odhad z malého vzorku tudíž může být velmi vychýlený. Tuto vlastnost teoreticky napravuje CVaR (o této hodnotě bude více zmíněno dále v alternativách) – ta ale naráží na další technické překážky.

Právě nízká vypovídací schopnost o extrémních ztrátách je kritizována nejčastěji – například v roce 2008 americký investor a zakladatel fondu *Greenlight Capital* David Einhorn přirovnal metodu VaR k airbagu, který funguje vždy, dokud nedojde k autonehodě (Einhorn, 2008). Stejně tak tvrdil, že metoda vede k podceňování extrémního rizika a tím vlastně k podpoře investic, které mohou být velice ztrátové (i když jenom při malé pravděpodobnosti). Tento přístup je v podstatě v rozporu s jinými postupy, jako je například Markowitzova teorie portfolií. Ta při minimalizaci rizika penalizuje velký rozptyl zisků a ztrát. Velkým problémem při tomto přístupu je navíc právě korelace trhů, která (jak bude ukázáno dále) je vyšší právě při velkých šocích. Tyto velice rizikové investice tedy často selhávají společně s dalšími a tím by mohla vzniknout finanční krize, která nebyla kvůli metodě VaR předvídatelná.

Mějme například hru, ve které hodíme dvěma kostkami. Pokud hodíme dvě jedničky, zaplatíme 1000 korun. Padne-li cokoliv jiného, dostaneme 20 korun. V takovém případě bude hodnota $\text{VaR}_{0,05}$ rovna +20 korun (s pravděpodobností 95 % nevyhrajeme méně než 20 korun). Tato hodnota ale nevypovídá nic o výhodnosti či bezpečnosti investice. Střední hodnota hry je přitom přibližně -8 korun. VaR samozřejmě nemá vypovídat nic o výhodnosti investice, ale právě jeho jednoduchá definice a interpretace k tomu může svádět. Tento zjednodušený příklad navíc ukazuje na slabinu při výkladu podstoupeného rizika. Investor, řídící se metodou VaR by mohl tuto investici považovat za bezpečnou (pokud by využíval 5% hladinu významnosti – při 1% už by extrémní ztráta byla zřejmá).

2.5. Postupy výpočtu

Ačkoliv vychází z Basel II povinnost používat 10denní 99% VaR pro odhad tržního rizika (Basel Committee on Banking Supervision, 2004), neobsahuje dohoda žádnou povinnou metodu výpočtu. Pro odvození hodnoty VaR se tak může využívat několik různých postupů. Velmi rozšířená je tzv. historická metoda. Ta pro odhad využívá historicky dosažených zisků/ztrát a kvantilová funkce se aplikuje na získané empirické rozdělení. V případě nedostatečného množství dat se dá také použít Monte-Carlo simulace.

Pro historický postup je však vyžadována velmi silná podmínka, že nedojde ke změně vývoje na trhu. Při měření rizika ale většinou chceme odhadnout právě chování v případě náhlých změn, takže tento odhad nemusí být pro risk management příliš užitečný. Value at Risk je v takovém případě spíše indikativní, než aby mohl být používán pro diverzifikaci portfolia.

Kromě mezidenního VaR se někdy používá i například měsíční nebo roční VaR. Vzhledem k tomu, že pro takové měření by se hůře sbírala data, využívá se k výpočtu aproximace z výsledku pro denní VaR, který se vynásobí odmocninou z počtu dní, pro které chceme VaR odhadovat.

Tento výpočet je založen na historických datech; je tedy automaticky potvrzen empirickými výsledky. Zároveň je jeho výpočet velmi jednoduchý, nenáročný i na výpočetní techniku. Tyto výhody jsou ale vyváženy právě nedostatečností pro zajištění proti extrémním ztrátám, k čemuž by míry rizika měly sloužit.

Další možností výpočtu je takzvaný variančně-kovarianční postup. Ten odhaduje rozdělení výnosů z vlastností jednotlivých aktiv, z nichž se portfolio skládá. Metoda často předpokládá normální rozdělení výnosů. Tato podmínka je na trhu používána často, i když empirické odhady ukazují, že rozdělení je ve skutečnosti spíše platykurtické, tzn. má tlustší konce než normální rozdělení (Springer Finance, 2007). To vede ke špatnému odhadu hodnoty VaR, neboť při předpokladu normality bude VaR nižší a tím podceňujeme skutečné riziko.

Díky předpokladu normality rozdělení je potřeba odhadnout pouze dva parametry – střední hodnotu (ta se odhadne jako průměrný výnos aktiv) a rozptyl (volatilitu). Rozptyl je tvořen nejen volatilitou jednotlivých aktiv, ale také jejich vzájemnou korelací. Pro odhad volatility se může využívat jednoduše historicky naměřená směrodatná odchylka, nevýhoda tohoto postupu však spočívá v tom, že volatilita tržních aktiv většinou není konstantní – spíše naopak, rozdělení bývá tzv. heteroskedastické (Martin & Klemkovsky, 1975).

Pokročilejší metody pro odhad volatility jsou regresní přístupy. Právě z důvodu heteroskedasticity se často používá metody GARCH. Ta zahrnuje pro volatilitu reakci na předchozí odchylku od střední hodnoty. To zjednodušeně znamená, že pokud dojde k většímu výkyvu, bude na nějakou dobu volatilita větší než obvykle. V praxi je to často vidět při výkyvech na trzích, kdy dochází k prudkým pohybům z důvodu zvýšení objemu obchodování. Použití tohoto modelu už klade větší nároky na výpočet, je ale velmi objektivní a pro odhady rizika daleko schopnější než prosté využívání historických výsledků portfolia. V zájmu zachování relativní jednoduchosti se většinou používá již zmíněný GARCH(1,1) – modely vyšších stupňů jsou náročnější na data. To vede k nutnosti používat ještě starší data, která nemusí správně popisovat aktuální stav na trhu (například jsou data měřena v době před krizí, před důležitou změnou v zákonech apod.) – snaha o zvýšení přesnosti tak může mít opačné následky.

Místo metody GARCH se také dá využít exponenciálních klouzavých průměrů (EWMA). Jejich výhodou je, že se dají nastavit váhy tak, aby čerstvější výsledky měly větší dopad a výsledek je tím užitečnější aktuálně. Flexibilita těchto vstupů má ale za následek snadné manipulování s výsledky a subjektivitu výsledného odhadu VaR. Problémem obou těchto metod ale je, že jsou schopny vysvětlovat nebo predikovat

chování pouze krátkodobě. To však pro potřeby VaR většinou stačí, protože hodnota se dá pravidelně přepočítávat a používá se převážně pouze 10-denní nebo měsíční odhad.

Další možností výpočtu je přímo odhad rozdělení zisků a ztrát aktiva. Využíváme zde znalosti historických hodnot, ale nevyžadujeme podmínku normality rozdělení dat (která často nemusí být splněna). Tím je postup daleko volnější a v případě porušení normality je tato metoda i přesnější. Zřejmou nevýhodou je vyšší náročnost na množství naměřených dat. To se stává problémem především při námi řešené situaci – pokud využíváme příliš mnoho dat, bereme v potaz i staré případy, od kterých se již chování trhu změnilo, a tedy dojdeme k nesprávným výsledkům. Pokud však použijeme méně dat, bude odhadu chybět přesnost a konzistence (tedy může dojít k nesprávnému odhadu z důsledku toho, že „máme smůlu“ na naměřená data).

Do této skupiny metod patří především tzv. jádrový odhad hustoty (v angličtině *kernel density estimation* – proto někdy nazývaný kernelovským odhadem). Tento postup spočívá v tom, že vytvoříme odhadovou funkci:

$$\hat{f}_h(x) = \frac{1}{nh} \sum_{i=1}^n K\left(\frac{x - x_i}{h}\right),$$

kde K označuje jádro. Jde o funkci, která je nezáporná a její integrál je 1 (to odpovídá hlavní vlastnosti hustoty). $h > 0$ je vyhlazovací parametr. Tento parametr je libovolný a jeho hodnota ovlivňuje přesnost a rozptyl odhadu. Pokud je tento parametr příliš nízký, funkce bude obsahovat velké skoky v hodnotách měření. Vysoký parametr potom způsobí to, že funkce bude zanedbávat tvar skutečné hustoty v zájmu velké „hladkosti“ funkce – tj. nízké hodnoty derivací. Optimální hodnota parametru se dá vybírat pomocí Fisherovy informační funkce (pokud tuto vlastnost známe alespoň přibližně pro odhadovanou hustotu) nebo střední integrované odchylky.

Postup jádrového odhadu je zjednodušeně takový, že pro každé měření (s hodnotou x_i) vytvoříme funkci pravděpodobnostní hustoty rozdělení, které má střední hodnotu v x_i a jeho rozptyl je závislý na h . Potom všechny tyto funkce sečteme a výsledek vydělíme počtem měření. Výsledná hustota tak bude mít integrál roven jedné a stejně tak se zachová nezápornost.

Při velmi nízkém množství dat nemusí být ani jádrový odhad optimální – jeho odhad nemusí splňovat některé očekávané vlastnosti rozdělení ztrát. Mezi ně patří například unimodalita, hodnota šikmosti, hladkost a další. Alternativním postupem

v takovém případě mohou být tzv. epi-spliny (Royset & Wets, 2014). Epi-spliny jsou využívány k odhadování hustot pomocí zdola polospojitéch funkcí. Předpona *epi-* pochází z pojmu epikonvergence (konvergence epigrafů), která je alternativou konvergence pro funkce, které nemusí být spojité (stačí pro ně právě dolní polospojitosť). Spliny jsou křivky využívané k prokládání dat, jejichž hlavní vlastností je spojitost derivací až do určitého řádu. Na rozdíl od splinů je možno pro epi-spliny relativně snadno vybrat vlastnosti které očekáváme od hledané funkce. Těmi může být kromě zmíněných například spojitost (většina funkcí hustot), monotonie (například exponenciální rozdělení), omezení Fisherovy informace (v podstatě nahrazuje hladkost) a další. Využitím této tzv. měkké informace (*soft information*) získáváme velmi dobré odhady hustot už pro malé množství dat. Se stoupajícím množstvím měření již odhady ztrácejí sílu a jejich konvergence ke správnému řešení je pomalejší než u jádrového odhadu.

Epi-spliny by tedy byly dobře využitelné v době po šoku na trhu, pokud předpokládáme, že efekt hromadění volatility je skutečný. Například bychom mohli v bodě porušení VaR začít měřit data od začátku a po naměření 5-10 dat si při odhadu rozdělení pomoci konstrukcí epi-splinu. Například (Royset & Wets, 2014) doporučují místo klasického epi-splinu použít dokonce exponenciální epi-spline. Důvodem je to, že většina běžně odhadovaných hustot patří do rodiny exponenciálních rozdělení (*exponential family*). Jde například o binomiální, normální, log-normální, exponenciální nebo gama/beta-rozdělení. Zejména normální a log-normální rozdělení se ve finančním světě vyskytují často. Exponenciální transformace epi-splinů by tedy měla vést k rychlejší konvergenci a tím bychom ještě snížili potřebné množství naměřených dat. Nevýhodou oproti kernelovskému postupu je však kromě pomalejší konvergence i vyšší výpočetní náročnost a citlivost na nesprávnost měkké informace. Metoda je navíc relativně nová a při odhadu VaR je potřeba i korelace, která by vyžadovala zobecnění metody pro vícerozměrná rozdělení. Tato oblast ale ještě nebyla dostatečně prozkoumána.

2.6. Alternativy VaR

Nedostatek, že hodnota VaR neříká téměř nic o dosažené ztrátě v případě, že k extrémní ztrátě dojde, se dá vyřešit malou změnou. Tou je, že v případě překročení hladiny α zjistíme průměrnou dosaženou ztrátu. Tím dostaneme CVaR (Conditional

Value at Risk, někdy také Expected Shortfall). Ten například pro hodnotu $\alpha=5\%$ ukazuje, k jaké dojde průměrně ztrátě v nejhorších 5 % případech. Tím je napraven problém VaR u rozdělení s těžkými chvosty. Tato hodnota totiž vypovídá o všech hodnotách extrémních ztrát.

CVaR je navíc koherentní míra rizika, neboť na rozdíl od VaR splňuje podmínku subaditivity. Problémem však je daleko vyšší nepřesnost při jeho výpočtu. Jak už bylo řečeno, ztráty za hodnotou VaR je těžké odhadovat z důvodu jejich vzácnosti. Výpočet CVaR tedy převážně vychází ze znalosti rozdělení a předpokladu, že opravdu jde o rozdělení ztrátové funkce. Malé posuny ve velikosti extrémních ztrát už mohou vést k nezanedbatelným změnám hodnoty CVaR. Po odhadnutí velikosti těchto nepravděpodobných ztrát by navíc čekala ještě náročnější výzva, a to odhad jejich pravděpodobností. Odhadovat pravděpodobnost například náhlého krachu na burze je ale v podstatě nemožné. Právě proto se mnohem častěji využívá méně vypovídající hodnota VaR. Pro potřeby risk managementu totiž stačí tato „cenzurovaná“ hodnota a místo práce s odhadem CVaR se používá očekávaná výnosnost portfolia (která je v podstatě rovna CVaR při na 0% hladině významnosti). Ta by totiž měla zahrnovat i extrémní ztráty, jejichž neznalost je v podstatě vyvážena neznámostí extrémních výnosů na druhé straně.

3. Modely s podmíněnou heteroskedasticitou

Při modelování finančních časových řad v této práci budeme využívat klasický předpoklad homoskedasticity. To znamená, že nepodmíněná volatilita bude konstantní. Bude tedy platit, že $E(\sigma_t)$ je konstantní pro všechna t . To dává smysl, neboť bez předchozích informací nemůžeme předem předpovědět změny ve volatilitě na základě toho, jaký je čas. Jak už jsme ale zmínili dříve, velmi důležitou součástí bude reakce na období zvýšené/snížené volatility v pohybech hodnot aktiv. To zajišťuje naopak podmíněná volatilita, která už bude vykazovat změny. Tyto řady proto budeme nazývat podmíněně heteroskedastické a hodnota $E(\sigma_t|F_{t-1})$, kde F_t značí σ -algebru v čase t (tedy v podstatě všechny informace do času t), bude odlišná od hodnoty $E(\sigma_t)$ a nebude ani konstantní.

Tato podmíněná heteroskedasticita umožňuje tvorbu daleko flexibilnějších řad, které navíc mohou využívat známé informace. To vede k velkému posílení především za účelem předpovědi – to je hlavní výhoda na trhu financí a risk managementu, neboť odhad budoucí volatility je jedna z jeho hlavních součástí; zejména pro výpočet hodnoty VaR, kterou se bude tato práce zabývat.

Jednou z možností, jak zapojit podmíněnou heteroskedasticitu do modelu časové řady, je modelovat volatilitu pomocí nějaké vnější proměnné. Mohli bychom například řadu modelovat jednoduše jako:

$$y_t = \varepsilon_t * x_{t-1} + \alpha,$$

kde α je střední hodnota řady y , vnější proměnnou jsme označili x a ε je řada se střední hodnotou 0 a nekonstantním rozptylem. Tím jsme vytvořili heteroskedastickou řadu, jejíž volatilita je závislá na nějakém vnějším faktoru. Tento postup ale není příliš využíván, neboť volatilita sama o sobě je málokdy přímo závislá na vnějším vlivu. Důvody měnící se volatility jsou šoky, které ve většině případů přicházejí pokaždé z jiného zdroje (náhlá změna poptávky, změna v diplomatických vztazích, legální změny, ...). Tyto změny je navíc i jednotlivě téměř nemožné nějak racionálně ohodnocovat a zapojovat do matematických modelů. Modely se proto nesnaží odhadovat prvotní šok. Jeho predikce vyžadují spíše subjektivní přístupy a matematika se zapojuje až do odhadu postupného rozpuštění a kolísání vzniklé volatility.

Můžeme proto navrhnout přístup, který odhaduje budoucí volatilitu časové řady přímo z její předchozí hodnoty. Možnou změnou oproti předchozímu návrhu je nahradit vnější proměnnou přímo hodnotou časové řady, tzn.:

$$y_t = \varepsilon_t * y_{t-1} + \alpha,$$

Tato řada by ovšem měla nepodmíněnou volatilitu rovnou buď 0 nebo nekonečno – ani jeden z těchto jevů by pro modelování finančních časových řad nebyl žádoucí, protože nepodmíněný rozptyl takové řady by měl být kladný a dobře definovaný.

Robert Engle proto ve své práci (Engle, 1982) navrhl následující úpravu modelu:

$$y_t = \varepsilon_t h_t^{\frac{1}{2}},$$

$$h_t = \alpha_0 + \alpha_1 y_{t-1}^2.$$

Posloupnost ε je standardní bílý šum, tedy je v čase nekorelovaný, má nulovou střední hodnotu a jednotkový rozptyl. Tento model již zahrnuje znalosti předchozích hodnot, veškerá informace je uchována v hodnotě y_{t-1} . Pro zobecnění modelu se dá použít větší zpoždění, čímž se v druhé rovnici objeví parametry a hodnoty y až do řádu p . Takto definovaný model se nazývá ARCH(p) (autoregressive conditional heteroskedasticity) a při zapojení předpokladu normality dat (ten zjednodušuje výpočty i využití) se dá také definovat následovně:

$$y_t | F_{t-1} \sim N(0, h_t),$$

$$h_t = \alpha_0 + \sum_{i=1}^p \alpha_i y_{t-i}^2.$$

Případně se dá do modelu přidat i drift, tedy změnit u rozdělení střední hodnotu na nenulovou. Ve finančním světě by touto střední hodnotou mohl být například bezrizikový výnos r . Tato změna se potom musí odrazit i v rovnici pro volatilitu, kde je nyní potřeba střední hodnotu odečíst. Stejně tak posloupnost h_t se dá modelovat i jiným způsobem, než je lineární regrese. Tyto postupy ovšem mohou zbytečně přidat na složitosti a pokud není předpoklad podložen nějakými důvody z praxe, není důvod ke změně.

Drift je ovšem velmi důležitý a je jedním ze zásadních výsledků matematických modelů ve financích. Většinou totiž chceme odhadovat dva základní parametry portfolia: jeho očekávaný výnos $= E(y_t | F_{t-1})$ a velikost podstoupeného rizika, které zde

budeme měřit rozptylem h_t . Pro odhad očekávaného výnosu bylo vytvořeno mnoho modelů, můžeme tedy využít například lineární regresi (založenou na nějakých vnějších vlivech), ARMA model nebo ho odhadnout pomocí modelu CAPM.

ARCH je tedy model, který má nějaké vhodné vlastnosti pro naši aplikaci: jeho nepodmíněná střední hodnota i nepodmíněný rozptyl jsou dobře definovány, můžeme spočítat i jeho podmíněnou střední hodnotu za pomoci známých informací, a jeho volatilita se mění s vyššími odchylkami historických hodnot od očekávané hodnoty. Při testování v praxi se ovšem při analýze autokorelačních funkcí velmi často ukazovalo, že vyžadovaný řád modelu ARCH je příliš vysoký a pro modelování podmíněného rozptylu by bylo lepší použít model ARMA. Zapojením tohoto přístupu vzniká model GARCH.

3.1. Definice modelu EWMA

Model EWMA, jak již název napovídá, při odhadu volatility využívá klouzavého průměru. Na rozdíl od klasického klouzavého průměru ovšem zahrnuje ve výpočtu všechna doposud naměřená data. Vliv starších dat je ale redukován, aby byl zachycen efekt hromadění volatility. Tímto způsobem se tedy zvýšená volatilita z důvodu šoků postupně rozpouští, ale její vliv nikdy není nulový (i když se k 0 přibližuje exponenciální rychlostí).

Matematicky je volatilita v čase t odhadnuta vzorcem:

$$\sigma_t^2 = (1 - \lambda)\varepsilon_{t-1}^2 + \lambda\sigma_{t-1}^2,$$

kde λ je koeficient rozpouštění (*decay factor*). Ten nabývá hodnot v intervalu od nuly do jedné. Pokud by byla jeho hodnota rovna jedné, odhadujeme volatilitu jako aktuální hodnotu – nikdy tedy neočekáváme změnu volatility. V případě nulové hodnoty odhadujeme volatilitu jako chybu poslední predikce. Hodnota tohoto parametru určuje rozložení vlivu měření na predikci volatility mezi nová/starší data. Čím je jeho hodnota bližší nule, tím vyšší je vliv nejnovějších odchylek v predikci a tím bude graf modelu volatilnější, s většími skoky mezi jednotlivými měřeními. Naopak pokud se hodnota bude blížit k 1, bude graf hladší, neboť dřívější data budou mít vyšší vliv na budoucí predikci. Tím pádem bude vliv šoků nižší.

