

Regime-afhankelijke risk parity-allocatie met GARCH–HMM voor dynamisch portefeuillebeheer

Julius Sepmeijer

Begeleid door: Dhr. De Kok

Vak: Economie en Wiskunde D

Inleverdatum: 26 november 2025

Samenvatting

Your abstract.

Inhoudsopgave

1	Inleiding	4
1.1	Achtergrond en motivatie	4
1.2	Probleemstelling	4
1.3	Doelstellingen van het onderzoek	5
1.3.1	Onderzoeksvragen en hypothesen	6
1.4	Opbouw van het profielwerkstuk	6
2	Theoretisch kader	7
2.1	De efficiënte markt hypothese	7
2.1.1	Adaptieve markt hypothese en behavioral finance	8
2.2	Traditionele portfolio-optimalisatie (Markowitz, mean-variance)	8
2.2.1	De wet van grote getallen en diversificatie	9
2.2.2	Termen uit de statistiek	10
2.3	Vergelijking Markowitz-portefeuille en risk parity	11
2.3.1	Overzicht van risk parity-strategieën	12
2.4	Tijdsvariërende volatiliteitsmodellen (ARCH/GARCH)	13
2.4.1	Een uitbreiding: het EGARCH-model	14
2.4.2	Schatting via maximale likelihood	15
2.5	Aanpassingen aan asset-allocatie na GARCH en risk parity: HMM's	16
2.5.1	Regime-switching modellen en Markov Chains	17
2.5.2	Verborgen regimes	18
2.5.3	HMM voor financiële tijdreeksen	19
2.6	Combinatie HMM-GARCH: regime-switching volatiliteitsmodellen	19
2.7	Integratie van GARCH-volatiliteit in de risk parity-aanpak	20
2.7.1	Aanpassing van risk parity-gewichten per regime	21
3	Methodologie	22
3.1	Onderzoeksoepzet	22
3.1.1	Onderzoeksaanpak: kwantitatieve backtesting	22
3.1.2	Dataselectie (markten, assets, periode, frequentie)	22
3.1.3	Afhankelijke variabelen in het model	23
3.1.4	Constructie van het portefeuillemodel (stappenplan)	24
4	Empirische resultaten	25
4.1	Beschrijvende statistieken van de data	27
4.2	Resultaten van de GARCH-schattingen	27
4.3	Resultaten van de HMM / regimeclassificatie	27
4.4	Prestaties van de GARCH-risk parity-portefeuille	27
4.5	Prestaties van HMM-regime-afhankelijke risk parity	27
4.6	Vergelijking met Markowitz- en andere benchmarkportefeuilles	27
4.7	Risicoprofiel, drawdowns en stabiliteit van gewichten	27
5	Robuustheidsanalyses	27
5.1	Alternatieve GARCH- en HMM-specificaties	27
5.2	Verschillende herallocatiefrequenties	27
5.3	Impact van transactiekosten	27
5.4	Subperiode- en stressperiode-analyses	27
6	Discussie	27
6.1	Interpretatie van de resultaten	27
6.2	Relatie met Markowitz, EMH en regime-switching	27
6.3	Praktische implementeerbaarheid van (HMM-)GARCH-risk parity	27
6.4	Beperkingen van het onderzoek	27

7 Conclusie en aanbevelingen	27
7.1 Antwoord op de onderzoeksvragen	27
7.2 Theoretische en praktische implicaties	27
7.3 Aanbevelingen voor vervolgonderzoek	27
Literatuurlijst	27
A Extra tabellen en figuren	27
B Technische details van GARCH- en HMM-schatting	27
C Overzicht van gebruikte code en algoritmes	27

1 Inleiding

1.1 Achtergrond en motivatie

Financiële markten laten plotselinge schommelingen zien met regimewisselingen. Wereldwijde aandelenmarkten kennen hoge en afwisselende risico's. Zo behaalde de MSCI World-index sinds eind jaren 70 een gemiddeld jaarrendement van circa 10,4% met een jaarlijkse volatiliteit van 14,8%, maar kende in crises maximale drawdowns van meer dan 50%. Zo was er in de periode van 2008-2009 een drawdown van 54,5% (XXX). Bovendien hebben globalisering en de toename van algoritmische trading ervoor gezorgd dat marktfases en trading steeds dynamischer en efficiënter werken. Dit heeft gezorgd voor een technische verandering in handelen.

Met de komst van kunstmatige intelligentie is handelen bovendien steeds complexer geworden. Politiek en economie raakt steeds meer met elkaar verweven. Grote techreuzen, zoals Nvidia en Palantir, zorgen voor een steeds grotere bijdrage aan de markten en zijn onmisbaar geworden. Zo heeft Nvidia een geschatte marktwaarde van 5 biljoen dollar; dat is bijvoorbeeld groter dan de BBP van heel Duitsland. Maar ook ING heeft als doel om duizend van de vijftienduizend banen in Nederland te schrappen en te automatiseren met kunstmatige intelligentie binnen één jaar. De vraag naar technische expertise wordt steeds groter. Waar vroeger fundamenteel investeren nog de beste strategie was, hebben grote firma's die technisch investeren nu het voortouw genomen. Dat was vroeger nog anders.

Zo heeft Renaissance Technologies, een bedrijf opgericht door de wiskundige Jim Simons, een gemiddeld jaarlijks rendement van 66% weten te realiseren. Rond de crisistijd in 2008 was er zelfs sprake van een enorme uitschieter van 152%. Dat kwam niet omdat ze meer wisten van business dan andere experts. Ze waren echter heel goed in één ding: wiskunde. Renaissance Technologies bestaat niet uit finance-experts, maar uit natuurkundigen en wiskundigen. De kwantitatieve analisten zijn heel goed in statistiek en stochastiek, en kunnen op die manier heel goed koersbewegingen analyseren en voorspellen, zonder naar macro-economische factoren te kijken. Ze keken puur naar de data. Dat was hun kracht.

Op die manier is de financiële industrie totaal gekanteld. Dat heeft ook gezorgd voor een verandering in markt-theorieën; theorieën die vroeger nog toepasbaar waren op de markt zijn nu onbetrouwbaar en ineffectief. Markowitz' idee dat de volatiliteit van markten constant is (XXX), gaat tegen resultaten van de markt in. Markten zijn dynamischer geworden en kunnen in bepaalde regimes van hoge volatiliteit of lage volatiliteit zitten. Ook zijn markten op sommige momenten efficiënt en op sommige momenten inefficiënt. Dit onderzoek gaat in op de wiskundige aspecten van koersbewegingen en gaat mee in de ontwikkelingen omtrent de technische vooruitgang. Zo richt het onderzoek zich op het ontwikkelen van een nieuw kwantitatief model voor de asset-allocatie door gebruik te maken van HMM's en een geïntegreerd GARCH-model. Hierover volgt in het vervolg een uitgebreidere uitleg.

1.2 Probleemstelling

Een verantwoording van het probleem is belangrijk voor het onderzoek, omdat daarmee de urgentie van het model wordt benadrukt. Het probleem is dan ook dat traditionele portefeuille-optimalisatiemethoden, zoals de klassieke Mean-Variance-benadering van Markowitz onvoldoende rekening houden met de dynamiek en onzekerheid in financiële markten. Rendementen zijn vaak niet normaal verdeeld en parameters zijn niet constant, maar stochastisch. Dit model zal in veel situaties werken, maar toch zijn er praktische complicaties. Zo is uit meerdere onderzoeken af te leiden dat het model in de praktijk tot tegenstrijdigheden heeft geleid (XXX). Zo komen extreme rendementen vaker voor dan het Mean-Variance model laat zien. Het model schiet dus te kort op veel vlakken. Er is een genuanceerdere benadering nodig. Er is behoefte aan een methode die deze onzekerheden meeneemt en rekening houdt met mogelijke regime wisselingen.