Epsilon značí inovace, tedy rozdíl mezi předchozí hodnotou časové řady a její poslední predikcí. Vliv na volatilitu tedy ukazuje zvýšení volatility v dobách po nějakém šoku, který vychýlil časovou řadu od původní predikce. Zde právě jde o vlastnost *volatility clustering*.

Při iterativním dosazování za volatilitu na pravé straně se vyjádření modelu dá také převést na následující tvar, který je vzorcem pro rychlý výpočet odhadu aktuální volatility:

$$\sigma^2 = \sum_{i=0}^{\infty} (1 - \lambda)\lambda^i (r_{t-i-1} - \bar{r})^2.$$

Z tohoto tvaru je zřejmé, že jde opravdu o exponenciálně vážený průměr, neboť součet jednotlivých vah λ_i je roven $\frac{1}{1-\lambda}$, což je potom vyváženo prvkem $(1 - \lambda)$ a součet vah je tak roven jedné. Nejvyšší váhu má vždy poslední inovace, tato váha je rovna právě $(1 - \lambda)$. Pokles těchto vah je potom exponenciální.

Samozřejmě jde pouze o teoretický vzorec, protože nemáme k dispozici nekonečnou historii naměřených volatilit. V horní hranici sumy přes i bychom tak měli $t-2$ a potom by nám přibyl prvek σ_0 , jehož vliv je v původním nekonečném tvaru limitně roven 0. Při výpočtu v praxi je jeho vliv sice nenulový, ale při dostatečném množství dat zanedbatelný. Inicializuje se většinou nějakým odhadem nepodmíněné volatility. Může jít například o směrodatnou odchylku naměřených dat. Místo škálovacího koeficientu $(1 - \lambda)$ by potom měl být ve vzorci skutečný součet použitých vah. V opačném případě totiž součet vah bude nižší než jedna a tím se poruší předpoklad nepodmíněné homoskedasticity, řada totiž ztratí martingalovou vlastnost a stane se z ní supermartingal. Při použití dostatečného množství dat ovšem rozdíl mezi skutečným součtem a škálovacím faktorem $(1 - \lambda)$ bude natolik nízký, že je možné ho zanedbat.

Tato metoda se pro měření a modelování volatility poprvé objevila v roce 1996 v rámci metodologie RiskMetrics, vyvíjené tehdy pod skupinou J.P. Morgan (RiskMetrics Technical Document, 1996). Vzhledem k relativní jednoduchosti a užitečnosti se postup pomocí této metody využívá dodnes. Výpočet nevyžaduje žádné optimalizační metody a tím je postup velmi snadný. Optimalizací se sice dá určit člen λ , který určuje míru snižování vlivu historické volatility s časem, v modelu RiskMetrics se ale nastavuje konstantní hodnota 0,94 pro obchodní účely (denní měření) a 0,97 pro investiční postup (měsíční měření). V našem postupu se pokusíme optimalizaci využít a najít takový člen λ , který bude pro naměřená data nejlepší.

Protože budeme odhadovat hodnotu VaR u portfolia složeného z většího množství aktiv, budeme kromě jednotlivých volatilit potřebovat odhadnout i korelace. Ačkoliv se dá použít konstantní odhad korelačních koeficientů pro zachování jednoduchosti výpočtu, pro zpřesnění se využívá metoda EWMA i pro odhad korelací mezi časovými řadami vývoje jednotlivých aktiv. Ty se odhadují velmi podobně jako volatilita a slouží k tomu následující vzorec:

$$\sigma_{12,t+1|t}^2 = \lambda \sigma_{12,t|t-1}^2 + (1 - \lambda) \varepsilon_{1,t} \varepsilon_{2,t},$$

který se dá opět přepsat na tvar:

$$\sigma_{12,t+1|t}^2 = (1 - \lambda) \sum_{j=0}^{\infty} \lambda^j \varepsilon_{1,t-j} \varepsilon_{2,t-j}.$$

Tím dostaneme odhady kovarianční koeficienty mezi jednotlivými řadami. Pro odhad korelací je nutno je ještě vydělit směrodatnými odchylkami těchto řad, tedy:

$$\rho_{12,t+1|t} = \frac{\sigma_{12,t+1|t}^2}{\sigma_{1,t+1|t}\sigma_{2,t+1|t}}.$$

Koeficient lambda ve vzorci pro odhad kovariance znovu značí *decay factor*, používá se stejná hodnota jako u odhadu volatility, tedy 0,94.

3.2. Definice modelu GARCH(p,q)

Model GARCH si nejdříve definujeme pro obecné parametry p, q ; ve zbytku práce se však budeme věnovat modelu GARCH(1,1). Ten již v sobě zahrnuje celou budoucnost, a ačkoliv zvýšení některého řádu může zpřesnit výsledky, cenou je výpočetní náročnost a může dojít i předimenzování parametrů. To je stav, kdy volitelných parametrů je příliš mnoho v porovnání s použitými daty – sice se zvyšuje přesnost na historických datech, ale predikční síla může zůstat stejná, nebo může dokonce dojít ke zhoršení. Více informací o definici a vlastnostech modelu GARCH může být nalezeno například v (Bollerslev, 1986)

V modelu GARCH je základem rovnice pro odhad střední hodnoty (ve finančním světě například logaritmických výnosů):

$$y_t = \mu_t + a_t,$$

kde μ_t je podmíněná střední hodnota y_t . Jak již bylo řečeno, tato hodnota může být modelována například vlastní řadou ARMA, může být nastavena jako konstanta – k té se většinou dojde pomocí CAPM modelu, v takovém případě by byla rovna $r_F + \beta^*(r_M - r_F)$, kde r_F je bezrizikový výnos, r_M výnos tržního indexu, a β je beta koeficient daného aktiva.

a_t je potom odchylka skutečného výnosu od očekávané hodnoty. Na rozdíl od klasických modelů střední hodnoty tato chyba nebude mít u modelu GARCH rozdělení $N(0, \sigma^2)$, ale bude pro ni platit následující rovnice:

$$a_t^2 = \sigma_t^2 z_t^2,$$

kde z_t je standardizovaný bílý šum (tedy nezávislé, stejně rozdělené veličiny s nulovou střední hodnotou a jednotkovým rozptylem) a h_t je právě podmíněná volatilita, která bude modelována právě již uvedeným modelem ARMA(p,q):

$$\sigma_t^2 = \alpha_0 + \sum_{i=1}^p \alpha_i \sigma_{t-i}^2 + \sum_{j=1}^q \beta_j a_{t-j}^2,$$

kde $\omega, \alpha_i, \beta_j$ jsou parametry, které budeme v modelech odhadovat. Vidíme, že v rovnici se objevuje jak předchozí hodnoty volatility (σ), tak i předchozí hodnoty inovací (a). V dobách zvýšené volatility se tak změna projeví přes hodnoty volatility i inovací. Jak už bylo řečeno, je to za účelem uvedení celé známé historie do výpočtu.

Tento model je velmi podobný modelu EWMA. Hlavním rozdílem je přidání stochastického prvku v podobě bílého šumu. Ten výrazně zvyšuje náročnost, neboť na rozdíl od deterministického modelu EWMA je nyní potřeba využívat simulace a podobné metody, jaké se používají při vyhodnocování stochastických modelů.

Takto je definován tzv. jednorozměrný GARCH. Ten je aplikovatelný pro jednu řadu a počítá pouze s její volatilitou. Protože však budeme chtít odhadovat chování několika různých řad, které mezi sebou mohou být (a nejspíše budou) korelovány, bude potřeba definovat ještě o něco složitější, vícerozměrný model.

Dále tedy bude model GARCH ještě mírně poupraven. Nebudeme totiž modelovat GARCH pro jedinou řadu, ale pro větší množství řad, a kromě jednotlivých volatilit se v neznámých parametrech objeví i korelace mezi volatilitami jednotlivých řad. Pokud jsou tyto korelace konstantní, přístup se nazývá CCC (constant conditional correlation) GARCH. V tomto postupu se volatility modelují pomocí matic jako součin směrodatných odchylek a korelací jednotlivých aktiv v portfoliu. Vznikne tak následující model:

$$y_t | F_{t-1} \sim N_n(0, \Sigma_t),$$

kde N_n značí n -rozměrné normální rozdělení se Σ_t jako maticí rozptylu. Tato matice je potom generována pomocí matic korelací a směrodatných odchylek následovně:

$$\Sigma_t = \sigma_{t|t-1} \rho \sigma_{t|t-1}$$

kde ρ značí matici korelací, která je symetrická, pozitivně semidefinitní a její prvky jsou mezi hodnotami -1 a 1 (z vlastností korelace). V tomto případě je matice korelací konstantní (proto CCC). $\sigma_{t|t-1}$ značí diagonální matici, jejíž prvky na diagonále jsou podmíněné směrodatné odchylky jednotlivých řad. Z indexu je zřejmé, že tato matice není v čase konstantní, ale bude modelována právě řadou ARMA, jako u klasického jednorozměrného GARCH. Model tedy bude vyžadovat ještě n rovnic ve tvaru:

$$\sigma_{k,t|t-1}^2 = \omega + \sum_{i=1}^p \alpha_i \sigma_{k,t-i|t-i-1}^2 + \sum_{j=1}^q \beta_j y_{t-j}^2$$

Pokud znovu budeme modelovat střední hodnotu řady y_t jiným způsobem než konstantní nulou, musíme znovu do modelu zapojit místo hodnot y_t inovace a_t . Nyní jsme tedy vytvořili první model, který je využitelný pro naše potřeby.

Problémem tohoto modelu je, že vyžaduje konstantní korelace v čase. Jak ale ukázaly empirické výzkumy a dá se snadno ukázat z historických dat, korelace mezi trhy se v časech krize zvyšují. Protože chceme modelovat právě reakce trhů na náhlé šoky, nemůžeme předpokládat matici podmíněných korelací jako konstantní. Pro tento případ vznikl model DCC (dynamic conditional correlation) GARCH. Ten již nepovažuje matici ρ za konstantní, ale modeluje ji další model. Engle a Sheppard ve své práci (Engle & Sheppard, Theoretical and Empirical properties of Dynamic Conditional Correlation Multivariate GARCH, 2001) navrhli matici modelovat ze tří různých částí: prvním zdrojem je matice nepodmíněných směřodatných odchylek z jednorozměrných modelů, druhou částí je náhodný vektor bílého šumu, generovaný z normálního rozdělení a poslední součástí jsou předchozí hodnoty korelační matice.

Tato matice potom splňuje důležitou podmínku pozitivní semidefinitnosti, navíc je při výpočtu upravena tak, aby splňovala i ostatní nutné podmínky pro korelační matici.

Pro zjištění parametrů modelu se ve většině případů využívá klasické metody maximální věrohodnosti. Po naměření hodnot tedy hledáme takové parametry, při kterých by byla tato kombinace měření nejpravděpodobnější. Protože v modelu DCC GARCH je množství potřebných parametrů vysoký a jsou na sobě vzájemně závislé kvůli provázanosti jednotlivých rovnic, je běžně používán dvoufázový přístup. Při něm se nejdříve odhadnou parametry pro jednotlivé rovnice. Ty se pak využijí ve druhé fázi, kdy se odhadnou vstupní parametry pro dynamický model korelační matice.

Další možnou metodou, kterou lze využít pro odhad parametrů, je bayesovský přístup. Ten využívá apriorní informace o rozdělení některých veličin, aby našel nejpravděpodobnější rozdělení odhadovaných parametrů. Jednou z velkých výhod bayesovského přístupu je jednoduchá možnost aktualizace výsledků. Při naměření dalších dat a tím získání nových informací lze získat lepší výsledky. Nevýhodou tohoto přístupu je velmi náročný analytický výpočet. Většinou se proto využívá simulačních postupů, nejznámějšími jsou takzvané metody Markov Chain Monte Carlo (MCMC), mezi nimi například Gibbsův výběrový plán nebo Metropolis-Hastingsův algoritmus. Tyto postupy umožňují generovat i ze složitých náhodných rozdělení za pomoci markovských řetězců. Pro bližší informace k postupu viz například (Rachev, Hsu, & Bagasheva, 2008).

Gibbsův výběrový plán je využíván pro mnohorozměrná rozdělení. Při jeho postupu se nejdříve vygeneruje náhodný startovací vektor. Jeho hodnoty se potom postupně aktualizují tak, že jednotlivé hodnoty vektorů se generují z rozdělení, které je podmíněno aktuální hodnotou vektoru. Máme tedy například vektor $\mathbf{X}$.

Začneme v bodě $X^{(0)} = (x_1^{(0)}, x_2^{(0)}, \dots, x_n^{(0)})$, který může být zvolen náhodně, z nějaké předchozí informace, nebo pomocí tzv. burn-in, kdy markovský postup proběhne velké množství řetězců za účelem vyřazení chyby z konečného výsledku (aby výsledné rozdělení nevykazovalo možnosti, které jsou způsobeny pouze špatnou volbou startovacího bodu).

V prvním kroku vygenerujeme $x_1^{(1)}$ z rozdělení $p(x_1^{(1)} | x_2^{(0)}, \dots, x_n^{(0)})$. V dalším kroku již využijeme aktualizovanou hodnotu a generujeme $x_2^{(1)}$ z rozdělení $p(x_2^{(1)} | x_1^{(1)}, x_3^{(0)}, \dots, x_n^{(0)})$. Tento krok potom opakujeme v libovolném množství, které musí být dostatečné pro zajištění co nejlepší konvergence. Pokud se iterace provede v příliš malém počtu, nebude rozdělení dostatečně blízké hledanému. Větší množství naopak může vést k vysoké operační a časové náročnosti.

Metropolis-Hastingsův algoritmus je obecnější verze Gibbsova výběrového plánu. Ten se totiž shoduje v tom, že aktualizuje bod vektoru. Tato aktualizace je poté ale prověřena tzv. pravděpodobností přijetí. Pokud je vygenerovaný bod pravděpodobnější než předchozí, je přijat. Pokud je méně pravděpodobný, je přijat právě s pravděpodobností přijetí – ta je rovna poměru jejich pravděpodobností.

Dalším rozdílem tohoto postupu je, že hustota, pomocí které se generují aktualizované body, může být téměř libovolná. Pokud je však špatně vybrána, může tento postup skončit nesprávným výsledkem, případně může běžet příliš dlouho a stále nevygenerovat hodnoty ze správného rozdělení.

3.2.1. Odhad pomocí maximální věrohodnosti

Při odhadu pomocí maximální věrohodnosti je nejdříve potřeba nalézt věrohodnostní funkci. Pro veličinu se standardním normálním rozdělením je věrohodnostní funkce ve tvaru:

$$L = \prod_{t=1}^T \frac{1}{(2\pi)^{\frac{n}{2}}} e^{-\frac{x^T x}{2}}$$

V našem modelu ale místo standardního rozdělení máme rozptylovou matici rovnou Σ .

To znamená, že $y_t | F_{t-1} = z_t \Sigma_t^{1/2}$. Potom z_t je již standardně normálně rozděleno.

Provedeme proto transformaci náhodné veličiny, jejíž jakobián je roven $|\Sigma_t^{1/2}|$. Výsledná věrohodnostní funkce tedy získává tvar:

$$L = \prod_{t=1}^T \frac{1}{(2\pi)^{\frac{n}{2}} |\Sigma_t^{1/2}|} e^{-\frac{y_t^T \Sigma_t^{-1} y_t}{2}}$$

Při výpočtu maximální věrohodnosti se z důvodu zjednodušení derivace nevyužívá přímo věrohodnostní funkce, ale její logaritmus. Tím se součin převede na součet a vznikne následující rovnice:

$$l = -\frac{1}{2} \sum_{t=1}^T (n \ln(2\pi) + \ln(|\Sigma_t|) + y_t^T \Sigma_t^{-1} y_t)$$

Nyní ve vzorci nahradíme výraz z definice korelační matice a získáme výsledek:

$$l = -\frac{1}{2} \sum_{t=1}^T (n \ln(2\pi) + 2 \ln(|\sigma_t|) + \ln(|\rho_t|) + y_t^T \sigma_t^{-1} \rho_t^{-1} \sigma_t^{-1} y_t)$$

Odhady parametrů by z tohoto výsledku byly stále velmi náročné, vzhledem k velkému množství násobených parametrů. DCC-GARCH ale umožňuje využívat výše zmíněný dvoufázový postup, ve kterém se v první fázi místo korelační matice dosadí matice identity I_t . Jde o čtvercovou matici, která má na diagonále 1 a všude mimo ni jsou hodnoty nulové. Tím zmizí všechny mimodiagonální prvky ve vzorci a funkce se zjednoduší na:

$$l = -\frac{1}{2} \sum_{t=1}^T (n \ln(2\pi) + 2 \ln(|\sigma_t|) + \ln(|I_t|) + y_t^T \sigma_t^{-1} I_t^{-1} \sigma_t^{-1} y_t) =$$

$$l = -\frac{1}{2} \sum_{t=1}^T (n \ln(2\pi) + \sum_{i=1}^n (\ln(\sigma_{i,t}^2)) + \frac{y_{i,t}^2}{\sigma_{i,t}^2})$$

Dolním indexem i se zde značí i -tý prvek vektoru, případně i -tý prvek na diagonále kovarianční matice. Díky nahrazení korelační matice jsou nyní jednotlivé rovnice na sobě výpočetně nezávislé, a proto lze úlohu řešit jako n jednorozměrných modelů GARCH. Potom zbývá dopočítat pouze parametry pro výpočet skutečné korelační matice. Využijeme definice DCC matice z (Engle & Sheppard, 2001). Ti navrhli znovu použít ARMA model (mírně upravený) s následujícími specifikacemi:

Nejdříve mějme matici Q_t , definovanou následovně:

$$Q_t = \left(1 - \sum_{m=1}^M a_m - \sum_{n=1}^N b_n\right) \bar{Q} + \sum_{m=1}^M a_m (\varepsilon_{t-m} \varepsilon'_{t-m}) + \sum_{n=1}^N b_n Q_{t-n}$$

$\bar{Q}$ zde značí nepodmíněnou kovarianci chyb ε spočtenou z první fáze postupu; M a N jsou řády ARMA modelu; a, b jsou hledané parametry. Po definici Q_t potom matici korelací σ_t získáme následovně:

$$\sigma_t = Q_t^{*-1} Q_t Q_t^{*-1}$$

Q_t^{*-1} v tomto případě značí diagonální matici, jejíž prvky jsou odmocniny z diagonálních prvků matice Q . To je provedeno za účelem standardizování matice a její úpravu, aby absolutní hodnota všech jejích prvků byla menší než 1. I korelační matice je tedy v podstatě modelována řadou ARMA, i když s drobnou úpravou.

Vzniká nám tak model, ve kterém se pomocí modelu ARMA modelují jak jednotlivé volatility řad, tak jejich vzájemná korelace. Navíc může být model ARMA použit i pro odhad střední hodnoty, čímž získáváme tři různé řady, vyžadující odhady parametrů.

Když už známe odhadnuté parametry z prvního kroku, vložíme je nyní zpět do původní věrohodnostní funkce, kde už matici korelací nebudeme nahrazovat identitou. Vznikne nám tak poslední potřebná sada rovnic:

$$l_2 = -\frac{1}{2} \sum_{t=1}^T (n \ln(2\pi) + 2 \ln(|\sigma_t|) + \ln(|\rho_t|) + y_t^T \sigma_t^{-1} \rho_t^{-1} \sigma_t^{-1} y_t)$$

$$l_2^* = -\frac{1}{2} \sum_{t=1}^T (\ln(|\rho_t|) + \varepsilon_t^T \rho_t^{-1} \varepsilon_t)$$

Využitím toho, že matice σ_t je známá a konstantní díky parametrům, zjištěným z první fáze postupu, a odstraněním všech konstant z funkce se nám podařilo zjednodušit funkci na l_2^* , jejíž maximalizací se dají dopočítat ARMA parametry potřebné k vymodelování matice korelací pro naši časovou řadu.