Concreet is het doel van een hedgefonds om de *sharpe ratio* zo hoog mogelijk te houden. Twee dingen zijn daarbij van belang: het portfolio moet zo min mogelijk risico nemen (1) en het rendement moet zo hoog mogelijk zijn (2). Vaak betekent minder risico ook minder rendement. Een fonds moet de perfecte verhouding vinden tussen risico en rendement. Hedgefondsen vinden het belangrijker om het risico bijna helemaal uit te sluiten tegen een normaal rendement, dan om veel risico te nemen met een hoog rendement. Zo zal een kwantitatieve analist direct 1 euro pakken

als er geen risico is, dan 100 euro met 10% kans op winst, als er één keer te kiezen is. Dat komt omdat analisten zoeken naar een bepaalde zekerheid tijdens handelen. Dat is lastig, want markten bewegen stochastisch en zijn daarom dynamisch. Kortom: de sharpe ratio moet zo hoog mogelijk zijn, door het rendement zo hoog mogelijk te houden tegen een lage variantie. De sharpe ratio is te formuleren op de volgende manier:

$$S = \frac{E[R_p] - R_f}{\sigma(R_p)},$$

waarbij $E[R_p]$ het verwachte rendement van het portfolio voorstelt en R_f het rendement zonder risico. Met de R_f wordt meestal de rente van een staatsobligatie bedoeld, omdat dat een rendement zonder risico is. Door ze van elkaar af te trekken, kom je op het totale deel van het rendement mét risico. Een portfolio kan dan wel 3% rendement hebben, maar dat stelt niks voor, omdat dat kapitaal ook in een staatsobligatie had kunnen zitten, waarbij de investeerder geen risico had hoeven te nemen. $\sigma(R_p)$ staat voor de variantie, of ook wel volatiliteit genoemd, van een asset of verzameling van assets binnen een portfolio. Onder dit begrip wordt de mate van schommelingen bedoeld binnen het portfolio P . Logisch is dat als een portfolio een hoge variantie heeft, dat dat minder gunstig is voor een hedgefonds. Dat komt omdat daarbij een grotere risico komt kijken en onzekerheid groter wordt. Concluderend, er volgt uit de vergelijking dat een lagere variantie en hogere $E[R_p] - R_f$ voor een grotere sharpe ratio zorgt. Een goede sharpe ratio ligt rond de 1, maar professionele hedgefondsen op de wereld halen een ratio van 3 (XXX). Dat is uitzonderlijk, want daarbij moet dus het rendement extreem hoog zijn in combinatie met de lage variantie. Dat zijn twee factoren die moeilijk zijn te combineren.

Met deze informatie kan een probleemstelling worden geformuleerd. Het is in de huidige markt steeds lastiger om de variantie van een verzameling assets te voorspellen, omdat technische innovatie tot een dynamische en complexe markt heeft geleid. Er is technische expertise vereist en de vraag naar nieuwe modellen die de markt begrijpen is enorm toegenomen. Door beroep te doen op oude modellen en wiskundigen kan een nieuw wiskundig model worden opgesteld dat de huidige markt beter begrijpt. Met de nieuwe AI (bubbel) trends zijn dit soort modellen onmisbaar. Oude modellen kunnen varianties niet goed voorspellen (XXX) en onbreken daarin in hun diepgang.

1.3 Doelstellingen van het onderzoek

Op basis van de probleemstelling kan een doelstelling voor het onderzoek worden opgesteld. Dit onderzoek wil GARCH-modellen combineren om de volatiliteit van een asset binne het portfolio te voorspellen en vervolgens op basis van risk parity een portfolio samen te stellen. HMM zorgen daarbij voor een nuance in de samenstelling, omdat dit kern-onderdeel van het model zich bezighoudt met het rendement van het portfolio. Het GARCH-onderdeel houdt zich bezig met de volatiliteit, het onderste deel van de breuk van de sharpe-ratio vergelijking. Het doel is dat hierbij de sharpe ratio rond de 1 zit. Om gericht onderzoek te kunnen doen, zijn er voorwaarden opgesteld.

Er is sprake van een eerlijke en bias-vrije opzet (1). Deze voorwaarde gaat eigenlijk over dat er niet vals gespeeld mag worden. Er mogen alleen dingen gebruikt worden die op dat moment bekend waren. Er wordt niet in de toekomst gekeken bij bijvoorbeeld backtesting. Er wordt alleen oude data gebruikt om het model te trainen en daarna andere, nieuwe data om te kijken of het model het goed doet. Ook zijn er consistente en realistische model- en portefeuillekeuzes (2). Er wordt duidelijk aangegeven hoe het model werkt en het wordt niet steeds veranderd om mooiere resultaten te krijgen. Een aandeel mag niet alles zijn, er wordt niet oneindig veel geld geleend en er wordt niet elke seconde gehandeld. Dat zorgt voor grotere transactiekosten en maakt het model complexer. Handelen per dag of maand zorgt voor een beter overzicht. Er gelden duidelijke grenzen voor gewichten, leverage en turnover. Tot slot bestaat er een transparante evaluatie en robuustheid (3). Resultaten worden gerapporteerd en transactiekosten worden meegenomen in het totale plaatje. De resultaten worden vergeleken met andere strategieën en er worden sensitiviteits- en robuustheidsanalyses uitgevoerd voor belangrijke aannames. Zo zou het model ook moeten werken als er kleine aanpassingen worden gedaan.

Zolang aan deze drie voorwaarden wordt voldaan, kan er een betrouwbare en valide conclusie worden gedaan. De resultaten worden op die manier verwerkt en vergeleken met andere modellen.

Bij een positief resultaat kan het model ook in de praktijk worden gebruikt om een bijdrage te leveren aan de technische ontwikkelingen.

1.3.1 Onderzoeksvragen en hypotheses

Nu de probleem- en doelstelling bekend zijn, kan om doelgericht onderzoek te doen een hoofdvraag worden geformuleerd. De hoofdvraag wordt op de volgende manier geformuleerd: ***“Hoe goed presteert een beleggingsstrategie die zich aanpast aan veranderende marktvolatiliteit, en is deze een aantrekkelijk alternatief voor sparen en obligaties wanneer we ook transactiekosten en praktische beperkingen meenemen?”*** Deze hoofdvraag gaat in op drie elementen. Allereerst onderzoekt het of het model überhaupt praktisch gezien lucratief kan zijn. Ook kan op deze manier worden vergeleken met bestaande modellen en houdt het rekening met de drie voorwaarden, zoals geformuleerd in paragraaf 1.3. Het is belangrijk dat deze drie kernelementen terugkomen in het onderzoek.

Op basis van deze hoofdvraag kunnen ook deelvragen worden geformuleerd om binnen het onderzoek een richting te vinden. Er is gekozen om de onderstaande selectie van deelvragen voor het onderzoek mee te nemen. Door de deelvragen uit te werken kan de hoofdvraag worden beantwoord.

1. *Welke theoretische concepten rond portefeuilletheorie, volatiliteit, risk parity en markregimes vormen de basis voor het ontwikkelen van een kwantitatief portefeuillemodel?*
2. *Hoe kunnen GARCH- en HMM-modellen worden gespecificeerd, geschat en beoordeeld zodat zij bruikbare schattingen van volatiliteit en markregimes leveren voor portefeuillebeheer?*
3. *Hoe kunnen volatilitets- en regime-inschattingen worden vertaald naar een risk parity-aanpak met regime-afhankelijke portefeuillegewichten en duidelijke rebalancing-logica?*

Voor deze drie deelvragen is gekozen omdat ze logisch aansluiten bij de opbouw en het doel van het onderzoek. De eerste deelvraag richt zich op de theoretische basis: voordat een model gemaakt kan worden, moet duidelijk zijn welke concepten uit portefeuilletheorie, volatiliteit, risk parity en markregimes relevant zijn. De tweede deelvraag focust op de kernmodellen (GARCH en HMM), omdat het essentieel is te onderzoeken of en hoe deze modellen betrouwbare inschattingen van volatiliteit en markregimes kunnen leveren. De derde deelvraag maakt tenslotte de brug naar de praktijk door te kijken hoe die inschattingen worden omgezet in concrete, regime-afhankelijke risk parity-portefeuillegewichten. Samen vormen de deelvragen zo een logische keten van theorie naar modellering en uiteindelijk toepassing in portefeuilleconstructie. De antwoorden op deze vragen leiden tot een betrouwbaar en valide resultaat.

De hypothese (H1) is dat een kwantitatief portefeuillemodel dat gebruikmaakt van volatilitetsinschattingen (GARCH) en markregimes (HMM) in een risk parity-structuur een hoger risico-gewogen rendement behaalt (bijvoorbeeld gemeten met de Sharpe-ratio) dan eenvoudige benchmarkstrategieën, zoals een buy-and-hold of een gelijkgewogen portefeuille, na aftrek van transactiekosten. De nulhypothese (H0) veronderstelt overigens dat er geen significant verschil in risico-gewogen rendement bestaat tussen het voorgestelde GARCH–HMM–risk parity-model en eenvoudige benchmarkstrategieën. Deze nulhypothese is een standaardaanname die weerlegt moet worden. In de eerste instantie gaan we ervan uit dat er geen effect of verschil is. Pas als er met data genoeg bewijs geleverd is, kan worden geconcludeerd dat H0 niet blijkt te kloppen en alternatief H1 aannemelijker is.