3.2.2. Bayesovský přístup

Základní rozdíl bayesovského a frekventistického přístupu spočívá v přístupu k odhadovaným parametrům. Ve frekventistickém přístupu máme nějaké podmíněné rozdělení $p(y|\theta)$ a považujeme parametr θ za fixní hodnotu, již se snažíme nalézt, nebo co nejlépe odhadnout. Na druhé straně bayesovský přístup využívá Bayesovy věty o podmíněných rozděleních, která ve vzorci vypadá následovně:

$$p(\theta|y, \alpha) = \frac{p(y|\theta)p(\theta|\alpha)}{p(y|\alpha)}$$

V tomto vzorci vidíme rozdělení parametru θ za podmínky naměřených dat y a tzv. hyperparametru α . To je parametr samotného rozdělení parametru θ , tomuto rozdělení se v bayesovské statistice říká apriorní. Cílem potom je získat rozdělení na levé straně, kterému se říká aposteriorní. Při využití věrohodnostních funkcí a vynechání hyperparametru lze vzorec přepsat na:

$$p(\theta|y) = \frac{L(y|\theta)p(\theta)}{\int L(\theta|y)p(\theta)d\theta}$$

Integrál ve jmenovateli je přitom číslo, které pouze standardizuje číselník tak, aby hustota pravděpodobnosti měla integrál roven 1. Tato „rovnost až na násobení konstantou“ se značí pomocí symbolu $\propto$. V našem vzorci tedy bude platit:

$$p(\theta|y) \propto L(y|\theta)p(\theta)$$

Toto zanedbání konstanty by se mohlo zdát jako přehnaný krok, v bayesovské statistice se ale právě i kvůli problémům s příliš náročným výpočtem tohoto integrálu využívá metod markovské simulace. Její použití umožňuje tuto škálovací konstantu vynechat a zabývat se pouze simulací ze skutečného hledaného rozdělení. Při postupu se tedy očividně využijí věrohodnostní funkce, které jsme již vyjádřili v předchozí kapitole.

Důležitou otázkou u bayesovského postupu je také využití správných apriorních rozdělení. Může se jednat o takzvané informativní apriorní rozdělení, ve kterém využíváme nějakou znalost ohledně možného výskytu parametru. Můžeme ale využít i tzv. neinformativní rozdělení. V takovém případě využíváme například rovnoměrné rozdělení a markovský řetězec potom musí dobře vygenerovat hledané aposteriorní rozdělení. Problémy potom mohou vzniknout, pokud je apriorní rozdělení špatně určeno (například rovnoměrné rozdělení nemusí být definováno na správném intervalu). Může

vést buď ke konvergenci k nesprávnému rozdělení, nebo může řetězec sice konvergovat správně, ale příliš pomalu.

3.2.3. Bayesovský odhad GARCH(1,1)

Při odhadu si nejdříve musíme vybrat apriorní rozdělení pro parametry $\alpha_0, \alpha_1, \beta$. Jako nejjednodušší a vhodné pro praxi se nabízí normální rozdělení. Koeficienty α_0, α_1 si nastavíme tak, aby mezi nimi byla možná korelace. Máme tedy rozdělení parametrů s hyperparametry $\Sigma_\alpha, \mu_\alpha, \mu_\beta, \sigma_\beta$:

$$p(\alpha) \propto N_2(\mu_\alpha, \Sigma_\alpha) * 1_{\{\alpha > 0\}}$$

$$p(\beta) \propto N(\mu_\beta, \sigma_\beta) * 1_{\{\beta > 0\}}$$

1 na konci obou vzorců značí indikátor toho, že jsou parametry kladné. To je jeden z našich předpokladů a tato změna se potom promítne v tom, že tyto parametry nebudou mít přímo normální rozdělení, ale pouze jeho useknutou verzi. To je potom vyváženo tím, že místo rovnosti se zde znovu objevuje proporcionalita. V modelu navíc předpokládáme nezávislost mezi parametry α a β . Pokud si pro zjednodušení zápisu vytvoříme vektor $\varphi = (\alpha_0, \alpha_1, \beta)$, můžeme aposteriorní hustotu psát ve tvaru:

$$p(\varphi | y) \propto L(y | \varphi) p(\varphi)$$

Tímto způsobem získáváme nejen střední hodnotu a směrodatnou odchylku, jako je tomu u frekventistické analýzy, ale dokonce celé podmíněné rozdělení jednotlivých parametrů. Protože jde často o rozdělení, která je náročné vypočítat analyticky, využívá se metod markovské simulace a celkové rozdělení je potom odhadnuto pomocí pravděpodobnosti výskytu jednotlivých hodnot v řetězci.

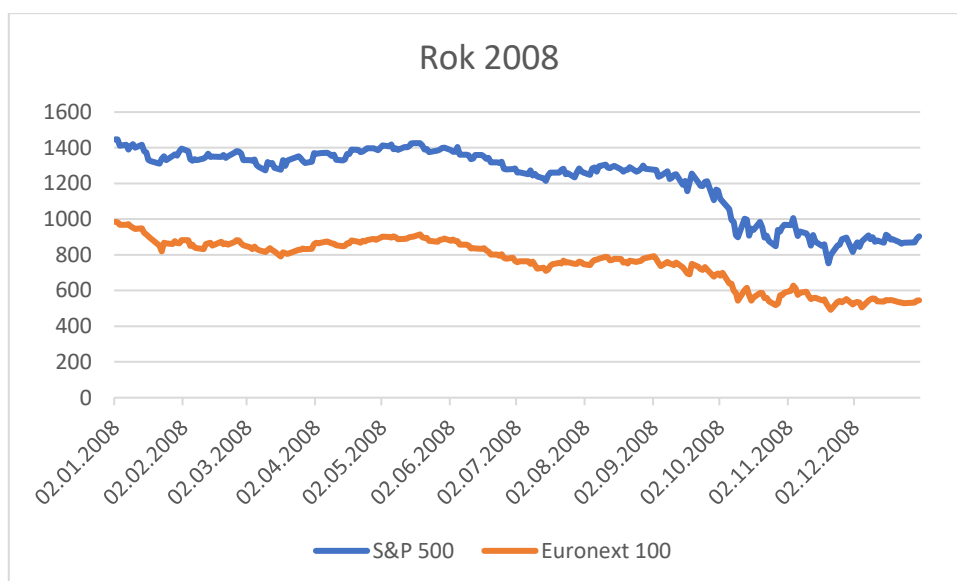
V této práci budou tyto tři parametry odhadnuty pro každou časovou řadu, tj. 9 parametrů. K tomu ještě bude potřeba odhadnout 2 parametry pro matici Q , která se objevuje v DCC-metodě. Budeme tedy generovat 11 markovských řetězců, což může mít za následek vyšší časovou náročnost výpočtu. Pokud ale dojde ke správnému výběru apriorních rozdělení, neměla by tato překážka být příliš velká.

4. Empirická studie

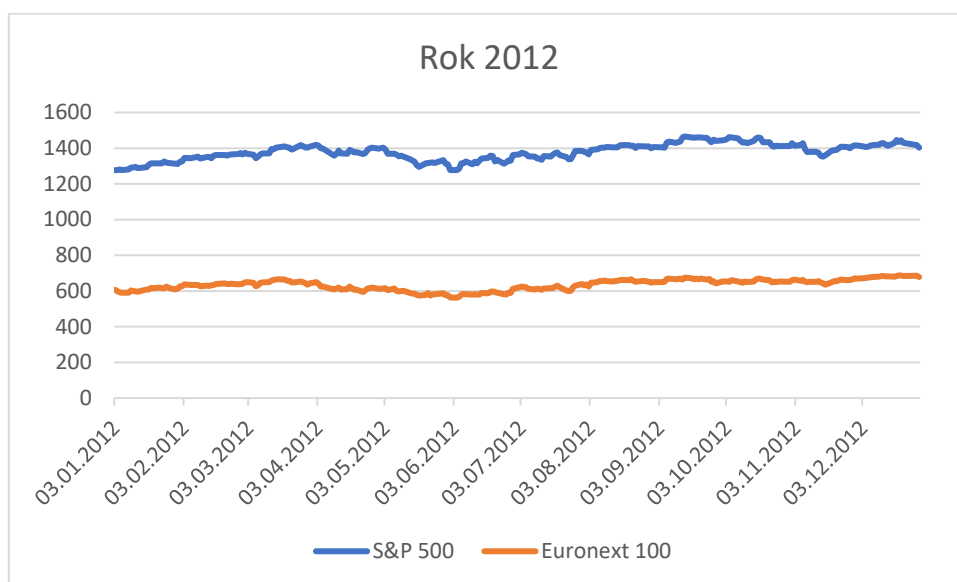
Tématem této práce bude aplikace předvedené teorie na měření korelace mezi světovými trhy. Ačkoliv v moderní době jsou trhy z velké části ovládány algoritmickým programováním a burzy jsou propojeny velmi rychlým elektronickým spojením, korelace mezi vzdálenějšími trhy stále nemusí být stoprocentní. Pro ukázkou budeme v této práci uvažovat tři velké trhy a počítat s jejich indexy. Srovnávanými indexy budou americký S&P 500, pařížský EURONEXT 100 a japonský Nikkei 225. Protože jde o velmi vyspělé a likvidní trhy, můžeme očekávat relativně nízkou volatilitu výnosů na indexech. Během šoků však právě tyto indexy dobře ukazují změny na trhu, a to jsou období, kdy by mezi trhy mohlo vést ke zvýšení volatility.

Protože jde o vyspělé trhy, i jejich provázanost by měla být vysoká. Mnoho mezinárodních firem je obchodovatelných na všech těchto trzích a pohyby indexů by proto mohly být dobře korelované. I proto byly trhy vybrány jako ukázkový příklad. Samozřejmě by se daly použít i další velké trhy, jako například indický, čínský nebo australský. Korelace už by zřejmě byla menší, ačkoliv ne nutně zanedbatelná. Na druhou stranu společné posílení korelace během šoků na trzích by se už nemuselo projevovat tak silně.

Můžeme si nejdříve názorně ukázat vliv krizového období na korelaci mezi trhy. K tomu budou využita data S&P 500 a Euronext 100 z roku 2008 (krizové období) a roku 2012 (který by měl být relativně klidný). Vývoj hodnoty těchto indexů v těchto letech je zobrazen na obrázcích 4-1 a 4-2.



Obrázek 4-1: Vývoj indexů S&P a EURONEXT v roce 2008



Obrázek 4-2: Vývoj indexů S&P a EURONEXT v roce 2012

Z obrázku 4-2 je zřejmá silná korelace během krize, relativně vysoká ale vypadá i v období klidu. Výběrové koeficienty korelace už srovnání ukazují jasně:

$$\rho_{2008}=0,977;$$

$$\rho_{2012}=0,843.$$

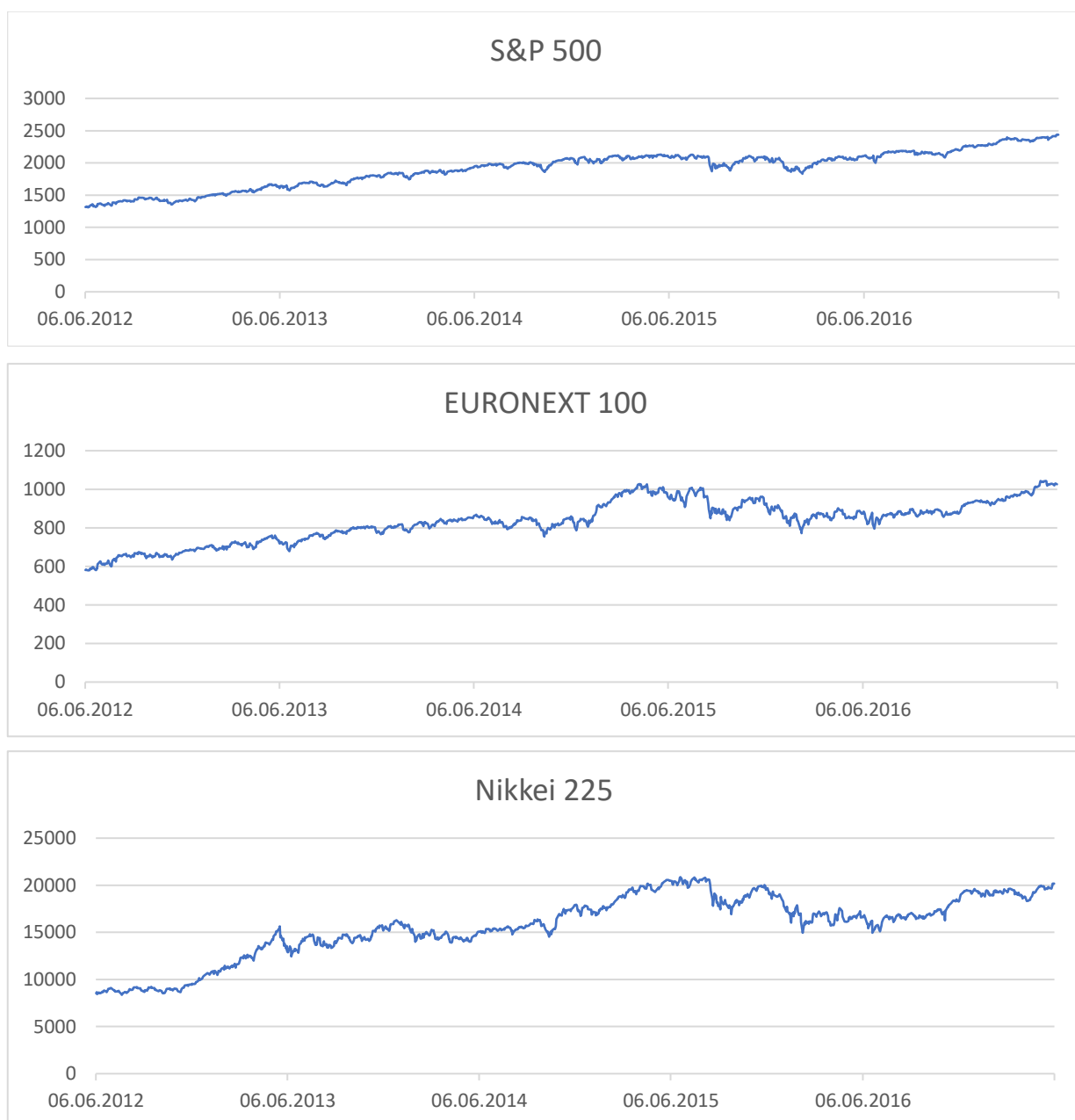
Na toto téma byly napsány mnohé studie (např. (Hyde, Bredin, & Nguyen, 2007)), tento příklad je zde pouze pro ukázkou a pro potvrzení hypotézy, která je jedním z hlavních důvodů zapojení metody GARCH jako vhodného modelu za účelem zachycení měnící se korelace.

Další věcí, kterou bude potřeba si uvědomit, je časový posun mezi trhy. Tokijská burza je otevřena od 09:00 do 15:00 japonského času (JST: +10), pařížský EURONEXT je otevřen k obchodování mezi 09:00 a 17:30 (CET: +1) tamního času a na amerických NYSE a NASDAQ se obchoduje od 09:30 do 16:00 (EST: -5). Při řešení korelace mezi trhy bude tedy problém se zvolením odpovídajícího data. Zatímco americké burzy zavírají ve 21:00 GMT, v Japonsku už zbývají jen dvě hodiny do začátku dalšího obchodního dne. V této práci se s problémem vypořádáme tak, že naměřená data indexu Nikkei posuneme o den dopředu (budeme tedy například srovnávat S&P ze dne 1.6.2017 a Nikkei ze dne 2.6.2017).

Kromě časového důvodu k tomuto posunu vede také kauzalita mezi burzami. I z důvodu objemů obchodování můžeme očekávat daleko vyšší vliv amerických burz na japonské trhy než opačným směrem. Posunuli jsme tedy v podstatě bod datového zlomu z Asie do Evropy a ve sledovaný den budou časy otevření burz: Paříž -> NY -> Tokyo. Využití tohoto postupu by měla podpořit i statistika, kdy očekáváme daleko vyšší korelaci při posunu Japonska o den. Pokud bychom chtěli použít nějaký komplexnější model, mohli bychom modelovat kauzalitu tak, že bychom váženým průměrem do modelu využili volatilitu japonských trhů ze dne před obchodováním na americké burze i po ní. Kauzalita bez posunu je však relativně nízká, takže změna by nebyla natolik silná, aby se výhoda takového složitějšího modelu nějak projevila.

4.1. Data

Data k vybraným indexům byla stažena z internetové stránky Yahoo Finance. Jde o zavírací ceny daných indexů za posledních 5 let, tedy od června roku 2012. Frekvence je denní (až na zřejmé výjimky víkendů). Vzhledem k různě definovaným pracovním svátkům v daných zemích musela být data ještě upravena tak, aby se ve všech časových řadách vyskytovaly pouze ceny ze stejných dní. Pokud tedy na některé z burz byl non-trading day, byla všechna data v souboru zanedbána. Po tomto kroku zbylo v souboru 1170 naměřených hodnot. Na obrázku 4-3 se můžeme podívat na vývoj hodnot jednotlivých časových řad.



Obrázek 4-3: Vývoj cen akciových indexů

Všechny tři grafy obsahují hodnoty v domácí měně, tedy postupně USD, EUR, JPY. To může vést k problémům z důvodu pohybů síly měn. Tyto pohyby jsou však silně provázány s pohyby akciového trhu, neměly by proto vyvolávat falešné signály. Již z grafů vidíme, že nejvyšší stabilitu vykazuje index S&P 500 (což se dá předpokládat). Protože hlavním tématem této práce je korelace, je důležité si všimnout podobností v pohybech grafů. Jasně je, že dlouhodobě se všechny trhy pohybují stejným směrem, a to nahoru.

Největší šok je vidět v druhém pololetí roku 2015. Přesněji jde o přelom srpna a září, kdy se na trzích setkala několik negativních faktorů. Jedním z nich byla kritická situace v Řecku, které 30. června vyhlásilo default na úvěr MMF. Hlavní příčinou pohybu byl ale pád akcií na šanghajské burze, který také začal již na začátku června. Došlo zde ke splasknutí bubliny, které mělo za následek několik krizových dní a burzovní index spadnul během několika týdnů až o 30 %. Z námi sledovaných indexů se logicky nejsilnější efekt objevil na tokijské burze, i na ostatních indexech jsou ale zřejmé změny. To by potvrdzovalo hypotézu zvýšených volatilit v časech krize a podobných vyšších pohybů na burze.

Abychom mohli data analyzovat, bude důležité nejprve transformovat hodnoty indexů na jejich výnosy. V práci budeme využívat logaritmické výnosy (hodnoty indexů jsou nutně vždy kladné, není tedy problém s definičním oborem logaritmu), definované následovně:

$$r_t = \ln(C_t) - \ln(C_{t-1})$$

Transformací tedy získáme mezidenní logaritmické výnosy, na které budeme aplikovat výše popsany model. Nejdříve provedeme vybrané analýzy těchto dat, abychom ověřili vhodnost daného postupu. Potřebovali bychom zjistit přibližné rozdělení těchto výnosů. Pro model GARCH mluví především případ leptokurtického rozdělení (tzv. rozdělení s těžkými chvosty).

V tabulce 4-1 si ukážeme důležité vlastnosti jednotlivých řad logaritmických výnosů:

	Průměr	Sm. Odch.	minimum	maximum	25-kvantil	75-kvantil
S&P 500	0,000528	0,00819	-0,0402	0,0447	-0,0032	0,0049
EURONEXT 100	0,000491	0,01104	-0,0696	0,0401	-0,0052	0,0064
Nikkei 225	0,000725	0,01463	-0,0825	0,0743	-0,0065	0,0087

Tabulka 4-1: Základní vlastnosti logaritmických výnosů

Naměřená data potvrzují, co bylo vidět z grafu 4-3. Tedy že japonské trhy jsou nejvolatilnější, zatímco americký index je oproti němu stabilní. Další věc, která byla zřejmá z grafu je, že všechny indexy dosahují v průměru kladného výnosu. Dále přejdeme k nejdůležitější části základní analýzy dat a tou bude odhad korelační matice pro jednotlivé páry indexů. Ta byla odhadnuta pomocí výběrových korelačních koeficientů.

	S&P 500	EURONEXT 100	Nikkei 225
S&P 500	1	0,6043	0,4073
EURONEXT 100	0,6043	1	0,3034
Nikkei 225	0,4073	0,3034	1

Tabulka 4-2: Korelace logaritmických výnosů

Tabulka 4-2 potvrzuje několik z našich očekávání. Ukazuje, že americké a evropské trhy jsou mezi sebou provázanější, než jsou jednotlivé trhy s tokijskou burzou. Všechny korelační koeficienty jsou navíc nejen kladné (takže pohyby na burzách jsou směřově sladěné), ale ještě k tomu relativně vysoké. Tím více vypovídající bude analýza pohybů těchto korelací. Pro zajímavost a potvrzení správnosti rozhodnutí posunout japonská data o jeden den si můžeme ukázat ještě jak by tabulka vypadala v případě, že bychom porovnávali data ze stejného dne (4-3):

	S&P 500	EURONEXT 100	Nikkei 225
S&P 500	1	0,6043	0,1868
EURONEXT 100	0,6043	1	0,3149
Nikkei 225	0,1868	0,3149	1

Tabulka 4-3: Korelace logaritmických výnosů po posunutí japonského indexu o 1 den

Je vidět velký rozdíl pro korelace mezi americkými a japonskými akcemi. Při porovnávání s Evropou však odhad korelace dokonce klesnul, ačkoliv rozdíl není velký. Evropský trh je totiž časově od Japonska vzdálen pouze 9 hodin, takže ačkoliv jejich obchodní dny se neprolínají, časový posun je přibližně rovnoměrný. Ve zbytku práce budeme počítat s původní, o den posunutou řadou.

Dalším krokem, který uděláme z důvodu zaměření práce, je posun výnosů. Jde nám především o modelování volatilit a nebudeme se tedy snažit najít co nejlepší model pro střední hodnotu výnosů. Taková práce by už byla nad rozsah našich cílů a kombinace nejlepšího modelu pro střední hodnotu i volatilitu může být námětem pro další studie. Všechny časové řady tedy posuneme o jejich střední hodnotu níže. Jedná se o integraci této řady. Ta nemá vliv na nepodmíněnou korelaci (jde o posun o konstantu). Hlavním cílem je, abychom částečně standardizovali řady tak, aby měly nulovou nepodmíněnou střední hodnotu. U takové řady je pak o mnoho jednodušší modelovat volatilitu. Modelujeme totiž řadu konstantní nulou. To pro surová data znamená následující:

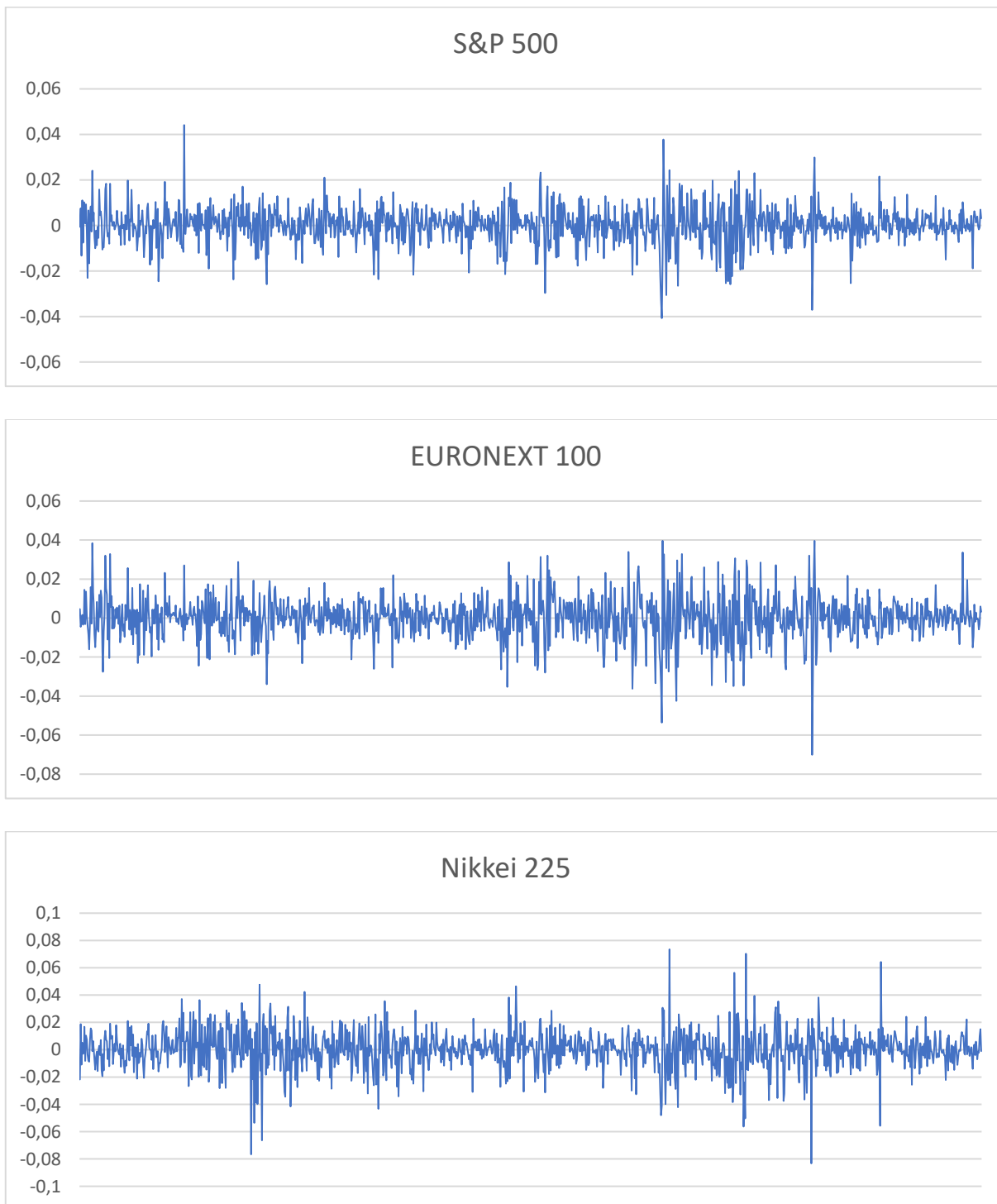
$$\ln x_t - \ln x_{t-1} = \mu,$$

$$e^{\ln x_t - \ln x_{t-1}} = e^\mu,$$

$$\frac{x_t}{x_{t-1}} = e^\mu.$$

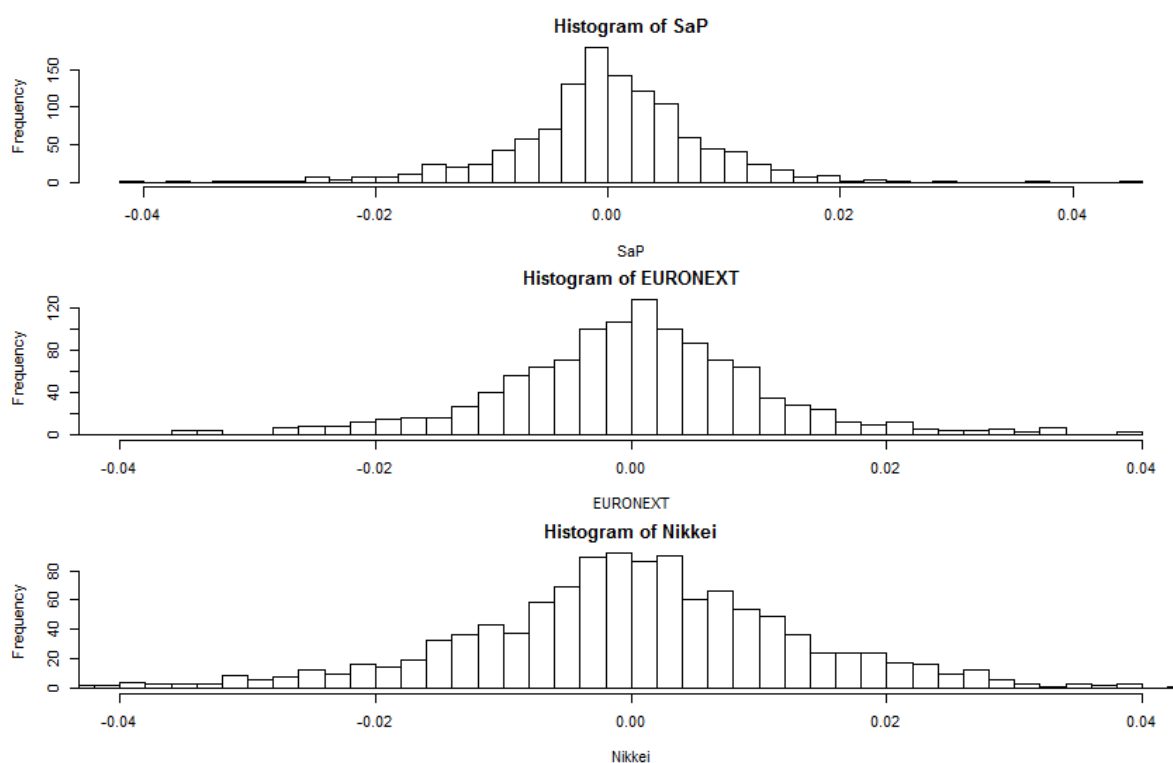
Modelujeme tedy vývoje cen indexů jako exponenciální řady. To v podstatě odpovídá spojitému úročení – to je velmi blízko tomu, co bychom na velmi likvidních burzách očekávali (pokud reinvestujeme zisky, což v podstatě děláme tím, že pozici držíme).

V další části tedy budeme analyzovat hodnoty těchto posunutých řad. Rádi bychom viděli přibližně normální rozdělení, ovšem s o něco těžšími chvosty. Ty by měly být způsobené hromaděním volatilit, neboť vyšší volatilita způsobuje vyšší pravděpodobnost zvýšených inovací. V následujících grafech (obrázek 4-4) vidíme časové řady těchto standardizovaných logaritmických změn:



Obrázek 4-4: Chování posunutých logaritmických výnosů

Grafy potvrzují převážně informace, které jsme viděli už z grafu surových dat. Velké změny volatilit během pádů na burzách, stejně tak i jejich hromadění. Zároveň vidíme časově blízká období zvýšené volatility mezi americkým a evropským trhem. Na japonském trhu se často vyskytuje zvýšená volatilita, která se zdánlivě neprojeví na ostatních trzích. To ukazuje na nižší provázanost těchto trhů a stejně tak na to, že západní trhy jsou vcelku zabezpečeny proti pádům na asijských trzích. Pro nás bude důležitější rozbor histogramů těchto odchylek (obrázek 4-5), neboť z nich bude lépe vidět jejich přibližné rozdělení. Ty ještě doprovodíme klasickými testy momentů.



Obrázek 4-5: Histogramy posunutých logaritmických výnosů

Z histogramů na obrázku 4-5 vidíme, že rozdělení budou velmi leptokurtická. Především u japonského trhu je graf celkem plochý, chvosty jsou o hodně těžší než u klasického normálního rozdělení. To by měla potvrdit hodnota koeficientu špičatosti, která je v případě leptokurtických rozdělení vyšší než u normálního rozdělení, tedy více než 0. Další důležitý parametr pro zkoumání je šikmost rozdělení. Ačkoliv v této vlastnosti není u metody GARCH takový rozdíl od normálního rozdělení, většina studií ukazuje, že šikmost pro finanční data bývá spíše záporná. To znamená, že těžší chvosty se v rozdělení vyskytují spíše v záporném směru. Pro akciové indexy to dává smysl,

neboť běžně dochází převážně k pomalému růstu, ale rychlým pádům. Vlastnosti šikmosti a špičatosti na finančních trzích je zkoumána například v (Jondeau & Rockinger, 2000).

Tabulka 4-4 ukazuje výsledky při výpočtech výběrových koeficientů šikmosti a špičatosti, tedy třetí a čtvrtý centrální moment. Upozorňujeme na fakt, že špičatost je uvedena oproti hodnotě normálního rozdělení, tedy už po odečtení 3 (*excess kurtosis*).

	Výběrová šikmost	Výběrová špičatost
S&P 500	-0,2359	2,7947
EURONEXT 100	-0,3065	2,9070
Nikkei 225	-0,2825	3,6322

Tabulka 4-4: Odhad šikmosti a špičatosti jednotlivých řad

Vysoké hodnoty potvrzují to, co jsme viděli na grafech. Špičatost dosahuje vysokých hodnot, především pro japonský trh. To poukazuje na výskyt hromadění volatility a znamená to správnost využití metody GARCH. Výběrová šikmost je nenulová. Často se z toho důvodu při aplikaci metody GARCH využívají jiná rozdělení pro inovace než klasické normální. Jde především o tzv. sešikmené Studentovo t-rozdělení (klasické Studentovo rozdělení je symetrické a má tedy nulovou šikmost, pokud je definována), které je právě za účelem aplikace na sešikmené časové řady upraveno. Výhodou Studentova t-rozdělení je také to, že má těžší chvosty než normální (Kercheval & Liu, 2011). Tato vlastnost je navíc v podstatě nastavitelná pomocí jeho jediného parametru, a to stupňů volnosti.

Stupně volnosti u Studentova t-rozdělení mohou ovlivnit jak hodnoty rozptylu a špičatosti, tak jejich existenci. Například rozdělení $t(1)$, které odpovídá tzv. Cauchyho rozdělení s parametry $(0;1)$, vůbec nemá definovanou střední hodnotu ani rozptyl. Má totiž natolik těžké chvosty, že tato střední hodnota nekonverguje. Přesněji platí, že Studentovo t-rozdělení má momenty definované až do počtu stupňů volnosti-1. Tedy $t(2)$ má definovanou střední hodnotu, $t(3)$ i rozptyl, atd. Pokud tedy chceme uvažovat konečnou špičatost, potřebujeme rozdělení s minimálně 5 stupni volnosti.

4.2. Odhad modelu pomocí EWMA

Jak již bylo řečeno, tento model je výpočetně jednodušší než GARCH, a proto jej vyhodnotíme jako první. K jeho výpočtu nám bude stačit program Microsoft Excel (použita byla verze 2016). Nejdříve budeme analyzovat možnosti při výběru koeficientu

lambda. Jak bylo uvedeno v teoretické části, jeho hodnota určuje rychlost rozpouštění velkých odchylek a také míru reakce na šoky. Podle RiskMetrics je doporučená hodnota 0,94 (RiskMetrics Technical Document, 1996); my se pokusíme optimalizační metodou ověřit, zda by nebylo lepší využít jinou hodnotu.

Pokud by námi nalezená hodnota byla velmi odlišná od doporučované, mohli bychom zpřesnit výsledky použitím této hodnoty. V případě nepatrného rozdílu od tohoto ovšem upustíme, protože takovou změnou bychom změnili přesnost pouze nepatrně za cenu jednoho stupně volnosti, což může mít za následek slabší schopnosti predikce na skutečných datech (i když při stavbě bude mít model lepší vlastnosti).

Dalším důležitým parametrem, který je úzce propojen s lambdou, je takzvaný *cutoff*. Jak již bylo řečeno, není v praxi možné využívat nekonečnou historii a stejně tak není praktické používat zbytečně stará data. Ta totiž byla naměřena za jiné situace na trhu a stejně ají velmi malé váhy, jejich započtení v modelu tedy pouze vede ke zvýšení výpočetní náročnosti. *Cutoff* tedy určuje hranici toho, jak stará data budeme v modelu používat. Tato veličina je v procentech a udává maximální součet nepoužitých vah.

V následujících tabulkách si ukážeme velikost některých vah při použití různých hodnot lambda a stejně tak jejich kumulovaný součet do daného času. V tabulce 4-5 jsou vidět váhy pro lambda 0,75:

měření	1	2	5	10	20	50	100
váha	0,25	0,1875	0,079102	0,018771	0,001057	1,8877E-07	1,07E-13
kumul.váha	0,25	0,4375	0,762695	0,943686	0,996829	0,999999434	~1

Tabulka 4-5: Kumulované exponenciální váhy pro lambda=0,75

Číslo v prvním řádku udává, jaký je časový posun mezi měřením a aktuálním datem. Hodnota 10 tedy znamená měření v čase $t-10$. Je vidět, že při takové hodnotě lambda (která je velmi nízká) by již prvních pět vah mělo celkový součet přes 75 %. RiskMetrics využívá cutoff na hladině 1 %, což by zde znamenalo, že bude použito pouze nejstarších 16 měření. Když použijeme doporučenou hodnotu 0,94, budou měření využívat váhy uvedené v tabulce 4-6.

měření	1	2	5	10	20	50	100
váha	0,06	0,0564	0,046845	0,03438	0,018517	0,00289345	0,000131
kumul.váha	0,06	0,1164	0,266096	0,461385	0,709894	0,954669273	0,997945

Tabulka 4-6: Kumulované exponenciální váhy pro lambda=0,94

Jak jsme již zmínili, vyšší hodnota λ znamená menší vliv posledních měření, a tedy hladší průběh grafu klouzavého průměru. To je zřejmé i zde, neboť vliv posledního měření je nyní přibližně pouze čtvrtinový. Navíc zatímco u předchozí hodnoty bylo 95 % překročeno už u 11. hodnoty, nyní je na 95 % součtu vah nutno použít 50 dat. Cutoff pro tuto hodnotu (a tedy také cutoff podle metodologie RiskMetrics) je u měření $t-74$.

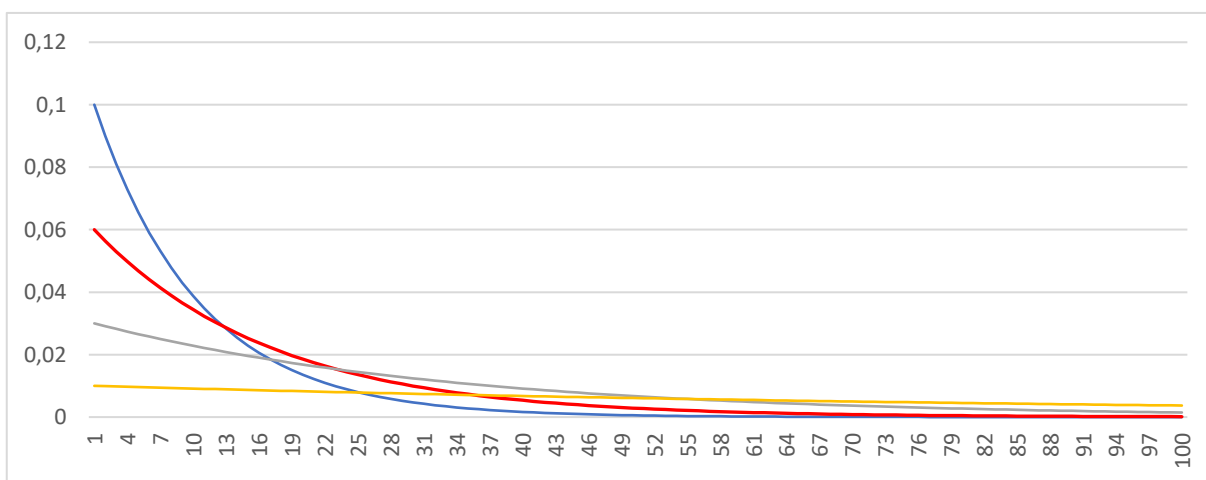
V tabulce 4-7 si ještě ukážeme stejná data u koeficientu pro investiční účely (0,97):

měření	1	2	5	10	20	50	100
váha	0,03	0,0291	0,026559	0,022807	0,016818	0,00674429	0,001471
kumul.váha	0,03	0,0591	0,141266	0,262576	0,456206	0,781934625	0,952447

Tabulka 4-7: Kumulované exponenciální váhy pro $\lambda=0,97$

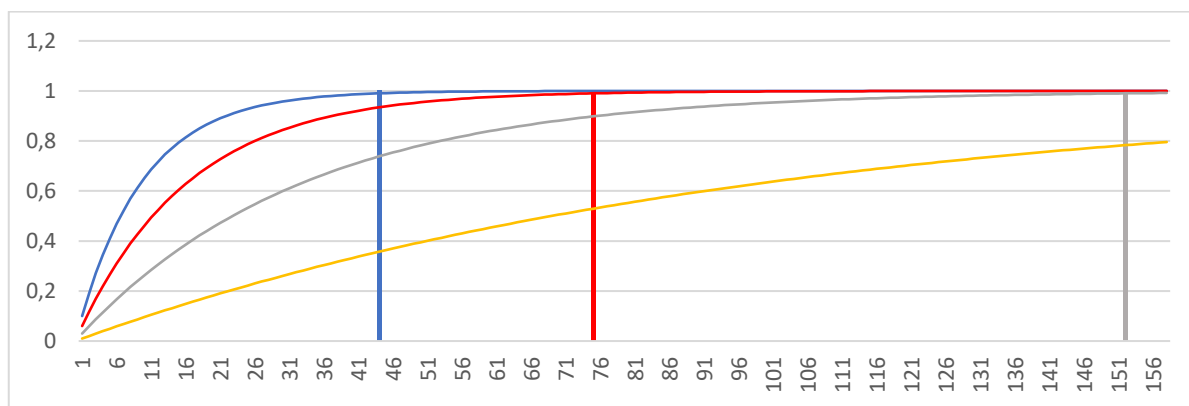
U tohoto koeficientu už by křivka proložená modelem byla velmi hladká. Vidíme, že efekt posledních 20 měření je menší než 50 %, což vede k tomu, že je volatilita z velké části modelována její nepodmíněnou střední hodnotou. Hranice 95% součtu vah je tentokrát překročena až u měření v čase $t-98$ a cutoff je dokonce 151-denní.