1.4 Opbouw van het profielwerkstuk

Het profielwerkstuk focust zich op het minimaliseren van het risico tijdens investeren. Het opgestelde model zal uiteindelijk bestaan uit een geïntegreerd Hidden Markov Model en een GARCH-onderdeel dat de volatiliteit voorspelt dat resulteert in een allocatie op basis van risk parity. Allereerst wordt de theorie over markten en informatie besproken, worden HMM's en GARCH modellen grondig onderzocht om uiteindelijk in het hoofdstuk Methodologie een eigen model te maken. De resultaten van het model worden besproken en vergeleken met andere modellen in de discussie. Tot slot zal de nulhypothese worden verworpen en wordt kort een conclusie gegeven.

2 Theoretisch kader

In dit hoofdstuk wordt het theoretisch kader geschetst dat nodig is om het verdere onderzoek te kunnen begrijpen en onderbouwen. Het hoofdstuk begint met een toelichting over hoe markten bewegen op basis van informatie. Later wordt ingegaan op de klassieke portefeuilletheorie van Markowitz, waarin de relatie tussen rendement, risico en diversificatie centraal staat. Vervolgens komen risk parity-strategieën, volatiliteit en verschillende risicomaten aan bod, evenals de efficiënte markthypothese en de vraag in hoeverre markten voorspelbaar zijn. Ten slotte wordt ingegaan op tijdsvariërende volatiliteitsmodellen, zoals ARCH/GARCH, en op het concept van marktregimes met behulp van Hidden Markov Models. Op basis van dit overzicht wordt antwoord gegeven op de deelvraag: welke theoretische concepten rond portefeuilletheorie, volatiliteit, risk parity en marktregimes vormen de basis voor het ontwikkelen van een kwantitatief portefeuillemodel? Aan het einde wordt de informatie samengevat, terwijl antwoord wordt gegeven op de deelvraag.

2.1 De efficiënte markt hypothese

De basis van asset-allocatie ligt bij het verspreiden van het eigendom van assets. Dat kan op een efficiënte manier worden gedaan als er bruikbare informatie aanwezig is. Op basis van deze informatie kunnen beslissingen voor het portfoliomanagement worden genomen. We gaan dan uit van een markt waarin alle beschikbare informatie volledig de markt reflecteert. Zo'n markt noemt men efficiënt. Dergelijke markten kunnen op twee manieren worden geanalyseerd: fundamenteel en technisch. Wanneer beleggers fundamenteel beleggen kijken ze naar fundamentele factoren in de analyses, waaronder de intrinsieke waarde van een bedrijf. Technische analyses focussen zich vooral op trends en statistiek. Ons onderzoek focust zich alleen op de technische kant en laat de fundamentele kant volledig buiten beschouwing. Daarover volgt geen verdere informatie.

In het artikel *Efficient Capital Markets: A Review of Theory and Empirical work* (XXX) gaat Eugene Fama in op de efficiënte markt hypothese (hierna: EMH). Hij begint zijn artikel door onderscheid te maken tussen verschillende soorten efficiënte markten¹. Zo kan er een zwakke vorm aanwezig zijn waarbij alleen historische data de markt reflecteert. Ook bestaat er een variant, de semi-sterke variant, waarbij de markt ook wordt beïnvloed door publiekelijke informatie, bijvoorbeeld door het nieuws. Tot slot bestaat er een sterke variant waarbij monopolistische instituten alle informatie hebben die volledig de markt beïnvloed, waaronder ook privé informatie. Hierbij bestaan geen factoren die toegevoegde waarde bieden voor de markt. Zo kan een markt in verschillende gradaties van efficiëntie zitten. Wanneer informatie volledig wordt verwerkt in markt, kan dit voor complicaties zorgen tijdens het investeren. Dat betekent dat informatie direct verwerkt wordt in de markt. Wanneer bijvoorbeeld nieuwe informatie bekend is, wordt dat direct verwerkt en kunnen investeerders te laat reageren. Dat zorgt voor uitdagingen binnen het portfoliomanagement. Het is moeilijk om systematisch bovengemiddelde rendementen te behalen.

Voor dit onderzoek en het te ontwikkelen GARCH-HMM-risk-paritymodel is de EMH om meerdere redenen relevant. Als markten (minstens) in zwakke of semi-sterke vorm efficiënt zijn, is het lastig om structureel extra rendement te behalen door simpelweg publieke informatie of historische koersen te gebruiken. In plaats daarvan verschuift de focus van het voorspellen van rendementen naar het modelleren en beheren van risico. Het in dit profielwerkstuk voorgestelde model probeert dan ook niet om systematisch misprijzingen te verwerken, maar om de tijdsvariërende volatiliteit en regimeverschuivingen in de markt zo goed mogelijk te kwantificeren. GARCH-modellen sluiten aan bij het idee dat rendementen op korte termijn moeilijk voorspelbaar zijn, terwijl volatiliteit wél een zekere structuur vertoont. Hidden Markov Models vullen dit aan door expliciet rekening te houden met verschillende marktregimes (bijvoorbeeld “rustig” versus “crisis”), zonder te veronderstellen dat deze regimes triviaal te voorspellen zijn op basis van publiek beschikbare informatie. Binnen de logica van de EMH is het dus zinvol om een model te bouwen dat zich richt op een efficiëntere risicotoedeling (via risk parity), in plaats van op het najagen van bovengemiddelde rendementen, omdat juist risicobeheersing en drawdownvermindering realistische en haalbare doelstellingen blijven in (semi-)efficiënte markten.

¹Dit onderscheid van efficiënte markten werd voor het eerst ter sprake gebracht door Harry Roberts.

2.1.1 Adaptieve markt hypothese en behavioral finance

Toch zegt de efficiënte markthypothese niet alles. Markten fluctueren vaak in efficiëntie en zijn dus niet constant efficiënt of inefficiënt. Dit idee sluit aan bij de *Adaptive Markets Hypothesis* (XXX) (hierna: AMH). In dit kader wordt de markt gezien als een adaptief systeem waarin beleggers zich gedragen als economische “soorten” die zich voortdurend aanpassen aan veranderende omstandigheden. In plaats van volledig rationele beleggers, zoals de EMH veronderstelt, gaat de AMH uit van begrensde rationaliteit, leerprocessen en competitieve selectie. Strategieën die in één periode succesvol zijn, kunnen in een andere periode hun werking verliezen, omdat andere beleggers zich aanpassen of omdat marktomstandigheden veranderen. Emoties en gedragfouten spelen daarmee een structurele rol in prijsvorming en marktdynamiek, waardoor efficiëntie tijdsafhankelijk wordt in plaats van absoluut. Dat resulteert onder andere in stochastische koersbewegingen.

Voor dit onderzoek en het te ontwikkelen model is de AMH bijzonder relevant. Als markten zich adaptief gedragen en de mate van efficiëntie over de tijd fluctueert, dan is het aannemelijk dat ook volatiliteit, correlaties en risicopremies regime-afhankelijk zijn. Dit sluit direct aan bij het gebruik van een Hidden Markov Model, waarin verondersteld wordt dat de markt zich in verschillende, niet-direct observeerbare toestanden kan bevinden (bijvoorbeeld een “rustig” regime en een “stress”- of “crisis”-regime). In combinatie met GARCH maakt dit het mogelijk om niet alleen tijdsvariërende volatiliteit te modelleren, maar deze ook expliciet te koppelen aan wisselende marktcondities en gedragsmatige factoren, zoals angst en hebzucht van beleggers. Vanuit het perspectief van de AMH is een statisch risk-paritymodel onvoldoende, omdat het impliciet uitgaat van relatief stabiele en rationele marktomstandigheden. Een GARCH-HMM-gebaseerde risk-paritybenadering sluit beter aan bij een wereld waarin markten zich aanpassen, beleggers niet volledig rationeel zijn en efficiëntie dynamisch en stochastisch is. Het model gaat mee in de doelstelling van het onderzoek.

2.2 Traditionele portfolio-optimalisatie (Markowitz, mean-variance)

De EMH en AMH beschrijven vooral hoe prijzen tot stand komen en in welke mate informatie en gedrag van beleggers al in koersen zijn verdisconteerd. Voor een belegger die vandaag een portefeuille moet samenstellen, is echter vooral de vraag relevant welke combinatie van assets past bij een gewenst risico- en rendementsprofiel. In de praktijk komt men dan al snel uit bij de klassieke mean-variance-benadering van Markowitz. Deze theorie gaat uit van verwachte rendementen, varianties en covarianties als input voor de portefeuillesamenstelling en levert een efficiënte grens op waaruit een belegger kan kiezen. Juist hier ontstaat een spanningsveld met de eerder besproken marktbeelden: als volatiliteiten en correlaties door regimes en gedrag sterk in de tijd veranderen, dan zijn de statische aannames van Markowitz problematisch. Dit vormt een directe aanleiding om verder te kijken dan het klassieke Markowitz-model en meer dynamische risicomodellen, zoals GARCH en HMM, te onderzoeken binnen de context van portefeuilleselectie.

In een van Markowitz’ belangrijkste artikelen over de mean-variance-theorie (XXX) begon hij met het onderverdelen van verschillende fasen in het proces van het kiezen van een portfolio. Zo bestaat dit proces volgens Markowitz uit twee fasen. De eerste fase gaat volgens hem over het observeren van een asset en het vormen van een bepaald geloof over deze asset in de toekomst. De tweede fase begint bij dat geloof over de asset in de toekomst en leidt tot het daadwerkelijk kiezen van een portfolio. Zijn onderzoek richt zich voornamelijk op deze tweede fase: het opstellen van een portfolio. Het gaat hier dus over de asset-allocatie. Daarnaast worden bepaalde voorwaarden besproken voor het onderzoek. Zo wordt genoemd dat een belegger zou moeten proberen het verwachte rendement te maximaliseren, ongeacht risico of spreiding. Markowitz wijst deze regel echter vrijwel direct af, omdat dit in strijd is met zijn hypothese om beleggingsgedrag te verklaren. Onderzoek laat namelijk zien dat beleggers zich niet alleen laten leiden door het hoogste rendement, maar ook rekening houden met risico en spreiding. Ook wordt een andere voorwaarde besproken: een hoger rendement is wenselijk en een hogere variantie is onwenselijk. Hij duidt dus op een voorwaarde die niet uitsluitend gefocust is op hoog rendement. Deze voorwaarde is veel realistischer en sluit goed aan op de werkelijkheid. Bovendien gaat het hand in hand met de doelstelling van het profielwerkstuk.

Ook gaat hij later nog een keer in op de eerste, verworpen voorwaarde. Zo zegt hij dat het maximaliseren van het gediscoteerde verwachte rendement geen diversificatie oplevert. *Gediscoteerd* impliceert het verwachte rendement, omgerekend naar vandaag. Dus een asset kan dan wel

110\$ waard zijn in de toekomst, maar omgerekend naar vandaag slechts 100\$. Hij geeft aan dat een investeerder al zijn geld zal investeren in de asset met de hoogste gediscoteerde waarde. Wanneer er twee assets zijn met de hoogste gediscoteerde waarde, dan zal het vermogen verdeeld worden, en is elke verdeling goed. Hij onderbouwt dat door uit te gaan van een portfolio met N effecten. Hij stelt een verwacht rendement r_{it} voor in tijd t met een asset i ; ook bestaat er een disconteringsfactor d_{it} die het rendement terugreken naar een waarde voor vandaag; en tot slot bestaat een relatieve bijdrage X_i van de asset die de investeerder stopt in het portfolio. Omdat er geen short sales zijn toegestaan, geldt voor alle posities dat $X_i \geq 0$. Hierbij staat X_i voor het (relatieve) bedrag dat in asset i wordt geïnvesteerd. Het totale gediscoteerde verwachte rendement van de portefeuille kan dan als volgt worden geformuleerd:

$$R = \sum_{t=1}^{\infty} \sum_{i=1}^N d_{it} r_{it} X_i = \sum_{i=1}^N X_i \left(\sum_{t=1}^{\infty} d_{it} r_{it} \right).$$

In deze uitdrukking staat r_{it} voor het verwachte rendement van asset i in periode t , en d_{it} voor de bijbehorende disconteringsfactor die het rendement in periode t terugreken naar tijdstip 0. De binnenste som $\sum_{t=1}^{\infty} d_{it} r_{it}$ berekent dus de contante waarde van alle toekomstige verwachte rendementen van asset i . De buitenste som over i telt vervolgens de bijdragen van alle assets bij elkaar op, gewogen met hun portefeuillegewicht X_i . Om dit compacter te noteren, wordt vaak eerst het gediscoteerde verwachte rendement per individuele asset gedefinieerd als

$$R_i = \sum_{t=1}^{\infty} d_{it} r_{it}$$

Dit geeft ons weer nieuwe informatie. R_i is onafhankelijk van X_i . Ook is bekend dat de som van alle X_i gelijk is aan 1 (namelijk 100%) en dat R een gewogen gemiddelde is van R_i met $X_i \geq 0$. Om R te maximaliseren, stellen we $X_i = 1$ met i een maximum R_i . Als verschillende R_a , met $a = 1, \dots, K$ maxima zijn, dan zal elke allocatie in het portfolio R maximaliseren. Dat betekent dat als er meerdere assets de hoogste gediscoteerde waarde hebben, dan maakt het niet uit hoe de spreiding is. Dat impliceert dat een gespreid portfolio in geen situatie meer gewaardeerd zou moeten zijn dan een niet-gespreid portfolio. Deze situatie wordt op de volgende manier geformuleerd:

$$R_{\max} = \max_X \sum_{i=1}^N X_i R_i \quad \text{zodat} \quad \sum_{i=1}^N X_i = 1, \quad X_i \geq 0,$$

$$\Rightarrow R_{\max} = \max_{1 \leq i \leq N} R_i \quad \text{bij} \quad X_i = 1 \text{ voor een } i \text{ met } R_i = \max_j R_j.$$

Uit het voorgaande blijkt dat, als je alleen kijkt naar welk aandeel het hoogste verwachte rendement heeft, je in theorie al je geld in dat ene aandeel zou stoppen. In zo'n benadering wordt een gespreide portefeuille dus niet beloond: één asset met het hoogste verwachte rendement wint altijd. Voor dit profielwerkstuk is dat belangrijk, omdat het laat zien dat we niet alleen naar rendement kunnen kijken. Het GARCH-model gaat juist over risico en schommelingen in de tijd (volatiliteit). Wil je GARCH echt gebruiken bij het kiezen van een portefeuille, dan moet risico dus een rol spelen in de beslissing. Pas als we naast rendement ook rekening houden met risico en spreiding, wordt een gespreide portefeuille aantrekkelijk en heeft het zin om een GARCH-model te gebruiken om de portefeuillegewichten te bepalen. Op die manier rijkt de sharpe-ratio zo hoog mogelijk.

2.2.1 De wet van grote getallen en diversificatie

Toch kiezen investeerders vaak voor diversificatie en daar is ook een grondige reden voor, aldus Harry. Zo bestaat er een regel die zowel diversificatie als maximalisatie van verwacht rendement aanbeveelt. Volgens deze regel verdeelt een belegger zijn vermogen over alle assetklassen die het hoogste verwachte rendement opleveren. De wet van grote getallen zou dan ervoor zorgen dat het werkelijke rendement van de portefeuille dicht bij het verwachte rendement in de buurt ligt. De regel sluit aan bij de expected returns-variance (hierna: EV) regel, die uitgaat van een portefeuille die zowel maximaal rendement als minimale variantie biedt. Harry geeft aan dat deze EV-regel niet geaccepteerd kan worden. Hij gaat daarbij in op de covariantie tussen de assets en

dat diversificatie niet alle variantie kan elimineren. Beleggers kunnen op die manier een trade-off maken, door bijvoorbeeld een deel van het risico te accepteren voor hoger rendement. Ook kunnen beleggers minder risico accepteren en een deel van het rendement opgeven. Op deze manier wordt er een nuance gebracht in het EV model van Markowitz. Het doel van ons profielwerkstuk is om zo min mogelijk risico te zoeken tegen een rendement dat aanzienlijk groter is dan R_f , zodat de sharpe ratio een aangename waarde bereikt ($S \geq 1$). Markowitz leert ons dat assets gekozen moeten worden op basis van verwacht rendement en variantie en dat daarin een goede verhouding gevonden moet worden. Toch mist er diepgang in het artikel van Markowitz op de gebieden van dynamica van markten. De volatiliteit is niet statisch, maar dynamisch (XXX).

2.2.2 Termen uit de statistiek

Markowitz begint ook met een introductie van enkele elementaire begrippen uit de statistiek. Zo heeft hij het over een variabele Y die afhangt van een kans. Hij gaat uit dat Y allerlei waardes $y_1, y_3, \dots, y_N$ kan aannemen. Hij stelt de kans van $Y = y_n$ gelijk aan p_n . Het verwachte rendement kan dan worden berekend met de volgende formule:

$$E(Y) = \sum_{i=1}^N p_i y_i$$

Dit is intuïtief te begrijpen. Neem bijvoorbeeld een eerlijke dobbelsteen met de uitkomsten 1, 2, 3, 4, 5, 6. De verwachtingswaarde $E(Y)$ van één worp is dan het gewogen gemiddelde van alle mogelijke uitkomsten, waarbij iedere uitkomst een kans van $\frac{1}{6}$ heeft. We krijgen dan:

$$E(Y) = 1 \times \frac{1}{6} + 2 \times \frac{1}{6} + 3 \times \frac{1}{6} + 4 \times \frac{1}{6} + 5 \times \frac{1}{6} + 6 \times \frac{1}{6} = 3,5.$$

Dit betekent niet dat de dobbelsteen ooit daadwerkelijk 3,5 zal laten zien, maar dat 3,5 het gemiddelde resultaat is dat we op lange termijn verwachten als we de dobbelsteen heel vaak zouden gooien. De verwachtingswaarde is dus een maat voor het gemiddelde gedrag van een willekeurige uitkomst. Precies hetzelfde principe kan worden toegepast op financiële markten. In plaats van dobbelsteenuitkomsten kijken we dan naar mogelijke rendementen van een asset, met bijbehorende kansen. De verwachtingswaarde vertegenwoordigt dan het verwachte rendement van een belegging, als gewogen gemiddelde van alle mogelijke scenario's.