V grafu 4-6 si můžeme znázornit vývoj těchto vah pro koeficienty λ 0,9 (modře), 0,94 (červená), 0,97 (šedá) a 0,99 (oranžová). Na grafu je zřejmě vidět, jak rychle by klesal vliv starších měření při použití koeficientu 0,9. Na druhou stranu u hodnoty 0,99 mizí efekt velmi pomalu.



Obrázek 4-6: Graf vývoje exponenciálních vah pro různé hodnoty λ

Pomalý pokles se projeví především na vyšším *cutoff time*, což znamená že při použití vyšší hodnoty lambda musíme vzít v úvahu daleko vyšší množství dat – to způsobuje růst výpočetní náročnosti a stejně tak to může vést k situacím, kdy používáme data z dob, kdy bylo odlišné prostředí na trhu. Různé hodnoty cutoff můžeme vidět na obrázku 4-7, zobrazujícím graf vývoje kumulativních vah:



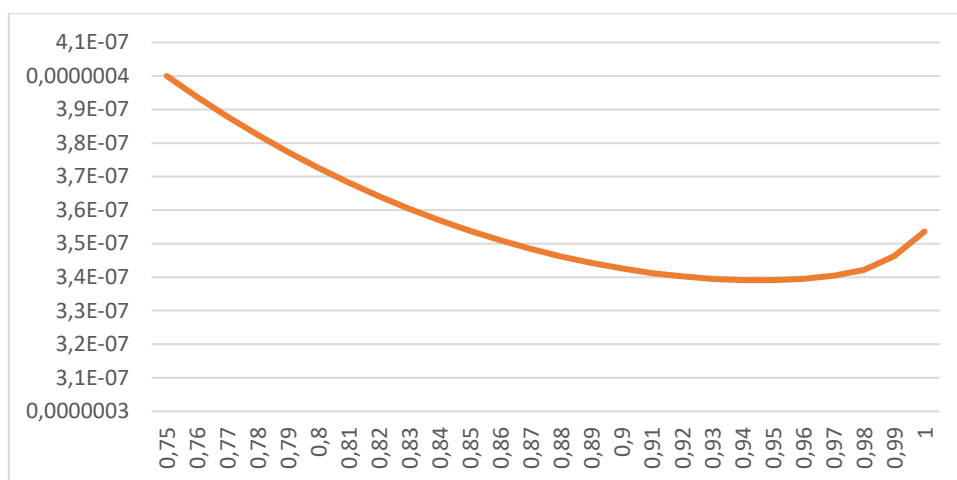
Obrázek 4-7: Graf vývoje kumulativních exponenciálních vah a jejich cut-off pro různá lambda

Svislými čarami jsou zde vyznačeny jednotlivé cutoff times. Pro oranžovou řadu (koeficient 0,99) by bylo pro 99 % vah potřeba použít dokonce 459 měření – to tedy na grafu vidět není.

Naším úkolem v další části bude otestovat hypotézu, že vhodná hodnota konstanty lambda je 0,94. Toho bude dosaženo pomocí následující optimalizační úlohy:

$$\min \sum (\varepsilon_t^2 - \hat{\varepsilon}_{t,\lambda}^2)^2,$$

kde se snažíme odhady volatility přiblížit skutečné inovaci (rozdílu mezi dosaženou hodnotou a dlouhodobým průměrným výnosem). Minimalizujeme tedy kvadrát rozdílu mezi skutečnou a odhadovanou inovací přes parametr lambda. K tomu byla využita procedura, která vypočetla tyto hodnoty pro hodnoty lambda od 0 do 1. V následujícím grafu (obrázek 4-8) jsou ukázány hodnoty v rozpětí mezi 0,75 a 1, protože při zahrnutí nižších hodnot není zřejmý tvar křivky ani umístění minima (hodnota účelové funkce je totiž řádově vyšší).



Obrázek 4-8: Graf čtvrcové chyby inovací v závislosti na koeficientu lambda

Vidíme, že optimum se podle očekávání pohybuje někde kolem hodnoty 0,94. Přesněji je optimum v hodnotě 0,945 – rozdíl je tedy minimální a můžeme používat doporučovanou hodnotu 0,94. Další možností by ještě bylo pro každou z řad používat vlastní koeficient lambda. To by mohlo být užitečné v případě, že by řady měly rozdílné reakce na náhlé šoky. Pokud by například japonský index vykazoval daleko silnější vlastnost hromadění volatility, jeho lambda by mohla být o mnoho blíže jedné než u klidnějších amerických aktiv.

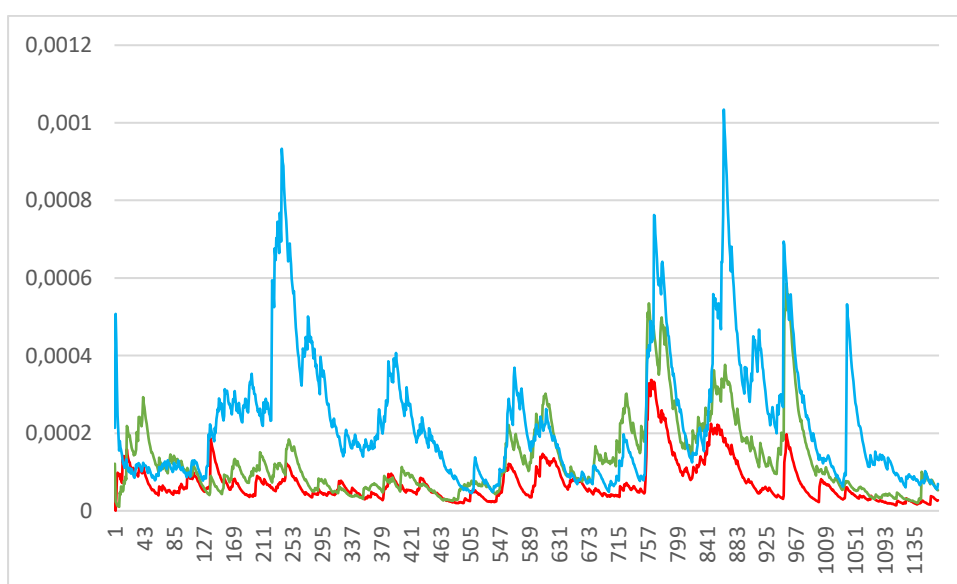
Použili jsme proto optimalizaci i na každou řadu zvlášť. Z výpočtu vyšly hodnoty účelové funkce pro jednotlivé koeficienty lambda následovně:

λ	S&P	EURONEXT	NIKKEI
0,9	2,4141E-05	8,08426E-05	0,0002959
0,91	2,4107E-05	8,06158E-05	0,0002945
0,92	2,4095E-05	8,04597E-05	0,0002935
0,93	2,4106E-05	8,0376E-05	0,0002927
0,94	2,4140E-05	8,03674E-05	0,0002923
0,95	2,4197E-05	8,04375E-05	0,0002921
0,96	2,4277E-05	8,05935E-05	0,0002924
0,97	2,4383E-05	8,08557E-05	0,0002930
0,98	2,4532E-05	8,13109E-05	0,0002945
0,99	2,4833E-05	8,24759E-05	0,0002979
1	2,5407E-05	8,55632E-05	0,0003027

Tabulka 4-8: Čtvrcová chyba odhadu v jednotlivých řadách pro různě zvolená lambda

Optimální hodnoty jsou tedy pro řady opravdu různé. Pouze u evropského indexu je optimální hodnota opravdu v 0,94. Pro S&P a Nikkei jsou hodnoty postupně 0,92 a 0,95. Další otázkou je, jestli se vyplatí využívat tři různé koeficienty λ . Výhodou tohoto postupu by byla vyšší přesnost modelu na trénovacích datech (tedy na datech, ze kterých odhadujeme model). Cenou za tuto přesnost je výpočetní náročnost, vyšší složitost výpočtu korelačních koeficientů a odebrání tří stupňů volnosti z modelu. Přesnost na trénovacích datech by tedy mohla být pouze zdánlivá a mohla by naopak vést ke snížení schopnosti predikce modelu, pokud by ve skutečnosti byly hodnoty λ proměnlivější nebo by hodnoty byly odlišné od těch, které jsme naměřili. Budeme proto pro všechny řady používat doporučenou hodnotu 0,94.

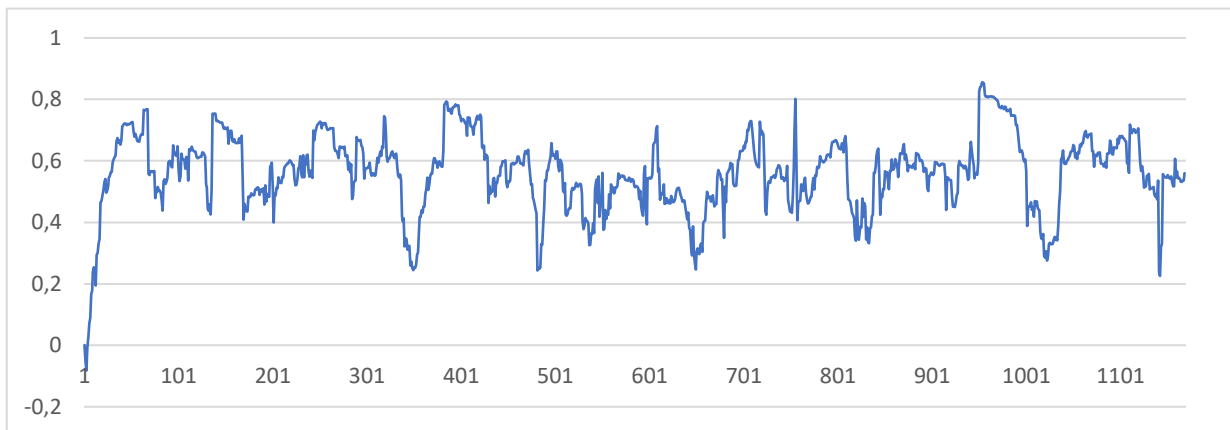
Při odhadu VaR tedy budeme nejdříve potřebovat zjistit hodnoty volatilit jednotlivých řad. Tento postup je nejjednodušší a nevyžaduje odhad žádných parametrů kromě zmíněné λ . Vývoj odhadů volatilit jednotlivých aktiv jsou ukázány na obrázku 4-9:



Obrázek 4-9: Odhady volatilit S&P (červeně), EURONEXT (zeleně) a Nikkei (modře)

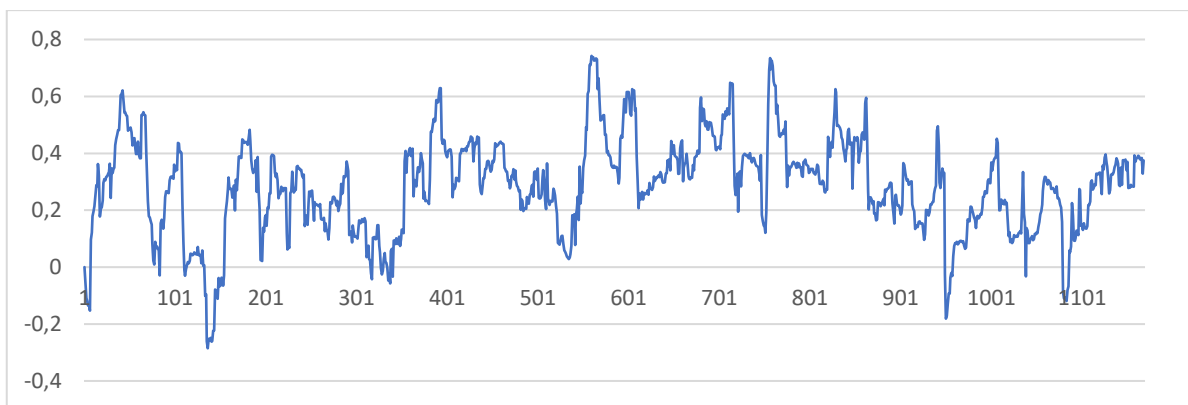
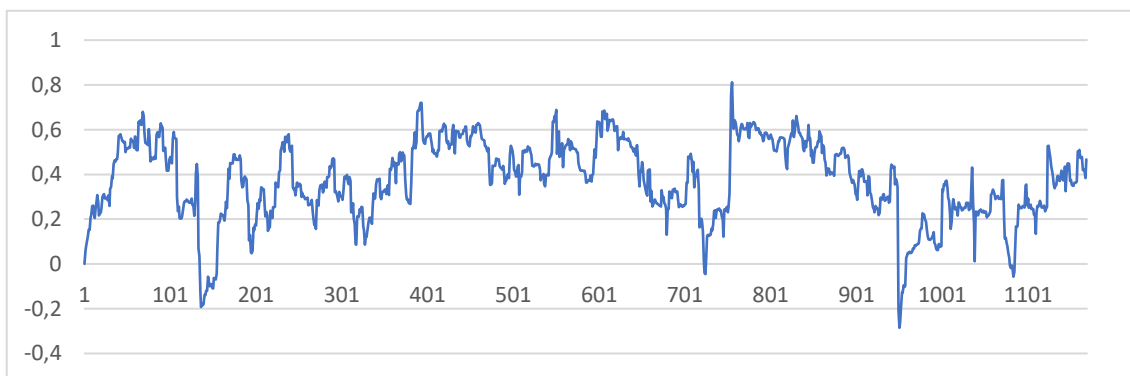
Z grafu 4-9 je zřejmých několik závěrů, které jsme již vyvodili dříve. Americký index je nejstabilnější, jeho volatilita je až na několik případů nejnižší ze všech aktiv. Na druhé straně je index NIKKEI, který je překonán volatilitou evropského indexu jen výjimečně. V modelu EWMA je navíc nápadně zachycena vlastnost hromadění volatility, když ve všech případech po velkém šoku dojde ke zvýšení volatility a ta potom postupně vyprchává. Zřejmá je i volatilita v těchto bodech.

Pro názornost si ještě ukážeme grafy jednotlivých korelací, které byly vypočteny pomocí metody popsané v kapitole teoretického základu EWMA:



Obrázek 4-10: Odhad vývoje korelace výnosů indexu S&P a EURONEXT pomocí metody EWMA

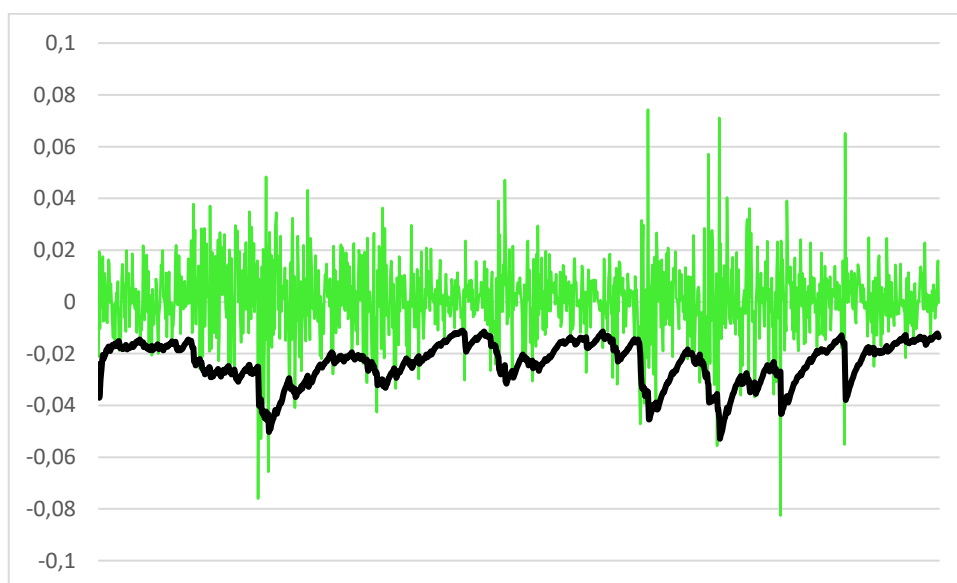
Na grafu korelace amerického a evropského indexu je vidět volatilní vývoj odhadu této volatility. Hodnota osciluje přibližně kolem 0,6 – to odpovídá původně odhadnuté korelaci těchto dvou aktiv. Zbylé dva grafy v podstatě vykazují podobné výsledky, pro ilustraci je ale ukážeme i tak:



Obrázek 4-11: Odhad vývoje korelace výnosů indexu S&P a Nikkei (nahore); EURONEXT a Nikkei (dole) pomocí metody EWMA

Vidíme, že i tyto korelace se pohybují kolem původně odhadnutých hodnot. Jsou ovšem volatilnější a v některých případech (i když výjimečně) dokonce poklesnou pod 0.

Po vygenerování řad odhadnutých volatilit a korelací již můžeme přistoupit k odhadu Value at Risk s využitím výsledků metody EWMA. Díky tomu, že jsme transformovali řadu logaritmických změn tak, aby měla střední hodnotu v 0, budeme hledaný kvantil modelovat tak, že 5% kvantil normálního rozdělení vynásobíme odmocninou z modelovaného rozptylu. Například pro japonský index tak dostaneme následující graf hodnot VaR:



Obrázek 4-12: Odhad vývoje Value at Risk pro logaritmický výnos indexu Nikkei

Na grafu 4-12 je vidět, že princip EWMA vytváří jakýsi dolní obal, jehož hloubka je přizpůsobena v případě šoků. Snaží se tak zachytit oblasti zvýšené volatility. Vzhledem k definici Value at Risk by k porušení této hranice (černě) mělo dojít v 5 % měření, v našich datech by to odpovídalo přibližně 58 případům. K tomu ve skutečnosti dojde při výpočtu na datech S&P. V případě EURONEXT bude takových překonání celkem 60 a u japonského indexu dokonce 61. V další části práce se pokusíme Value at Risk modelovat pomocí složitějších postupů v zájmu přesnějších výsledků.