Naast de verwachtingswaarde is het echter ook belangrijk om te weten hoe sterk de uitkomsten rond dit gemiddelde schommelen. Twee beleggingen kunnen hetzelfde verwachte rendement hebben, maar toch heel verschillend risico laten zien. Hiervoor gebruiken we de variantie: een maat voor hoe ver de uitkomsten gemiddeld van de verwachtingswaarde af liggen. Er bestaat een standaardformule voor de variantie. Deze formule kan op de volgende manier worden geformuleerd:

$$V(Y) = \sum_{i=1}^N p_i (y_i - E(Y))^2$$

Hierbij weten we dat het totale rendement van de portefeuille afhangt van de gewichten a_i en de individuele rendementen R_i van de verschillende assets. Elk gewicht a_i geeft aan welk deel van het totale vermogen in asset i wordt belegd. Hoe groter het gewicht van een bepaalde asset, hoe groter de bijdrage van het rendement van die asset aan het totale rendement. Het verwachte rendement van de portefeuille kan daarom worden berekend op basis van de afzonderlijke verwachte rendementen van de assets. In plaats van het rendement van de hele portefeuille in één keer te benaderen, kijken we eerst naar het verwachte rendement van iedere asset afzonderlijk en nemen vervolgens een gewogen gemiddelde. Dit levert de bekende relatie op:

$$E(R_p) = \sum_{i=1}^N a_i E(R_i),$$

waarbij $E(R_p)$ het verwachte rendement van de portefeuille is, $E(R_i)$ het verwachte rendement van asset i en $\sum_{i=1}^N a_i = 1$ (het totale belegde vermogen is 100%). Ter illustratie: als 60% van het

vermogen wordt belegd in een asset met een verwacht rendement van 5% en 40% in een asset met een verwacht rendement van 3%, dan is het verwachte portefeuillerendement gelijk aan

$$E(R_p) = 0,6 \cdot 0,05 + 0,4 \cdot 0,03,$$

In een portefeuille met meerdere assets is het niet genoeg om alleen naar de afzonderlijke varianties (risico's) van de rendementen te kijken (XXX). Minstens zo belangrijk is de vraag hoe deze rendementen zich samen bewegen. Dit gezamenlijke gedrag wordt vastgelegd in de covariantie tussen twee assets. Als $\text{Cov}(R_i, R_j) > 0$ dan hebben de twee assets de neiging om in dezelfde richting te bewegen. Als $\text{Cov}(R_i, R_j) < 0$ dan bewegen ze gemiddeld genomen juist in tegenovergestelde richting. Bij een covariantie dichtbij nul is er weinig lineair verband tussen de twee rendementen. Wanneer we met N verschillende assets werken, is het onhandig om alle covarianties één voor één op te schrijven. In plaats daarvan verzamelen we alle varianties en covarianties in één object: de covariantiematrix. De covariantiematrix Σ wordt dan gedefinieerd als

$$\Sigma = \text{Cov}(R) = \begin{pmatrix} \text{Cov}(R_1, R_1) & \text{Cov}(R_1, R_2) & \cdots & \text{Cov}(R_1, R_N) \\ \text{Cov}(R_2, R_1) & \text{Cov}(R_2, R_2) & \cdots & \text{Cov}(R_2, R_N) \\ \vdots & \vdots & \ddots & \vdots \\ \text{Cov}(R_N, R_1) & \text{Cov}(R_N, R_2) & \cdots & \text{Cov}(R_N, R_N) \end{pmatrix}.$$

Hierbij staan de elementen op de diagonaal, $\text{Cov}(R_i, R_i)$, gelijk aan de variantie van asset i : $\text{Cov}(R_i, R_i) = \text{Var}(R_i)$. De elementen buiten de diagonaal geven de covariantie tussen twee verschillende assets i en j weer. In portefeuilletheorie speelt de covariantiematrix een centrale rol. Stel dat de gewichten van de portefeuille in een vector $\mathbf{a}$ staan:

$$\mathbf{a} = \begin{pmatrix} a_1 \\ a_2 \\ \vdots \\ a_N \end{pmatrix}, \quad \text{met} \quad \sum_{i=1}^N a_i = 1.$$

We zijn nu geïnteresseerd in het risico van de hele portefeuille. Dit risico drukken we uit in de variantie van de portefeuille. De gewichten van de portefeuille staan in de vector $\mathbf{a}$ en alle varianties en covarianties staan in de matrix Σ . Met deze notatie kan de variantie van het portefeuillerendement compact worden geschreven als

$$\text{Var}(R_p) = \mathbf{a}^\top \Sigma \mathbf{a}.$$

Deze formule laat zien dat het portefeuillerisico niet alleen afhangt van de individuele risico's (varianties) van de assets, maar ook van de onderlinge samenhang (covarianties). Juist door assets te combineren die niet perfect met elkaar meebewegen (lage of negatieve covariantie), kan het totale risico van de portefeuille lager worden dan de simpele som van de individuele risico's. Dit is de kern van diversificatie.

In modellen zoals dat van Markowitz, en ook in meer geavanceerde modellen met GARCH en HMM, vormt deze matrix de basis voor het berekenen en beheersen van het portefeuillerisico. De EV-regel van Markowitz gaat ook in op andere - vooral complexe - situaties die niet relevant zijn en niet aansluiten op de doelstelling van dit onderzoek. De statistische hulpmiddelen die Markowitz aan de lezer reikt, sluiten echter aan bij de doelstelling van het begrip investeren en diversificatie. Deze informatie wordt dan ook meegenomen in het maken van het model, waarbij het model wordt vergeleken met andere modellen.

2.3 Vergelijking Markowitz-portefeuille en risk parity

Voor het maken van het model is het cruciaal om het verband tussen risk parity en de ideologie van Markowitz te begrijpen. Daarom wijdt deze paragraaf zich toe aan het uitpluizen van de urgentie van de twee theorieën en hoe deze theorieën samen gebruikt kunnen worden.

Markowitz' idee gaat vooral over het risico dat je neemt en hoe je het verwachte rendement dan kan maximaliseren of als je een bepaald verwacht rendement hoe je de variantie kan minimaliseren. Deze theorie hangt dus heel erg af van verwachte rendementen en wanneer deze in de praktijk een beetje afwijken, kunnen er al grote verliezen optreden. Als er kleine fouten in de input zijn,

zijn er ook niet optimale wegingen aanwezig. Het portefeuille is instabiel. *Risk parity* aan de andere kant gaat over een verdeling van assets in een portefeuille zodat elke asset evenveel risico bedraagt (XXX). Zo zal een asset met een hoge variantie een kleinere bijdrage krijgen binnen het portefeuille. Een asset met een lage variantie krijgt een grotere bijdrage. Geen enkele belegging mag het risico domineren. Risk parity gebruikt dus alleen volatiliteiten en correlaties en geen verwachte rendementen.

Bij risk parity hanteer je vaak een vaste doelvolatiliteit (bijv. 10% per jaar) (XXX) en stuur je mechanisch op risico: als de gemeten of verwachte volatiliteit oploopt, schaal je alle posities terug om rond die target te blijven. Dat maakt het risicobeheer regelgebaseerd. Moderne risk-paritystrategieën denken daarbij in risicofactoren (zoals value, momentum, duration, equity-beta, inflatie) in plaats van losse producten, en actief risicobeheer betekent dan het herverdelen van risicobudget tussen deze factoren.

De doelstelling van het profielwerkstuk was een S groter of gelijk aan 1. Door risk parity toe te passen wordt naar het hoofddoel gestreven: het risico minimaliseren. Toch is het ook belangrijk om, zoals Markowitz' theorie erkent, te kijken naar het verwachte risico. Dan moet er dus wél een sterke analyse zijn voor het verwachte rendement, omdat kleine fouten al snel tot grote verliezen kunnen leiden op de lange termijn. Risk parity als fundament is dus het meest veilig en sluit aan bij het profielwerkstuk. Er wordt gekozen voor een model met een geïntegreerd HMM dat kleine aanpassingen doet aan het risk parity portefeuille en dus ook de potentie op stijgen of dalen van een asset in rekening brengt. Hoe klein deze aanpassingen zullen zijn en hoe deze berekend worden, wordt in het hoofdstuk Methodologie uitgewerkt.

2.3.1 Overzicht van risk parity-strategieën

Er zijn verschillende soorten risk parity-strategieën. In deze paragraaf worden alle strategieën besproken en wordt afgewogen welke strategie het beste bij dit onderzoek past.

Het doel van klassieke risk parity is een gelijke risico bijdrage per asset binnen het portfolio (XXX). Hierbij wordt meestal alleen een long-only portefeuille opgebouwd over enkele brede asset classes. Vaak bestaat er leverage op defensieve assets, zoals obligaties, om een doelvolatiliteit te behalen. Deze vorm van risk parity staat bekend als equal risk contribution (hierna: ERC): elk onderdeel van het portefeuille levert een even groot deel risico. Voor een portefeuille X met n assets geldt dan: $RC_1 = RC_2 = \dots = RC_n$. Het risico is gelijk verdeeld over het portefeuille, zodat de totale variantie van het portefeuille wordt geminimaliseerd. Met de formule voor deze totale variantie is de logica achter de klassieke risk parity theorie logisch af te leiden. Deze formule is kan worden opgeschreven als:

$$\sigma_P^2 = \sum_{i=1}^n \sum_{j=1}^n w_i w_j \sigma_{ij}$$

Uit de formule is af te leiden dat als σ_{ij} groter is, dat de σ_P^2 ook groter wordt. Om dit effect te compenseren moeten de bijdragen kleiner zijn. Op deze manier kan de theorie ook wiskundig worden afgeleid. De risk parity theorie vormt een goede basis voor risico-minimalisatie.

Deze theorie kijkt dus naar de variantie en covariantie. Een manier is om alleen te kijken naar de variantie; deze manier wordt naïef genoemd (XXX), omdat het uitsluitend de variantie meet, maar niet de covariantie. De kernregel waarbij $RC_1 = RC_2 = \dots = RC_n$ wordt geprobeerd te bereiken, alleen het zal nooit aan deze voorwaarde voldoen. Daarvoor moet ook de covariantie meegerekend worden, omdat assets met elkaar samenhangen. Voor deze naïeve variant geldt:

$$w_i = \frac{\frac{1}{\sigma_i}}{\sum_{j=1}^n \frac{1}{\sigma_j}}$$

De inverse van alle varianties binnen het portefeuille wordt bij elkaar opgeteld en wordt in de noemer gezet van de breuk. Bovenaan in de teller komt de volatiliteit van asset X_i . Een extreme variant hiervan, waarbij grote volatiliteiten meer worden afgekeurd, wordt genoteerd in de volgende

vorm:

$$w_i = \frac{\frac{1}{\sigma_i^2}}{\sum_{j=1}^n \frac{1}{\sigma_j^2}}$$

Zoals duidelijk te zien is door de tweede macht, worden aan grote varianties nog minder bijdrage geleverd en aan kleine varianties juist meer. In onzekere tijden zou dit model aannemelijker zijn, omdat een lage totale variantie dan gewild is. Dan is de kans op extreme drawdown kleiner.

Een aannemelijker model is een waarbij ook de covariantie in rekening wordt genomen. De voorwaarde voor risk parity met $RC_1 = RC_2 = \dots = RC_n$ wordt dan vaker bereikt dan bij vorige varianten. Er wordt een volledige covariantiematrix Σ meegenomen. Deze covariantiematrix komt overeen met die van hoofdstuk 2.3. Voor de gewichten geldt: $w = (w_1, \dots, w_n)^T$ en voor de portefeuille variantie: $\sigma_P = \sqrt{w^T \Sigma w}$. Voor de marginale risicobijdrage van asset i geldt (XXX):

$$MCR_i = \frac{\partial \sigma_P}{\partial w_i} = \frac{(\Sigma w)_i}{\sigma_P}$$

Marginale risicobijdrage is het risico dat je loopt als je een gewicht van een asset een klein beetje verhoogd. De vergelijking is een partieel differentiaalvergelijking waarbij de verandering van σ_P tegen de verandering in w_i gezet wordt. Voor de risicobijdrage van asset i geldt het volgende:

$$RC_i = w_i \cdot MCR_i = w_i \frac{(\Sigma w)_i}{\sigma_P}$$

Voor deze bijdrage van asset i zou dus een Python code gerund moeten worden, dat dit per asset berekent. Eerst moet de MCR_i worden geschat en daarna kan deze input worden ingevuld in de bovenstaande formule. w_i kan worden berekent aangezien RC_i moet voldoen aan de volgende voorwaarde van ERC voor risk parity: $RC_1 = RC_2 = \dots = RC_n$.

Deze covariantie variant voor de risk parity voldoet aan de ERC voorwaarde en sluit aan bij het doeleinde van het profielwerkstuk. Het doel is namelijk om de totale variantie van een portefeuille zo laag mogelijk te houden en door de covariantie mee te nemen zal dit doel beter bereikt worden dan wanneer alleen de variantie wordt meegenomen.

2.4 Tijdsvariërende volatiliteitsmodellen (ARCH/GARCH)

Het portfolio wordt dus verdeeld op basis van risk parity met de covariantie en variantie. Door deze volatiliteit te voorspellen - voor bijvoorbeeld de volgende maand - kan het risico voor het portefeuille van de volgende maand worden geminimaliseerd. Dat gaat mee in de doelstelling van het onderzoek. Een manier om deze volatiliteit te voorspellen is aan de hand van ARCH- en GARCH-modellen. Dit onderdeel van het profielwerkstuk focust zich op deze modellen.

In veel klassieke modellen wordt verondersteld dat het rendement van een asset een constante variantie heeft, zoals in artikelen van Markowitz. In de praktijk is dit echter niet realistisch: markten kennen periodes van relatieve rust, afgewisseld met periodes van stress en hoge volatiliteit (XXX). Dit verschijnsel staat bekend als *volatility clustering*: grote koersschokken treden vaak in clusters op, net als rustige dagen. Om dit gedrag beter te kunnen beschrijven, zijn zogenoemde ARCH- en GARCH-modellen ontwikkeld. Deze modellen vormen het fundament van dit onderzoek.

ARCH staat voor Autoregressive (hangt af van vorige periodes), Conditional (op basis van het verleden) en Heteroskedasticity (variantie is niet constant). Oftewel: een model waarin voorwaardelijke variantie door de tijd heen verandert en afhangt van historische fouten. De ARCH-modellen zijn voor het eerst geïntroduceerd door Robert F. Engle in een artikel uit 1982 (XXX), oorspronkelijk om de variantie van inflatie in het VK te modelleren. Daarna zijn allerlei uitbreidingen ontwikkeld, waarvan GARCH (XXX) de bekendste is. Bij het ARCH-model beschouwen we een dagelijkse reeks van rendementen r_t , die als volgt kunnen worden geformuleerd:

$$r_t = \mu + \varepsilon_t,$$

waarbij μ het (eventueel kleine) gemiddelde rendement weergeeft en ε_t de onverwachte schok op tijdstip t is. In plaats van een constante variantie aan te nemen, veronderstelt een GARCH-model

dat de variantie van ε_t in de tijd varieert. Het gemiddelde van deze variantie is nul en het heeft een standaard uitwijking binnen σ_t^2 . De foutterm is dus stochastisch verdeeld en niet deterministisch:

$$\varepsilon_t \sim (0, \sigma_t^2),$$

Hierbij is σ_t^2 dus de voorwaardelijke variantie is, gegeven alle beschikbare informatie tot en met tijdstip $t - 1$. Er wordt dus alleen gekeken naar tijdstip $t - 1$ om een besluit mee te nemen. Het meest gebruikte model in de praktijk is het GARCH(1,1)-model. Dit GARCH-model specificeert de voorwaardelijke variantie met $t - 1$ als

$$\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2,$$

met de volgende voorwaarden: $\omega > 0$, $\alpha \geq 0$, $\beta \geq 0$ en $\alpha + \beta < 1$. Hierbij is ω het langetermijngemiddelde niveau van de volatiliteit; α meet hoe sterk de volatiliteit reageert op nieuwe schokken in het rendement (de impact van ε_{t-1}^2); β meet hoe persistent de volatiliteit is: hoe groot is de invloed van de vorige volatiliteit σ_{t-1}^2 . Als $\alpha + \beta$ dicht bij 1 ligt, dan keert de volatiliteit slechts langzaam terug naar haar langetermijngemiddelde na een schok. Dit sluit goed aan bij empirische observaties op financiële markten, waar perioden van hoge volatiliteit vaak langer aanhouden (XXX). Door de rest van de gegevens in de vergelijking te verwerken komt met op de volatiliteit op tijdstip t . Dat is dan een schatting van de volatiliteit. Door risk parity toe te passen komt met op een portefeuille waarvoor elke asset X_i een ERC geldt.