4.3. Odhad GARCH(1,1) pomocí ML

V první části empirické studie budeme odhadovat parametry v metodě GARCH pomocí metody maximální věrohodnosti. Jak již bylo řečeno v teoretické části, budeme využívat dvoufázové metody. V první fázi se odhadnou parametry jednotlivých volatilit, ve druhé bude odhadována matice dynamických korelací. Tohoto výpočtu bude dosaženo pomocí matematického programu R a jeho package *ccgarch* (Nakatani, 2014). Ten vrátí výsledky pro následující formy rovnic, které odpovídají definicím GARCH(1,1), uvedeným výše:

$$y_t = H_t^{1/2} \epsilon_t,$$

$$H_t: \quad h_{ii,t} = \omega_i + \alpha_i * y_{i,t-1}^2 + \beta_i h_{ii,t-1},$$

$$h_{ij,t} = \frac{q_{ij,t} \sqrt{h_{ii,t} h_{jj,t}}}{\sqrt{q_{ii,t} q_{jj,t}}}$$

$$Q_t = (1 - a - b) \mathbf{R} + a u_{t-1} u'_{t-1} + b Q_{t-1},$$

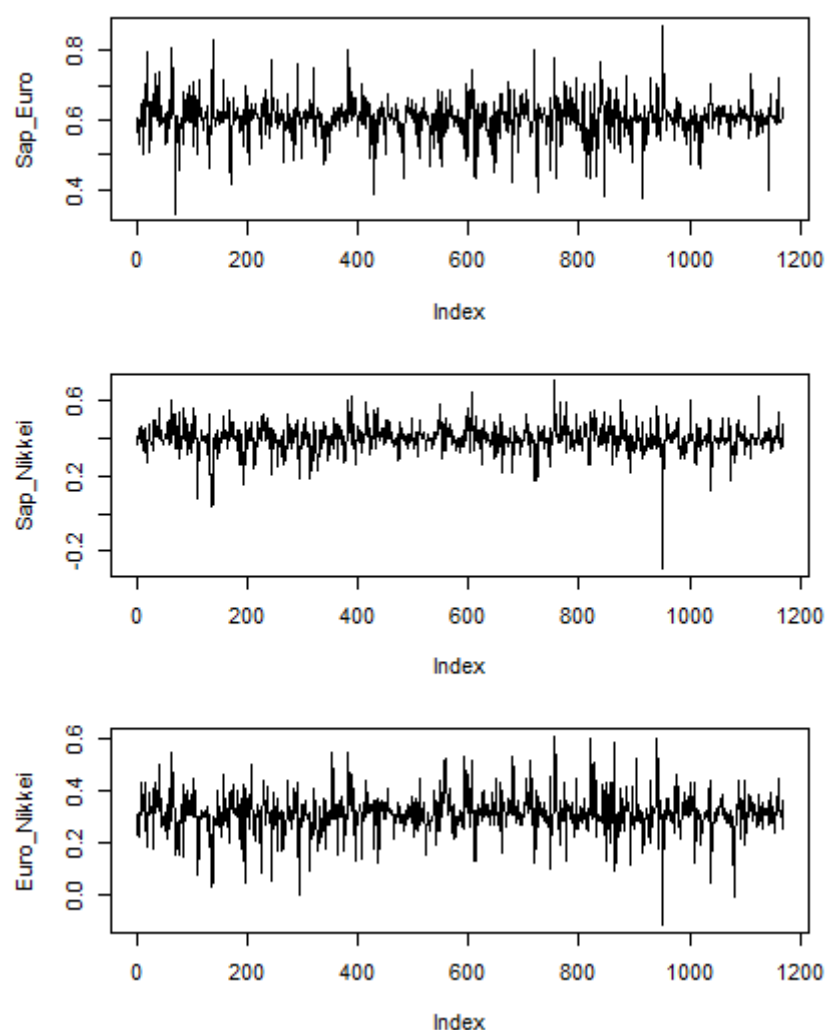
kde ϵ_t je standardizovaný bílý šum, u_t jsou inovace a $\mathbf{R}$ je matice nepodmíněných kovariancí u_t . Pokud nyní spustíme výpočet pro naměřená data (posunutá logaritmické výnosy), dostaneme maximálně věrohodné odhady, které jsou v tabulce 4-9:

	odhad	Směrodatná odchylka
ω_1	1,002*10 ⁻⁵	3,47*10 ⁻⁶
ω_2	8,29*10 ⁻⁶	5,286*10 ⁻²
ω_3	9,74*10 ⁻⁶	8,718*10 ⁻²
α_1	0,20	2,6*10 ⁻⁶
α_2	0,15	0,03027
α_3	0,13	0,04048
β_1	0,650	3,697*10 ⁻⁶
β_2	0,780	0,0279
β_3	0,830	0,0362
a	0,0825	0,0259
b	0,3095	0,2500

Tabulka 4-9: Výsledné parametry ML odhadu DCC GARCH

U parametrů ω , které se v rovnicích objevují jako konstantní část volatility, je zajímavá především vysoká směrodatná odchylka u indexů pařížské a japonské burzy. Vzhledem k tomu, že tato chyba je o čtyři řády vyšší než samotný odhad, můžeme tyto dva koeficienty považovat za nulové. Je zajímavé, že pro první řadu (S&P 500) vyšly všechny koeficienty s velmi nízkou směrodatnou odchylkou, zatímco u ostatních řad je odchylka nezanedbatelná.

Z koeficientů alfa a beta vidíme, že odhad další volatility bude složen z větší části z prvku předchozí volatility (AR část), ale i předchozí inovace se v rovnici projeví relativně vysokou měrou. Ve druhé fázi (rovnice DCC) byly vypočteny koeficienty a a b . Znovu je u nich vidět relativně vysoká směrodatná odchylka, ale zdá se, že využití metody GARCH je oprávněné. Dále si ukážeme pro tuto práci hlavní výstup této procedury, tedy grafy vývoje jednotlivých korelací v čase.



Obrázek 4-13: Vývoj korelací, shora dolů: S&P/EURONEXT, S&P/Nikkei, EURONEXT/Nikkei

Z vymodelovaných korelací není příliš zřejmé žádné ARMA chování. Ačkoliv dochází k relativně velkým výkyvům korelací, nejsou patrná žádná období zvýšené korelace, které bychom očekávali například během pádu na čínských burzách. Připomínáme výběrové koeficienty korelace, odhadnuté pro celý soubor jsou (v grafech shora dolů):

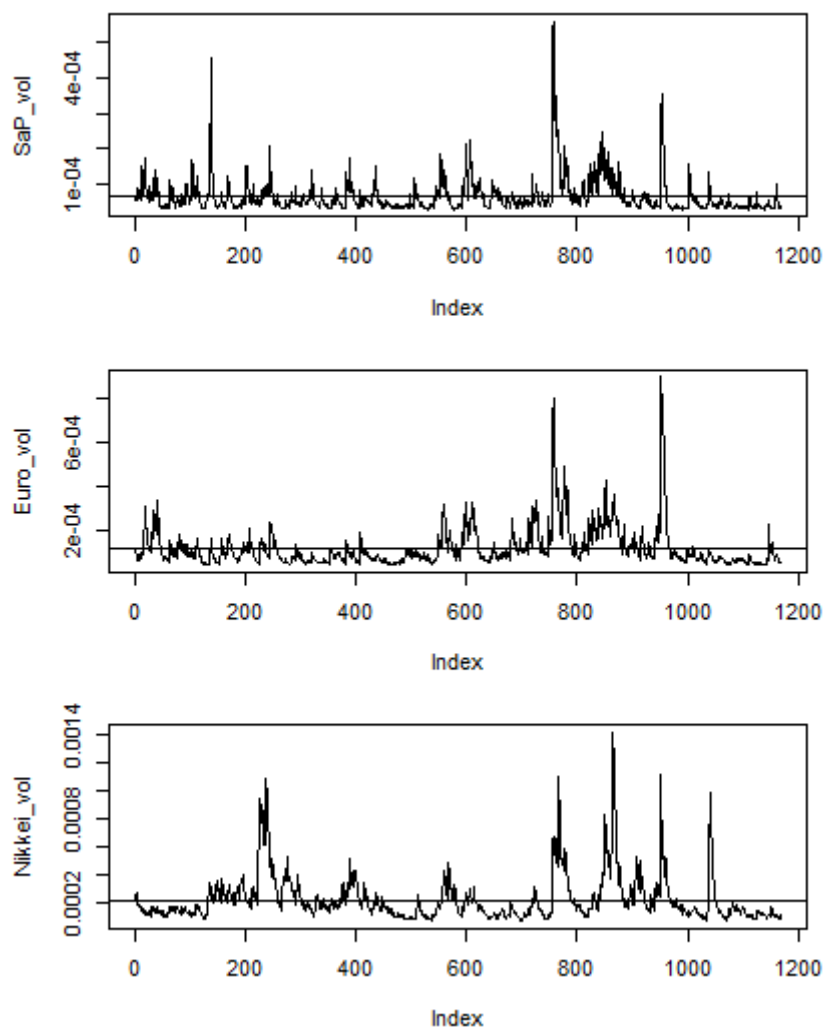
$$\text{Corr}(\text{S\&P}, \text{EURONEXT}) = 0,6043;$$

$$\text{Corr}(\text{S\&P}, \text{Nikkei}) = 0,4073;$$

$$\text{Corr}(\text{EURONEXT}, \text{Nikkei}) = 0,3034.$$

Kolem těchto hodnot se podle očekávání koeficient opravdu pohybuje, s občasnými odchylkami. Ve druhém a třetím grafu vidíme dokonce jeden skok až do záporných hodnot. Jde zřejmě o událost, která silně ovlivnila japonský trh, ale neměla vliv na západní trhy. Mohlo ale naopak jít o událost na japonském trhu, která měla na západní trhy (vzhledem k našemu posunu o den) vliv až následující den, takže zdánlivě mohlo vést k opačnému pohybu. Jde ovšem o ojedinělý úkaz.

Výstupem procedury je také hodnota volatility jednotlivých indexů. Na následujících grafech si můžeme všimnout již zmíněných období zvýšené volatility. Vzhledem k relativně nízké hodnotě odhadnutého b se ovšem vymodelovaná hodnota volatility snižuje celkem rychle. Zřejmě proto se tyto šoky neprojeví na grafech korelací. Je na nich totiž příliš mnoho dat a tyto skoky nejsou příliš patrné.



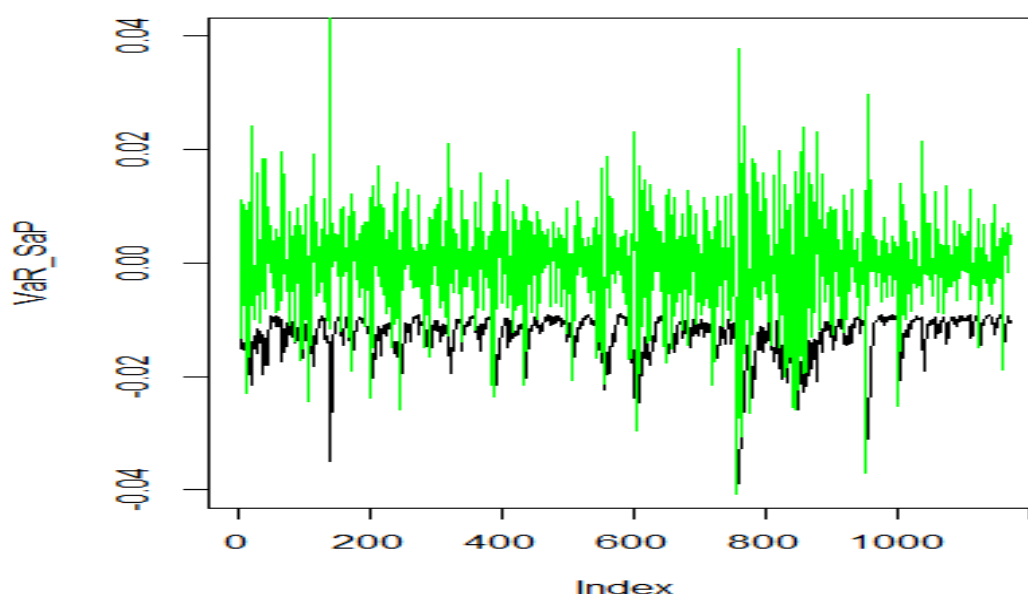
Obrázek 4-14: Vývoj odhadu volatility (shora: S&P, EURONEXT, Nikkei); horizontální čára značí celkový odhad

Grafy na obrázku 4-14 ukazují zejména zvýšená období volatility a zpomalené vyrovnávání a návrat ke stálé hodnotě volatility. Horizontální čára ukazuje původní odhad konstantního rozptylu. Vidíme, že zde jsou dočasné odchylky o mnoho vyšší, to by mělo zpřesnit odhad Value at Risk. Navíc je vidět, že tyto odchylky jsou nejčastější a také trvají nejdéle u japonských aktiv. Dalo by se tedy říct, že index Nikkei má nejvyšší volatilitu volatility.

Nyní už můžeme přejít k odhadu Value at Risk. Abychom dobře využili použitá data, budeme hledat hodnotu $\text{VaR}(0,05)$. Pokud je model dobře připraven, mělo by k překročení této hodnoty docházet přibližně v 5 % případů.

4.4. Odhad Value at Risk pomocí ML

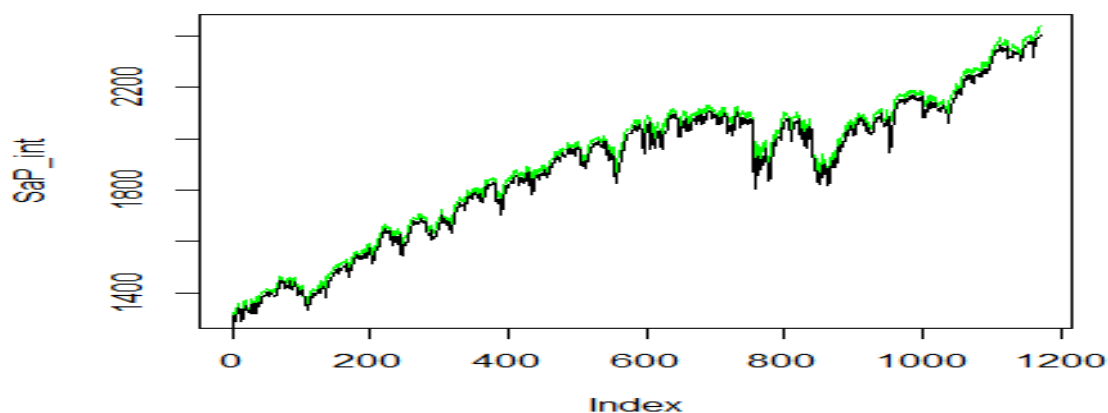
Jak již bylo uvedeno, odpovídá Value at Risk maximální ztrátě, ke které může dojít na 5% hladině významnosti (tedy 95% hladině spolehlivosti). Pro nás to bude znamenat minimální logaritmický výnos, ke kterému může s touto pravděpodobností dojít. Protože predikci modelujeme konstantní nulou, půjde jednoduše o kvantil normálního rozdělení s nulovou střední hodnotou a rozptylem, jehož řadu jsme vymodelovali. Hodnota VaR v každém bodě tedy bude vypočtena jako 5%-kvantil standardního normálního rozdělení, vynásobený odmocninou z hodnoty odhadnuté volatility v daném čase. Na obrázku 4-15 si tedy ukážeme grafy vymodelovaných hodnot VaR a skutečně dosažených hodnot jednotlivých indexů:



Obrázek 4-15: Vývoj odhadu Value at Risk indexu S&P (černě) v porovnání se skutečnými hodnotami (zeleně)

Černá čára zde značí VaR, zatímco zelená skutečně dosažené výsledky indexu S&P 500 v daný den. Zejména je zde vidět vliv inovací v modelu GARCH, neboť vždy po porušení VaR (nebo samozřejmě při nadměrném výnosu) dojde k velkému zvýšení odhadované volatility a tím i k posunu hodnoty VaR níže. Rádi bychom ještě ověřili hypotézu, že k tomuto porušení dochází v 5 % případů. Pokud by k němu docházelo příliš často, značilo by to podcenění volatility a nevhodnost modelu z důvodu podstupování vyššího rizika, než je naměřeno. Opačná chyba (příliš málo porušení) by poukazovala na přílišnou opatrnost modelu, která by pro investora mohla znamenat přehnanou opatrnost a tím by vedla ke zbytečně nízkým ziskům v zájmu diverzifikace. Protože Value at Risk se většinou definuje spíše pro skutečnou hodnotu než výnosy,

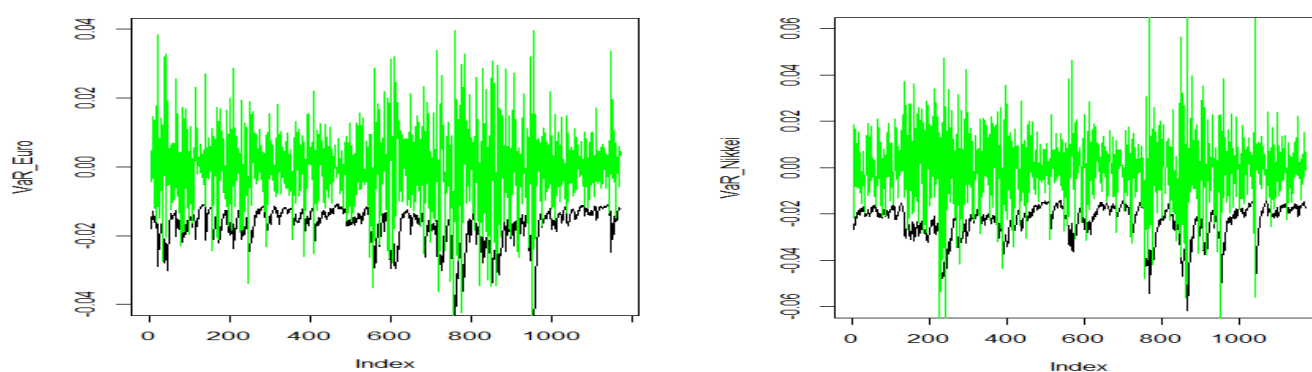
můžeme si ještě pomocí dříve uvedené exponenciální regrese ukázat, jak vypadá VaR ve skutečných hodnotách indexu S&P:



Obrázek 4-16: Value at Risk pro hodnotu indexu S&P (černě) v porovnání se skutečným vývojem (zeleně)

Z tohoto grafu vidíme, jak blízko tato „obálka“ skutečně je. Volatilita samotného indexu totiž není příliš vysoká, ani hodnota Value at Risk tedy nebude nijak extrémní. Jak vidíme z odhadu pro logaritmický výnos, maximální odhadovaná změna při dané pravděpodobnosti oproti předchozí hodnotě bude přibližně $e^{-0,04}$, tedy maximální mezidenní VaR jsou přibližně 4% aktuální hodnoty indexu.

Dále si v grafech ukážeme ještě odhady pro ostatní akciové indexy (vlevo Euronext, vpravo Nikkei) – znovu je vidět daleko vyšší volatilita, což vede i k vyšším VaR:



Obrázek 4-17: Vývoj odhadu Value at Risk pro výnosy indexu (černě) v porovnání se skutečným vývojem (zeleně): EURONEXT (vlevo) a Nikkei (vpravo)

Z důvodu vysoké volatility změn nejsou grafy na obrázku 4-17 příliš přehledné. Je však zřejmých několik období, kdy z důvodu dočasného zvýšení volatility dojde k velké změně u hodnoty VaR.