Om duidelijk te maken hoe een GARCH(1,1)-model in de praktijk werkt, nemen we een eenvoudig voorbeeld dat lijkt op de dagelijkse volatiliteit van een brede aandelenindex, zoals de S&P 500 of de MSCI World-index. Stel dat de parameters als volgt zijn:

$$\omega = 0,000002, \quad \alpha = 0,08, \quad \beta = 0,90.$$

Hieruit volgt dat

$$\alpha + \beta = 0,98,$$

wat betekent dat de volatiliteit zeer persistent is. Dit is een typische uitkomst in empirische studies van aandelenrendementen. Veronderstel bovendien dat de markt zich in een rustige periode bevindt, met een dagelijkse standaarddeviatie (volatiliteit) van ongeveer 1%. Dat betekent dat

$$\sigma_{t-1} = 1\% \Rightarrow \sigma_{t-1}^2 = 0,01^2 = 0,0001.$$

Stel dat de onverwachte schok in het rendement op dag $t - 1$ klein was, bijvoorbeeld 0,5%, dus $\varepsilon_{t-1} = 0,5\% = 0,005$. De volatiliteit op tijdstip t volgt dan uit het GARCH(1,1)-model:

$$\sigma_t^2 = \omega + \alpha \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2 = 0,000002 + 0,000002 + 0,00009 = 0,000094.$$

De bijbehorende standaarddeviatie of volatiliteit is dan:

$$\sigma_t = \sqrt{0,000094} \approx 0,97\%.$$

Het antwoord verschilt dus niet extreem veel vergeleken met de dag ervoor. Dat is logisch, want de parameter α is klein en β is groot. De volatiliteit van tijdstip t hoort dus niet erg te verschillen ten op zichte van $t - 1$. In het model kan een Python code worden opgesteld dat deze berekening voor elke maand doet en vervolgens op basis van risk parity w_i herverdeeld.

2.4.1 Een uitbreiding: het EGARCH-model

In financiële markten zien we vaak het *leverage effect*: negatieve rendementen resulteren vaak in een grotere variantie en positieve rendementen zorgen voor een lagere variantie (XXX). GARCH-modellen kunnen dat asymmetrische gedrag niet goed vangen. GARCH reageert namelijk symmetrisch op schokken; positieve en negatieve schokken van dezelfde grootte hebben hetzelfde effect op de volatiliteit. Zoals bekend, kan het standaard GARCH (1,1) model op de volgende manier worden geformuleerd:

$$\sigma_t^2 = w + \alpha \varepsilon_{t-1}^2 + \beta \sigma_{t-1}^2$$

Hieruit komt deze theorie van symmetrie duidelijk naar voren:

$$\sigma_t^2(+5\%) = \sigma_t^2(-5\%)$$

Een positieve schok heeft een even groot effect als een negatieve schok. Het enige wat duidelijk naar voren komt, is hoe groot de schok is; niet welke invloed een negatieve schok ten op zichte van positieve schok heeft. In financiële markten is het echter wel belangrijk om daar nog onderscheid tussen te maken.

Een EGARCH model neemt deze asymmetrie echter wel mee (XXX). EGARCH staat voor exponentiële GARCH. Het is een variant van het GARCH-model. In plaats van σ_t^2 zelf modelleer je de logaritme van de variantie:

$$\log(\sigma_t^2).$$

Hierdoor kan de logaritme elke waarde aannemen, maar door het exponent zal σ_t^2 een positieve waarde aannemen. Er hoeven geen waarden als $\alpha \geq 0$ te worden afgedwongen. EGARCH houdt ook rekening met het feit dat negatieve schokken meer effect hebben op de volatiliteit dan positieve. Voor de genormliserde schok (XXX) geldt:

$$z_t = \frac{\varepsilon_t}{\sigma_t}$$

Dat is de fout gedeeld door de huidige volatiliteit. Als $z_t \sim N(0,1)$, heeft hij gemiddelde 0 en variantie 1. Een veelgebruikte vorm van EGARCH (1,1) is:

$$\log(\sigma_t^2) = w + \beta \log(\sigma_{t-1}^2) + \alpha(|z_{t-1}| - E|z_{t-1}|) + \gamma z_{t-1}.$$

Het logaritme heeft natuurlijk een bereik van $[0, \infty]$ en daardoor is σ_t^2 altijd positief, zonder beperkingen zoals $\alpha, \beta \geq 0$. De term $\alpha(|z_{t-1}| - E|z_{t-1}|)$ modelleert hoe de grootte van de schok de volatiliteit beïnvloedt. Als $|z_{t-1}|$ groter dan gemiddeld is, dan neemt $\log(\sigma_t^2)$ toe. Dan neemt de volatiliteit dus ook toe. $E|z_{t-1}|$ is een constante; over het algemeen is dat $\sqrt{2/\pi}$ bij een normaal verdeelde z_t . De formule kan in dit geval eenvoudiger worden opgeschreven:

$$\log(\sigma_t^2) = w + \beta \log(\sigma_{t-1}^2) + \alpha(|z_{t-1}| - \sqrt{2/\pi}) + \gamma z_{t-1}.$$

γz_{t-1} zorgt voor de asymmetrie in het model. Als $\gamma < 0$, dan hebben negatieve schokken een groter effect op de volatiliteit dan positieve schokken van dezelfde grootte; hierbij komt dus het leverage effect kijken. Positieve schokken hebben dan een kleiner effect op de volatiliteit dan negatieve. Tot slot komt ook nog de $\beta \log(\sigma_{t-1}^2)$ in de formule voor. β bepaalt hoeveel de volatiliteit afhangt van de vorige schommelingen. Hoe dichter $|\beta|$ bij 1 komt, hoe langzamer de volatiliteit terugkeert naar normaal niveau. Er gaat dan tijd overheen om weer terug te keren naar de oorspronkelijke volatiliteit.

Op deze manier zorgen de extra parameters voor een nuance in het normale GARCH(1,1)-model. Juist voor ons model is een geïntegreerd EGARCH-model beter dan een GARCH-model, omdat juist de nuance in de gevolgen van schokken zo belangrijk is. Door deze vergelijking met historische data in een Python code te verwerken kan de volatiliteit worden voorspeld. Schatting gebeurt meestal via maximale likelihood; software doet dit proces, maar het is zwaarder dan een simpele AR-schatting. Over de maximale likelihood volgt in de volgende paragraaf meer.

2.4.2 Schatting via maximale likelihood

Voor ons model kan een GARCH- of -EGARCH model worden gebruikt. In allebei de modellen zitten parameters die zo goed mogelijk geschat moeten worden. Deze schattingen worden gedaan via software (Python), maar toch is het belangrijk om de theorie achter deze berekeningen te begrijpen. Een MLE is een *Maximum Likelihood Estimator*. Een MLE is een parameterwaarde die de geobserveerde data het meest waarschijnlijk maakt, gegeven een statistisch model (XXX). Het schatten van een EGARCH-model wordt gedaan met (quasi-)maximum likelihood (XXX). In deze paragraaf wordt de theorie kort toegelicht. Een diepe uitleg over deze MLE bij EGARCH's is onnodig, want Python kan met een econometrische bibliotheek deze parameters schatten en al het moeilijke werk doen. Toch volgt hieronder een intuïtieve uitleg over deze schatting, zodat het model goed begrepen kan worden. Want dan kunnen resultaten beter worden geïnterpreteerd.

In dit onderzoek gebruiken we een EGARCH(1,1)-model om de tijdsvariërende volatiliteit van rendementen te beschrijven. Laat $\{r_t\}_{t=1}^T$ een reeks (log-)rendementen zijn. We nemen aan dat elk rendement kan worden geschreven als een gemiddeld rendement plus een onverwachte schok:

$$r_t = \mu + \varepsilon_t,$$

waar μ een constante is en ε_t de foutterm voorstelt. Deze foutterm wordt verder opgesplitst in een grootte en een willekeurige factor:

$$\varepsilon_t = \sigma_t z_t.$$

Hier staat σ_t voor de (onbekende) volatiliteit op tijdstip t en z_t is een gestandaardiseerde willekeurige term met gemiddelde nul en variantie één. Intuïtief kun je σ_t zien als de “grootte van de schok” en z_t als de richting en relatieve sterkte. Het kenmerkende van een EGARCH-model is dat niet de variantie zelf, maar de logaritme van de variantie wordt gemodelleerd. Dit is in het vorige hoofdstuk al toegelicht. Daardoor blijft de variantie automatisch positief. In een EGARCH(1,1)-model schrijven we:

$$\ln(\sigma_t^2) = \omega + \beta \ln(\sigma_{t-1}^2) + \alpha(|z_{t-1}| - E|z_{t-1}|) + \gamma z_{t-1},$$

waar ω , β , α en γ parameters zijn die we willen schatten. Deze vergelijking zegt in woorden dat de huidige volatiliteit afhangt van (1) het niveau in de vorige periode, (2) de grootte van de vorige schok, en (3) het teken van de vorige schok. Op die manier kan het model bijvoorbeeld asymmetrisch reageren op slechte en goede nieuwsberichten (grote negatieve rendementen kunnen een ander effect hebben dan grote positieve).