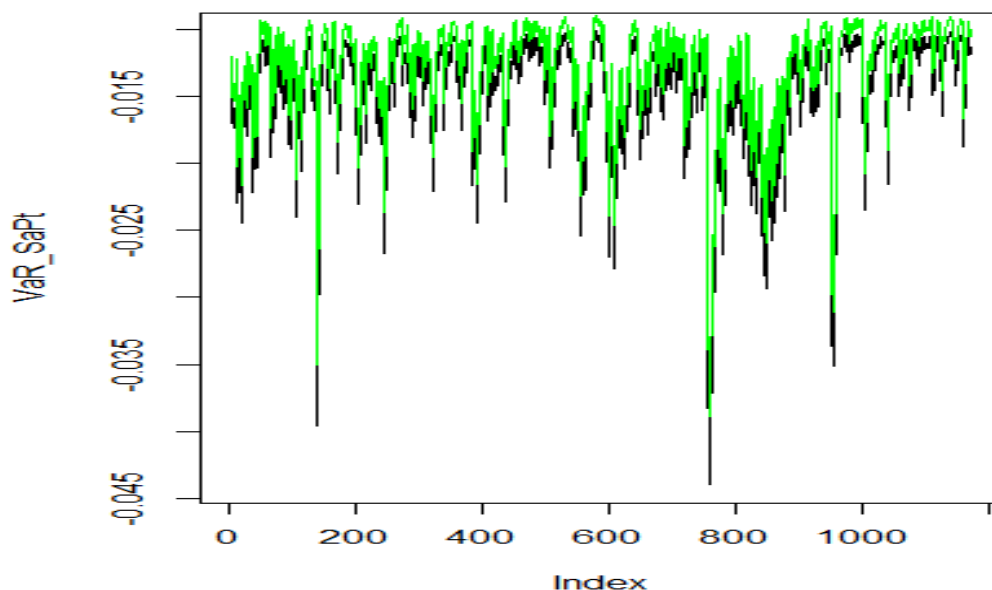
Z grafu je vidět, že VaR u indexu S&P a Euronext dosahovala v nejvolatilnějších obdobích hodnot kolem -0,04. Na japonském trhu ovšem hned několikrát došlo i k prolomení hodnoty -0,05. Ještě empiricky ověříme správnost těchto odhadů VaR podle tabulky 4-10:

	počet dat	počet porušení	EWMA	očekávaný počet
S&P 500	1169	68	58	58,40
Euronext 100	1169	61	60	58,40
Nikkei 225	1169	64	61	58,40

Tabulka 4-10: Počet porušení hodnoty Value at Risk v modelu DCC GARCH, porovnání s EWMA

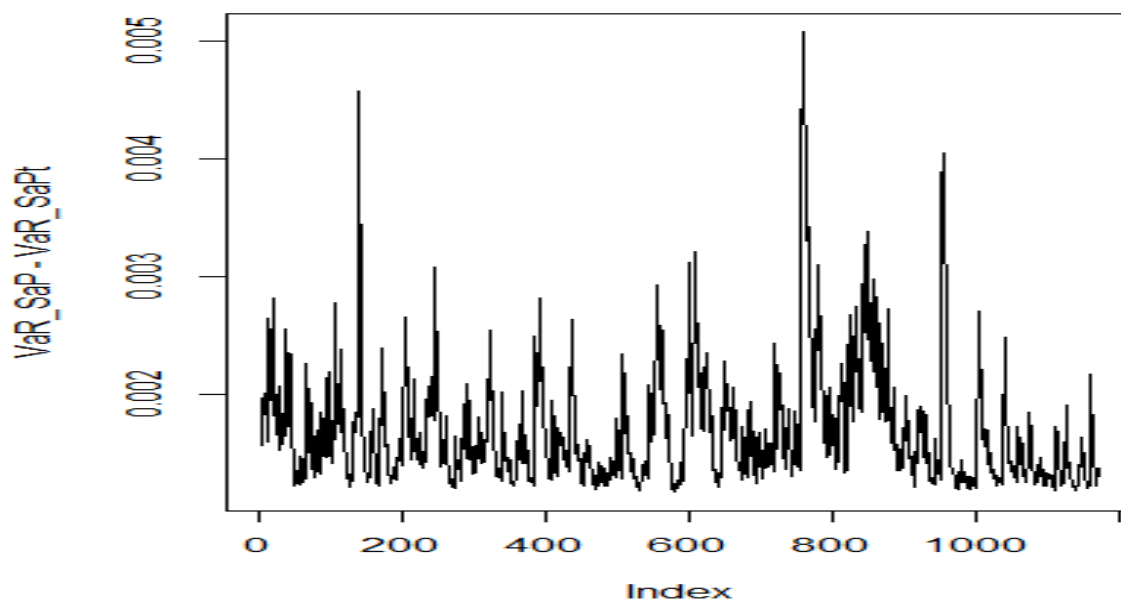
Tabulka ukazuje u všech odhadů příliš vysoký počet porušení. To by poukazovalo na to, že normální rozdělení není vhodné ani pro modelování inovací. Bylo by tedy lepší pro modelování hodnoty VaR použít nějaké rozdělení s těžšími chvosty, nejběžněji se používá Studentovo t-rozdělení. Při použití Studentova t-rozdělení ovšem vzniká další problém. Jako jeden ze vstupů totiž přibývá počet stupňů volnosti.

Protože používáme toto rozdělení z důvodu snahy zachytit vlastnost těžkých chvostů, můžeme použít takový počet stupňů volnosti, aby byla špičatost obou rozdělení téměř rovna. Špičatost t-rozdělení v závislosti na stupních volnosti je rovna $\frac{6}{\nu-4}$ (pokud má více než 4 stupně volnosti, jinak je špičatost rovna nekonečnu – chvosty už jsou tak těžké, že potřebný integrál nekonverguje). Protože nám vycházely empirické koeficienty špičatosti kolem hodnoty 3 (a počet stupňů volnosti musí být samozřejmě celé číslo), zdá se nejlepší využít rozdělení s 6 stupni volnosti, kdy t-rozdělení má teoretickou špičatost rovnou 3. Pokud budeme modelovat VaR tímto rozdělením, dostaneme například pro S&P 500 následující graf (porovnáváme s předchozím odhadem VaR pomocí normálního rozdělení):



Obrázek 4-18: Porovnání odhadu VaR s předpokladem normálního rozdělení (zeleně) a t(6)-rozdělení (černě)

V grafu 4-18 je černě vyznačen odhad VaR pomocí t-rozdělení, zeleně je původní odhad pomocí normálního rozdělení. Rozdíl není zřejmý, pouze je vidět, že t-rozdělení má opravdu těžší chvosty, což vede k opatrnějšímu odhadování VaR. Ta je ve všech bodech nižší, jak ukáže následující graf rozdílů mezi těmito odhady. Na něm jsou vyobrazeny rozdíly mezi odhadem normálním a Studentovým rozdělením. Kladná čísla tedy ukazují, že odhad t-rozdělením vždy vychází nižší, a tedy určitě nebude počet porušení vyšší:



Obrázek 4-19: Velikost rozdílů mezi odhady s využitím normálního a Studentova rozdělení

Další graf hodnot pro tuto VaR by nebyl dostatečně vypovídající, můžeme si tedy akorát v tabulce porovnat počty porušení hranice tvořené odhady VaR:

	t-rozdělení(6)	t-rozdělení(15)	normální rozdělení	očekávaný počet
S&P 500	47	58	68	58,40
EURONEXT 100	39	52	61	58,40
Nikkei 225	42	57	64	58,40

Tabulka 4-11: Počet porušení VaR pro různá rozdělení

Vidíme, že u t-rozdělení s 6 stupni volnosti jsou chvosty zřejmě až příliš těžké, neboť k porušení dochází daleko méně, než by v ideálním případě mělo. Pro zajímavost je v tabulce ještě přidáno t-rozdělení s 15 stupni volnosti. U toho dochází k počtu porušení, jaké bychom od modelu chtěli. Toto rozdělení má ale špičatost rovnu pouhých $\frac{6}{11}$ a nemáme žádný důvod předpokládat, že je vhodné takové rozdělení používat. Nebudeme proto t-rozdělení s 15 stupni volnosti využívat.

4.5. Odhad pomocí bayesovského přístupu

Jak již bylo řečeno, hlavní výhodou bayesovského přístupu je využití předem známých informací. Ty jsou uchovány v apriorních rozděleních pro jednotlivé parametry. Výsledkem potom není jen odhad střední hodnoty a směrodatné odchylky jako u klasického frekventistického přístupu, ale díky simulaci parametrů pomocí markovských řetězců je vygenerováno celé pravděpodobnostní rozdělení těchto parametrů.

Při dostatečném množství simulací potom za předpokladu správného použití apriorních rozdělení získáváme maximálně věrohodné rozdělení všech parametrů. Hlavní nevýhodou je výpočetní náročnost. Ta je kvůli nutnosti simulace přes několik tisíc iterací, především z časového hlediska, vysoká. V našem modelu se navíc nachází 12 parametrů, které je nutné odhadnout. Samotný výpočet tedy i pro vcelku nízký počet iterací zabere několik minut. Nezapomínejme ale, že výpočet pomocí maximální věrohodnosti je také velmi náročný a v zájmu využitelnosti musí využívat numerické postupy optimalizace.

K výpočtu parametrů pomocí bayesovského postupu bylo znovu využito programu R, tentokrát byl využit balíček `bayesDccGarch` (Fiorucci, Ehlers, & Louzada, 2016). Ten provádí právě generování řetězců pro jednotlivé parametry za využití zmíněných metod Markov Chain Monte Carlo. Přesněji jde o jednoblokový Metropolis-Hastingsův postup s náhodnou procházkou. Problémem metody je stanovení vhodného množství iterací. Jde o vyvažování mezi přesností rozdělení a časovou náročností procedury. My jsme v proceduře použili 10 000 iterací, při zvýšení jejich počtu už se výsledek příliš neměnil. Budeme tedy považovat výsledky za dostatečně přesné a předpokládat, že řetězec dokonvergoval ke stacionárnímu rozdělení.

Výsledky této procedury jsou následující:

	odhad	Směrodatná odchylka
ω_1	$9,544 \cdot 10^{-6}$	$2,688 \cdot 10^{-6}$
ω_2	$1,089 \cdot 10^{-5}$	$4,510 \cdot 10^{-5}$
ω_3	$1,730 \cdot 10^{-5}$	$1,014 \cdot 10^{-4}$
α_1	0,2013	0,02613
α_2	0,1409	0,02974
α_3	0,1228	0,02702
β_1	0,6763	0,03967
β_2	0,8030	0,04270
β_3	0,8395	0,03287
a	0,07979	0,01821
b	0,25540	0,1801

Tabulka 4-12: Výsledné odhady parametrů pro DCC GARCH pomocí bayesovského přístupu

Ještě ukážeme srovnání výsledků metody maximální věrohodnosti a bayesovské. I přes malé rozdíly ve většině proměnných se především u proměnných ω o několik řádů snížila směrodatná odchylka. To může ukazovat vyšší přesnost při využívání metody markovských řetězců. V tabulce 4-13 si ukážeme pouze porovnání středních hodnot odhadů:

	bayesovský odhad	ML odhad
ω_1	$9,544 \cdot 10^{-6}$	$1,002 \cdot 10^{-5}$
ω_2	$1,089 \cdot 10^{-5}$	$8,29 \cdot 10^{-6}$
ω_3	$1,730 \cdot 10^{-5}$	$9,74 \cdot 10^{-6}$
α_1	0,2013	0,20
α_2	0,1409	0,15
α_3	0,1228	0,13
β_1	0,6763	0,650
β_2	0,8030	0,780
β_3	0,8395	0,830
a	0,07979	0,0825
b	0,25540	0,3095

Tabulka 4-13: Srovnání výsledků metody DCC GARCH v závislosti na postupu výpočtu

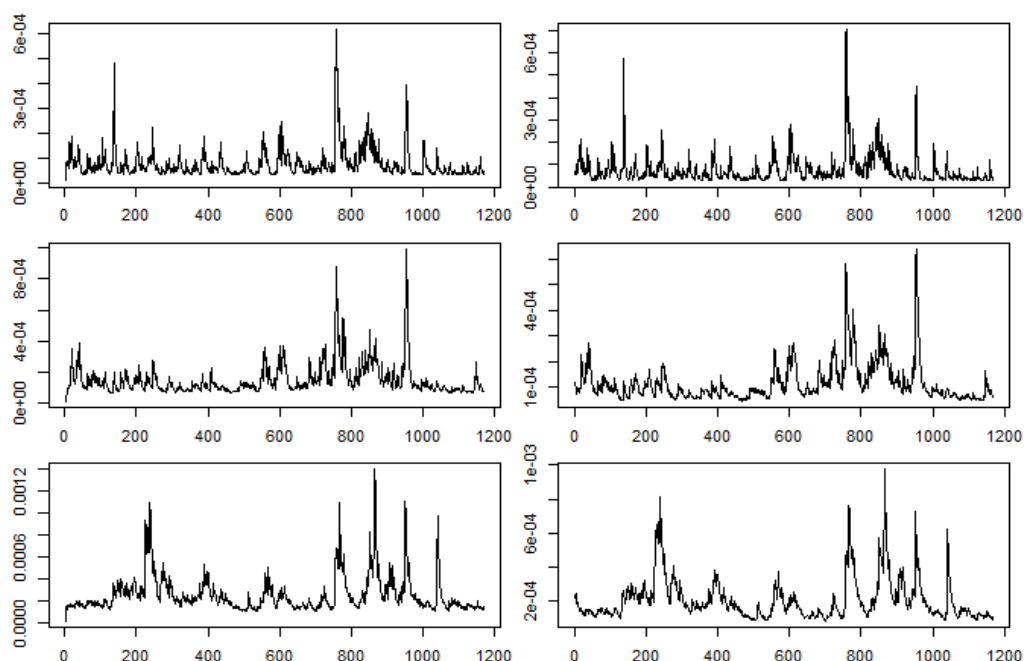
Odhady u všech parametrů vychází velmi podobně, řádové rozdíly se nenacházejí nikde. Pro zajímavost můžeme ještě srovnat výsledné hodnoty logaritmické maximální věrohodnosti obou modelů:

$$L_{ML}=11\,492,47.$$

$$L_B=11\,570,69.$$

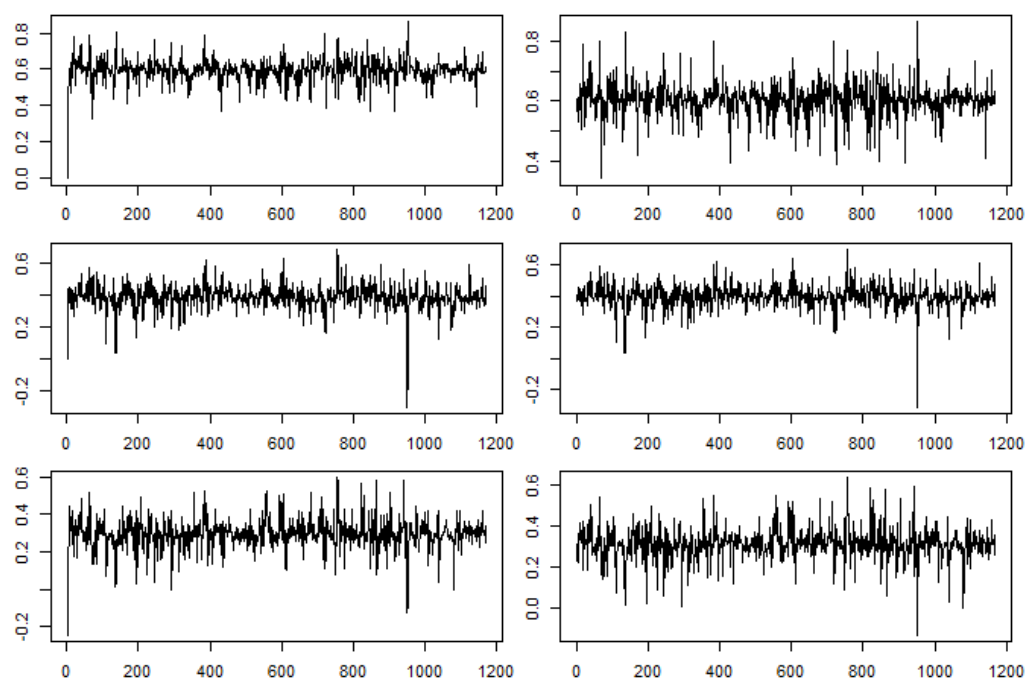
Vidíme tedy, že u modelu s odhadem pomocí maximální věrohodnosti se nepodařilo optimalizačnímu algoritmu nalézt tak věrohodné výsledky, jako se to podařilo u bayesovského algoritmu. Může to být z důvodu nedokonalosti algoritmu, nebo nesprávného výběru startovacích hodnot u algoritmu, které mohly vést k nalezení pouze lokálního maxima.

Nyní přejdeme znovu k vykreslení vymodelovaných volatilit. Výsledky se od původního odhadu téměř neliší, předvedeme si je proto v grafech na obrázku 4-20 hned vedle sebe:



Obrázek 4-20: Odhady volatilit (shora: S&P, EURONEXT, Nikkei) - vlevo bayesovský přístup, vpravo ML

Vlevo na obrázku 4-20 jsou modely, které využívají odhadů z bayesovského modelu, vpravo vidíme předchozí postup. Je vidět, že odhady jsou téměř identické. To ukazuje na fakt, že i výsledné odhady parametrů byly velmi blízké. Dalším důvodem je princip ARMA modelů. Modelujeme totiž z velké části pomocí předchozí inovace, a tedy hlavním jevem, viditelným z grafu, bude reakce na šoky. To samé potom bude platit i při modelování korelací mezi jednotlivými indexy. Znovu je můžeme porovnat s předchozími výsledky:



Obrázek 4-21: Odhady korelací (shora: S&P/EURONEXT, S&P/Nikkei, EURONEXT/Nikkei) - vlevo bayesovský přístup, vpravo ML

Zde je znovu vidět v podstatě stejné chování výsledků obou metod.

4.6. Portfoliový přístup

Hlavní silou využití metody DCC GARCH by měla být možnost modelování vzájemných korelací a jejich využití při odhadu budoucího rozdělení zisků aktiv. Zatím jsme však model aplikovali na investory, kteří drží celou svoji hodnotu v jediném aktivu. Tím však ztrácíme možnost modelování VaR pomocí odhadovaných korelací a mohli bychom téměř bez oslabení postupu využívat klasický GARCH (jediným využitím korelací v předchozím postupu bylo, že jsme mohli odhadovat pohyb vlastních aktiv podle pohybů na ostatních trzích – nevyužili jsme tím ale možnost diverzifikace). V další části práce proto metodu aplikujeme na portfolio, které bude složeno ze všech námi uvažovaných aktiv. V reálném případě bude takový případ daleko užitečnější. V praxi by bylo portfolio tvořeno větším množstvím aktiv (tato portfolio by navíc mohla být tvořena i z jiných aktiv, než jsou tržní indexy), v této práci nám pro názornost příkladu bude stačit portfolio tvořené aktivy které jsme si už uvedli výše.

Protože korelace mezi našimi aktivy není téměř nikdy záporná (k tomu došlo pouze v jednom měření), nebude možné provádět imunizaci portfolio pomocí diverzifikace. V reálném příkladu bychom mohli zařadit do portfolio aktiva, která by byla záporně korelována. Mohlo by jít například o krátký prodej aktiv s daným indexem jako podkladovým aktivem. Různým rozdělováním portfolio můžeme pouze regulovat podstoupené riziko za cenu očekávaného výnosu. Bylo by možné například v každém bodě provést optimalizační úlohu, kde bychom se snažili maximalizovat očekávaný zisk za podmínky nepřekročení nějaké hodnoty VaR. Tím bychom získali alternativu ke klasickému Markowitzovu přístupu, kdy minimalizujeme riziko za předpokladu udržení minimálního určeného výnosu.

Mějme tedy investora, jehož portfolio má hodnotu \$700 000. Z toho \$100 000 bude investováno v S&P a po \$300 000 bude vloženo v evropském EURONEXT a japonském tržním indexu. V průběhu výpočtů se zřejmě budou procentuální účasti měnit z důvodu jejich různého vývoje a průběhu zhodnocování – na počátku je portfolio procentuálně rozloženo přibližně na 14,3 %; 42,9 % a 42,9 %. Vzhledem k rychlejšímu růstu, ke kterému došlo na japonském trhu, bude na konci zastoupení jednotlivých indexů v portfoliu postupně (13,0 %; 37,5 %; 49,5 %) a hodnota portfolio vzroste z původních 700 tisíc dolarů na 1,4 milionu.

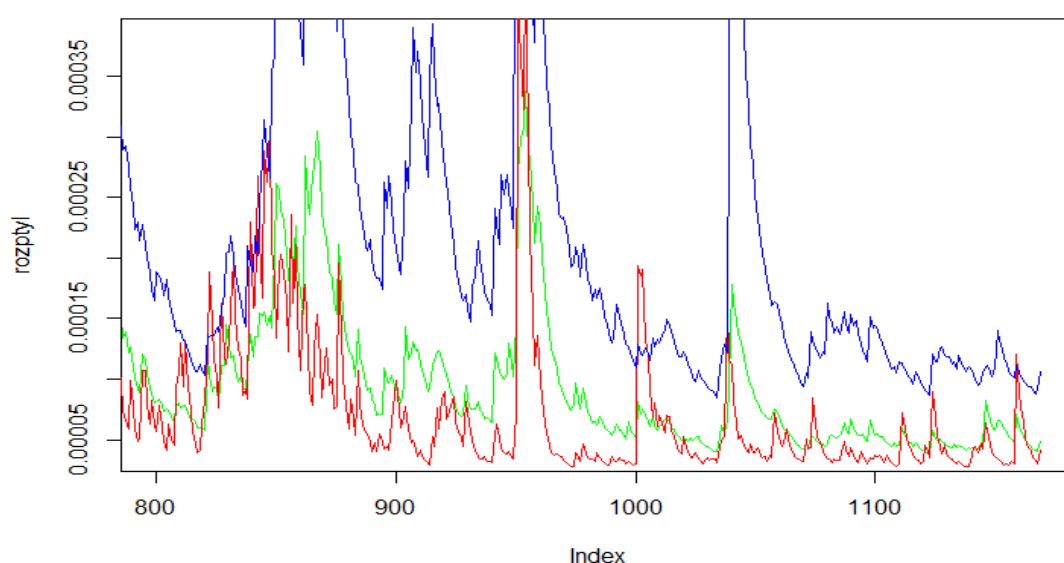
Po odhadnutí korelací a rozptylů pro aktiva můžeme rozptyl pro portfolio složené ze tří aktiv vypočítat následovně:

$$\sigma_p^2 = \omega_S^2 \sigma_S^2 + \omega_E^2 \sigma_E^2 + \omega_N^2 \sigma_N^2 + 2\omega_S \omega_E \sigma_S \sigma_E \rho_{SE} + 2\omega_S \omega_N \sigma_S \sigma_N \rho_{SN} + 2\omega_E \omega_N \sigma_E \sigma_N \rho_{EN},$$

kde ω značí podíl daného aktiva na celkové hodnotě portfolio. Vzhledem k různým pohybům jejich cen se musí tato hodnota v každém čase přepočítávat. Vzhledem k tomu, že při výpočtech používáme řady hodnot upravené o jejich střední hodnotu, budou podíly aktiv na konci rovny jejich podílům na začátku – v průběhu se však budou měnit. Tím se samozřejmě bude měnit i celková volatilita portfolio. To tedy znamená, že všechny hodnoty ve vzorci pro rozptyl portfolio je nutné počítat v každém čase zvlášť a vzorec by v podstatě měl vypadat následovně:

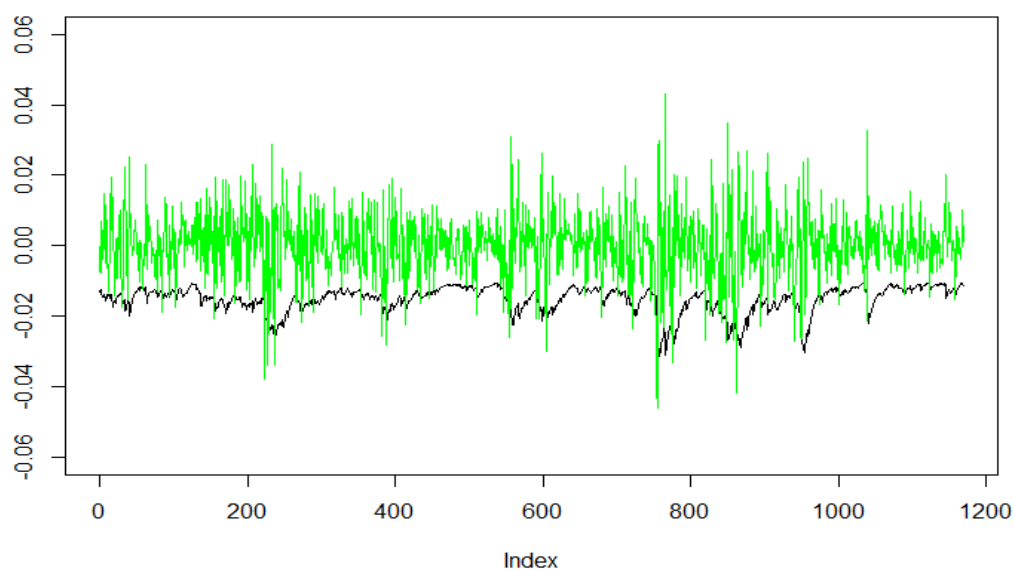
$$\sigma_{p,i}^2 = \omega_{S,i}^2 \sigma_{S,i}^2 + \omega_{E,i}^2 \sigma_{E,i}^2 + \omega_{N,i}^2 \sigma_{N,i}^2 + 2\omega_{S,i} \omega_{E,i} \sigma_{S,i} \sigma_{E,i} \rho_{SE,i} + 2\omega_{S,i} \omega_{N,i} \sigma_{S,i} \sigma_{N,i} \rho_{SN,i} + 2\omega_{E,i} \omega_{N,i} \sigma_{E,i} \sigma_{N,i} \rho_{EN,i}.$$

Po vypočtení rozptylu portfolio můžeme odhadovat VaR stejným postupem jako v předchozích kapitolách. Nejdříve si však můžeme porovnat posloupnost odhadovaného rozptylu s rozptylem jednotlivých složek. Předpokládali bychom, že rozptyl se bude pohybovat někde mezi jednotlivými rozptyly. Pro názornost si v grafu 4-22 zobrazíme pouze posledních 370 měření a ukážeme si kromě portfolio pouze volatilitu pro americký a japonský index.

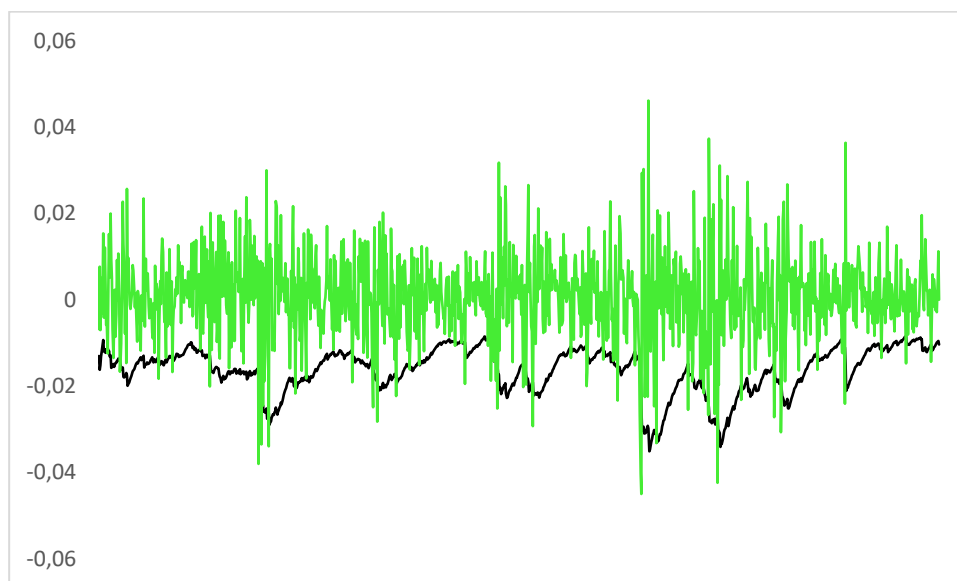


Obrázek 4-22: Odhad volatility portfolio (zeleně) v porovnání s S&P (červeně) a Nikkei (modře)

Zelená křivka na obrázku 4-22 ukazuje celkový rozptyl portfolia a podle předpokladu se většinou nachází mezi křivkami pro americký index (červeně) a Nikkei (modře). Vzhledem ke kombinaci portfolia je zde vyšší efekt *volatility clustering* než u stabilních amerických akcií. Zajímavé je, že v několika bodech se křivka rozptylu celého portfolia dostala pod volatilitu amerických akcií – došlo tedy i k částečné diverzifikaci. Po vymodelování Value at Risk můžeme znovu tuto řadu vykreslit v grafu 4-23.



Obrázek 4-23: Odhad VaR portfolia (černě) v porovnání se skutečnými výnosy (zeleně): DCC GARCH



Obrázek 4-24: Odhad VaR portfolia (černě) v porovnání se skutečnými výnosy (zeleně): EWMA

V porovnání s metodou EWMA (obrázek 4-24) je z grafu vidět, že GARCH se snaží být přesnější a tvoří menší „obálky“ pro odhad hodnoty, zatímco přizpůsobování v metodě EWMA je mnohem pomalejší a tím pádem dává opatrnější odhady. I tak jsou jejich grafy v reakci na šoky velmi podobné a oba tyto postupy se dobře vyrovnávají se zvýšenou volatilitou na trhu a obdobími vysoké korelace. V tabulce 4-14 ještě porovnáme výskyt porušení hodnoty VaR, pro ověření vhodnosti výpočtu:

GARCH (normální rozdělení)	GARCH t(6) rozdělení	EWMA	5 %
65	45	58	58,40

Tabulka 4-14: Počet porušení VaR portfolia pro jednotlivé přístupy

Z tabulky 4-14 vidíme, že nejbližší očekávanému počtu porušení hranice VaR je metoda EWMA (která je přesná). GARCH využívající normální rozdělení je na hodnotě 65, což odpovídá 5,6 % a při využití Studentova t-rozdělení se 6 stupni volnosti se dostaneme na 45 výskytů, tedy pouhých 3,9 %. Oba tyto výsledky jsou ve většině situací stále přijatelné, ale zřejmě je zde postup pomocí metody EWMA přesnější.

5. Závěr

Metoda GARCH, která je založena na zahrnutí *volatility clusteringu*, anomálie vedoucí k obdobím s vyšší volatilitou na světových trzích, se ukázala jako užitečná pro odhad Value at Risk pro námi vybraná aktiva. Díky zobecnění běžných metod, které si vynucují předpoklad konstantního rozptylu, je právě tato metoda vhodná pro statistické odhady i v obdobích, kdy dochází k náhlým změnám na trhu. Metodu GARCH jsme navíc v této práci posílili jejím dalším zobecněním. Metoda dynamických korelací je užitečná převážně při studování aktiv, které jsou závislé na podmínkách trhu. V práci jimi byly tržní indexy – právě proto, že v podstatě již z definice je jejich korelace se stavem na trhu vysoká.

Modelování korelací potom bylo zapojeno na tyto indexy ze tří světových trhů – těmi byly: americký index S&P 500, evropský index EURONEXT 100 a japonský NIKKEI 225. Ty v sobě zahrnují největší společnosti jednotlivých trhů a tím ukazují stav na trhu. Velikost těchto firem a moderní infrastruktura ve světových financích vede k tomu, že vývoj na trzích je někdy velmi úzce provázán. Tato korelace však není konstantní, ale roste při zvýšení volatility na trzích. Při uklidnění indexů a malých změnách naopak tyto korelace klesají. Při odhadu Value at Risk by potom nemuselo být vhodné počítat s konstantní korelací. Proto byl zapojen model DCC GARCH, který modeluje jak volatilitu, tak korelace pomocí autoregresních časových řad.

V práci byla předvedená metoda DCC GARCH. Její zásadní nevýhodou je výpočetní náročnost – nejde pouze o složitost postupu a pracnost implementace, ale i o časovou náročnost samotného výpočtu. I při využití na malé portfolio 3 aktiv, ukázané pro názornost v článku, byla náročnost relativně vysoká. Pokud bychom tedy chtěli konstruovat optimální portfolio na základě hodnoty VaR pomocí simulací metodou DCC GARCH, mohli bychom narazit na nemožnost dokončení výpočtu v rozumném čase.

Při odhadu parametrů metody DCC GARCH byly využity dva různé postupy. Prvním z nich byl odhad parametrů pomocí metody maximální věrohodnosti. Ten využívá maximalizace věrohodnostní funkce. Druhým postupem byl bayesovský odhad, který při odhadu parametrů využívá metody Monte Carlo pomocí markovských řetězců, tj. je provedeno několik tisíc simulací, na jejichž základě jsou odhadnuty parametry, které nejpřesněji odhadují chování časové řady. Výhodou této metody je, že při vyšším počtu dat neroste časová náročnost tolik jako u věrohodnostní metody. Časová

náročnost je však i její nevýhodou, protože při nižším množství dat (jako bylo například v naší práci) mohou být simulace zbytečně náročné. Pokud jsou navíc špatně zvoleny počáteční body simulace, nemusí metoda ve výjimečných případech konvergovat ke správnému řešení.

Oba postupy pro odhad parametrů DCC GARCH vedly k velmi podobným výsledkům, došli jsme proto k závěru, že není potřeba porovnávat užitečnost obou. Rozhodnutí o využití metody by tedy bylo přímo na uživateli a mělo by záviset na množství a složitosti použitých dat.

K porovnání užitečnosti metody GARCH byla jako benchmarková metoda využita EWMA, doporučovaná v metodice RiskMetrics jako vhodný nástroj výpočtu VaR. Tato metoda je výpočetně mnohem jednodušší a použitelná i při větším množství dat. Její nevýhodou by však mohl být nedostatečný důraz na vyšší korelaci a riziko extrémních ztrát.

To se však v práci neukázalo, protože při testování počtu překročení hodnoty VaR podle obou metod bylo lepších výsledků dosaženo metodou EWMA. A to jak při testování hodnot jednotlivých aktiv, tak u složitějšího portfolia. Metoda GARCH sice v přesnosti příliš nezaostávala, ale vzhledem k její náročnosti se v této situaci neosvědčila. Je pravděpodobné, že v případech s daleko proměnlivějšími korelacemi a volatilitou by její přesnost byla vyšší a mohla by i přesáhnout metodu EWMA. To se ale v našem případě nedokázalo a užitečnost tedy musí být předmětem dalšího zkoumání.

Malé rozdíly mezi jednotlivými metodami vedou k závěru, že by bylo praktické při odhadu v zájmu opatrnosti využívat obě metody a porovnávat jejich výsledky (pokud by taková věc byla výpočetně realistická). V práci jsme však zmínili několik nedostatků samotné metody Value at Risk a vyšší přesnost tohoto odhadu nemusí být příliš užitečná pro lepší odhad rizika.

Citovaná literatura

- (1996). Načteno z RiskMetrics Technical Document:
<https://www.msci.com/documents/10199/5915b101-4206-4ba0-ae2-3449d5c7e95a>
- Basel Committee on Banking Supervision. (2004). *Basel II: International Convergence of Capital Measurement and Capital Standards: a Revised Framework*. Načteno z Bank for International Settlements: www.bis.org/publ/bcbs107.htm
- Bollerslev, T. (1986). Generalized autoregressive conditional heteroscedasticity. *Journal of Econometrics*, stránky 307-327.
- Danielsson, J. e. (2005). Subadditivity Re-Examined: the Case for Value-At-Risk.
- Einhorn, D. (červenec 2008). Private Profits and Socialized Risk.
- Engle, R. F. (1982). Autoregressive Conditional Heteroscedasticity with Estimates of the Variance of United Kingdom Inflation. *Econometrica*, 987-1008.
- Engle, R. F., & Sheppard, K. (Říjen 2001). Theoretical and Empirical properties of Dynamic Conditional Correlation Multivariate GARCH.
- Fiorucci, J. A., Ehlers, R. S., & Louzada, F. (2016). bayesDccGarch: The Bayesian Dynamic Conditional Correlation GARCH Model. Načteno z <https://CRAN.R-project.org/package=bayesDccGarch>
- Hyde, S., Bredin, D., & Nguyen, N. (květen 2007). Correlation Dynamics between Asia-Pacific, EU and US.
- Jondeau, E., & Rockinger, M. (2000). *Conditional Volatility, Skewness, and Kurtosis: Existence and Persistence*.
- Kercheval, A. N., & Liu, Y. (2011). Risk Forecasting with GARCH, Skewed t Distributions, and Multiple Timescales. V *Modelling High-Frequency Data in Finance*. John Wiley & Sons.
- Martin, J. D., & Klemkovsky, R. C. (leden 1975). Evidence of Heteroscedasticity in the Market Model. *The Journal of Business*, stránky 81-86.
- Nakatani, T. (2014). ccgarch: An R Package for Modelling Multivariate GARCH Models with Conditional Correlations. Načteno z <http://CRAN.r-project.org/package=ccgarch>
- Rachev, S. T., Hsu, J. S., & Bagasheva, B. S. (2008). *Bayesian Methods in Finance*. New Jersey: John Wiley & Sons.
- Royset, J. O., & Wets, R. J. (2014). From Data to Assessments and Decisions: Epi-Spline Technology.
- Sandoval, L. (leden 2012). Correlation of financial markets in times of crisis. *Physica A: Statistical Mechanics and its Applications*, stránky 187-208.
- Springer Finance. (2007). Statistical Properties of Financial Market Data. *Financial Modelling Under Non-Gaussian Distributions*.

Seznam tabulek

Tabulka 4-1: Základní vlastnosti logaritmických výnosů	38
Tabulka 4-2: Korelace logaritmických výnosů	39
Tabulka 4-3: Korelace logaritmických výnosů po posunutí japonského indexu o 1 den	39
Tabulka 4-4: Odhad šikmosti a špičatosti jednotlivých řad	43
Tabulka 4-5: Kumulované exponenciální váhy pro $\lambda=0,75$	44
Tabulka 4-6: Kumulované exponenciální váhy pro $\lambda=0,94$	44
Tabulka 4-7: Kumulované exponenciální váhy pro $\lambda=0,97$	45
Tabulka 4-8: Čtvercová chyba odhadu v jednotlivých řadách pro různě zvolená λ	47
Tabulka 4-9: Výsledné parametry ML odhadu DCC GARCH	51
Tabulka 4-10: Počet porušení hodnoty Value at Risk v modelu DCC GARCH	57
Tabulka 4-11: Počet porušení VaR pro různá rozdělení	59
Tabulka 4-12: Výsledné odhady parametrů pro DCC GARCH pomocí bayesovského přístupu ..	61
Tabulka 4-13: Srovnání výsledků metody DCC GARCH v závislosti na postupu výpočtu	61
Tabulka 4-14: Počet porušení VaR portfolia pro jednotlivé přístupy	68

Seznam grafů

Obrázek 2-1: Ilustrace Value at Risk	10
Obrázek 4-1: Vývoj indexů S&P a EURONEXT v roce 2008	35
Obrázek 4-2: Vývoj indexů S&P a EURONEXT v roce 2012	35
Obrázek 4-3: Vývoj cen akciových indexů	37
Obrázek 4-4: Chování posunutých logaritmických výnosů	41
Obrázek 4-5: Histogramy posunutých logaritmických výnosů	42
Obrázek 4-6: Vývoje exponenciálních vah pro různé hodnoty λ	45
Obrázek 4-7: Vývoje kumulativních exponenciálních vah a jejich cut-off pro různá λ	46
Obrázek 4-8: Čtvercové chyby inovací v závislosti na koeficientu λ	47
Obrázek 4-9: Odhady volatilit S&P, EURONEXT a Nikkei	48
Obrázek 4-10: Odhad vývoje korelace výnosů S&P a EURONEXT pomocí metody EWMA	49
Obrázek 4-11: Odhad vývoje korelace výnosů indexů pomocí metody EWMA	49
Obrázek 4-12: Odhad vývoje Value at Risk pro logaritmický výnos indexu Nikkei	50
Obrázek 4-13: Vývoj odhadu korelací	52
Obrázek 4-14: Vývoj odhadu volatility	54
Obrázek 4-15: Vývoj odhadu Value at Risk pro výnos indexu S&P	55
Obrázek 4-16: Value at Risk pro hodnotu indexu S&P	56
Obrázek 4-17: Vývoj odhadu Value at Risk pro výnosy indexu	56
Obrázek 4-18: Porovnání odhadu VaR s předpokladem normálního rozdělení $t(6)$ -rozdělení ..	58
Obrázek 4-19: Velikost rozdílu mezi odhady s využitím normálního a Studentova rozdělení ..	58
Obrázek 4-20: Odhady volatilit - bayesovský přístup, ML	63
Obrázek 4-21: Odhady korelací - bayesovský přístup, ML	64
Obrázek 4-22: Odhad volatility portfolia v porovnání s S&P a Nikkei	66
Obrázek 4-23: Odhad VaR portfolia v porovnání se skutečnými výnosy: DCC GARCH	67
Obrázek 4-24: Odhad VaR portfolia v porovnání se skutečnými výnosy: EWMA	67