Omdat z_t gestandaardiseerd is, wordt vaak aangenomen dat z_t voorwaardelijk normaal verdeeld is met gemiddelde nul en variantie één. Gegeven alle informatie tot en met tijdstip $t - 1$ (het verleden) is de verdeling van r_t dan normaal met gemiddelde μ en variantie σ_t^2 . In woorden: als we het verleden kennen, dan weten we hoe de verdeling van het volgende rendement eruitziet, en die verdeling hangt af van de parameters en de eerder waargenomen schokken.

De parameters van het EGARCH-model worden geschat met maximum likelihood (MLE). Dit betekent grofweg dat we die waarden van de parameters kiezen die het meest waarschijnlijk maken wat we in de data hebben waargenomen. Formeel wordt hiervoor een log-likelihoodfunctie opgebouwd op basis van de veronderstelde verdeling van de foutterm en de EGARCH-structuur voor de volatiliteit. Deze log-likelihood wordt vervolgens numeriek gemaximaliseerd (XXX).

In de praktijk voeren we deze stappen uit met behulp van de `arch`-bibliotheek in Python. De onderzoeker specificeert het type model (EGARCH) en de orde van het model (hier 1,1), waarna de software automatisch:

- de reeks voorwaardelijke varianties σ_t^2 construeert op basis van de EGARCH-vergelijking;
- de bijbehorende waarschijnlijkheid van de geobserveerde data berekent (de log-likelihood);
- een numerisch optimalisatie-algoritme gebruikt om de parameterwaarden te vinden die deze waarschijnlijkheid maximaal maken.

De uitkomst van dit proces zijn schattingen van de modelparameters, samen met standaardfouten en diagnostische maten die we gebruiken om de geschiktheid van het model te beoordelen. Op deze manier kan het EGARCH model worden geïntegreerd en gebruikt voor een passend resultaat.

2.5 Aanpassingen aan asset-allocatie na GARCH en risk parity: HMM's

Het fundament is bekend: de volatiliteit kan voorspeld worden met een GARCH model en risico kan geminimaliseerd worden op basis van risk parity. Elke asset X_i krijgt een bepaalde bijdrage w_i binnen het portefeuille. Toch is het belangrijk om ook het verwachte rendement in rekening te brengen, want ook het verwachte rendement is onderdeel van de sharpe ratio. Om $S \geq 1$ te krijgen moet ook hier naar gekeken worden. In dit hoofdstuk wordt gekeken naar HMM's en wordt uitgepluisd hoe deze modellen ons model kan verbeteren.

Stochastische processen zijn het tegenovergestelde van deterministische processen (XXX). Deterministisch betekent dat een uitkomst vrijwel altijd gelijk is wanneer er sprake is van een bepaalde set beginvoorwaarden. Oftewel: de situatie is niet zomaar willekeurig, maar heeft altijd een bepaald resultaat bij een specifieke oorzaak. Bij stochastiek gaat het echter over hevige onzekerheid

en willekeurige fluctuaties. Marktbewegingen fungeren op deze stochastische manier. Daar zijn verschillende redenen voor. Zo zijn data, renteaanpassingen, geopolitieke spanningen, nieuws en onverwachte gebeurtenissen per definitie onvoorspelbaar. Menselijke keuzes zijn vaak irrationeel door oncontroleerbare en onvoorspelbare emoties, wat zorgt voor willekeur in de markt. Dat is terug te lezen in de AMH uit 2.1.1. Sommige bewegingen in de economie zijn te voorspellen, maar het resultaat van deze gebeurtenissen is niet met volledige zekerheid te voorspellen.

Een markt is dus stochastisch vanwege de toevallige, onzekere variabelen die een bepaalde kansverdeling volgen. Zo is er een bepaalde kans dat een aandeel omhoog of omlaag beweegt. De toestand van een activum staat centraal, waarbij een bepaalde variabele de waarschijnlijkheid van verschillende uitkomsten weergeeft. De toestand van een variabele op tijdstip t kan worden voorgesteld aan de hand van de stochastische variabele X_t (XXX). Een stochastisch proces kan worden geformuleerd als een collectie:

$$\{X_t \mid t \in T\},$$

waarbij de elementen stochastische variabelen zijn en T een random variabele is. Binnen zo'n collectie zijn drie begrippen belangrijk:

- **Uitkomstenruimte Ω :** dit is de verzameling van alle mogelijke uitkomsten van een stochastisch proces. Als voorbeeld kan een dobbelsteen worden gebruikt. De uitkomstenruimte voor een zeszijdige dobbelsteen is $\Omega = \{1, 2, 3, 4, 5, 6\}$.
- **σ -algebra $\mathcal{F}$:** dit is de verzameling van gebeurtenissen waarvoor we een kans kunnen berekenen. Een gebeurtenis is een deelverzameling van Ω . Voor de dobbelsteen kan een gebeurtenis bijvoorbeeld zijn: “de uitkomst is hoger dan 2”. Bij stochastische processen gaat men ervan uit dat de mogelijke uitkomsten (gebeurtenissen) gedefinieerd zijn door de σ -algebra; deze liggen dus theoretisch vast.
- **Kansfunctie P :** aan elke gebeurtenis in de σ -algebra wordt een kans toegekend. Zo kan aan de gebeurtenis “bij een dobbelsteen is de uitkomst groter dan 2” de kans $P = \frac{4}{6}$ worden toegekend. Met deze notatie is het ook eenvoudig te begrijpen dat $P(\Omega) = 1$, want alle mogelijke uitkomsten samen hebben kans 100%.

Zoals hierboven beschreven, kan een stochastisch proces worden aangeduid als een collectie $\{X_t\}$ in een kansruimte $(\Omega, \mathcal{F}, P)$ met random variabelen over een set T . Dit is een formeel kader waarin drie elementen voorkomen, waarin berekeningen kunnen worden gedaan die met kansberekening te maken hebben. Daarnaast komt nog de meetbaarheid van S kijken. Dit betekent dat er sprake is van een ruimte S waarin het proces zich plaatsvindt met een σ -algebra, waardoor we gebeurtenissen in S kunnen meten. Dit is een benadering voor stochastische processen en deze theorie is toepasbaar op financiële markten.

2.5.1 Regime-switching modellen en Markov Chains

Al deze informatie kunnen we meenemen en toepassen op het volgende soort stochastische proces dat vaak wordt gebruikt bij het modelleren van financiële markten: het Markov-proces (XXX). Het Markov-proces is een specifiek stochastisch proces waarbij de toekomstige toestand alleen afhangt van de huidige toestand en niet van het verleden. Oftewel:

$$P(X_{t+1} = x_{t+1} \mid X_t = x_t, X_{t-1} = x_{t-1}, \dots, X_0 = x_0) = P(X_{t+1} = x_{t+1} \mid X_t = x_t)$$

Aan de hand van deze theorie zou de waarde van een aandeel vandaag alleen uitmaken voor de waarde van het aandeel morgen. Daarbij is het belangrijk om bepaalde overgangskansen te kunnen berekenen, bijvoorbeeld tussen de dagen. Dit wordt formeel gegeven door:

$$P(s, t; x, y) = P(X_t = y \mid X_s = x),$$

waarbij tijdstip t later is dan tijdstip s . In feite wordt berekend wat de kans is dat X_t gelijk is aan y op tijdstip t , wanneer X_s gelijk is aan x op tijdstip s . Als voorbeeld nemen we het aandeel X op tijdstip s , met een waarde van €100. Met de formule wordt berekend hoe groot de kans is dat aandeel X op tijdstip t bijvoorbeeld €110 waard is. Daarvoor is historische data nodig.

Dit proces kent veel verschillende toepassingen in de markt. De Markov-keten staat dan ook als basis voor modellering binnen financiële markten. Strikt genomen zijn marktbevingen natuurlijk niet volledig Markov, omdat de koers in werkelijkheid ook afhangt van het verleden. Toch worden Markov-processen vaak gebruikt, omdat ze veel eenvoudiger te analyseren en te simuleren zijn. Bovendien is de Markov-benadering meestal op korte tijdsintervallen redelijk accuraat, omdat dan alleen hoeft te worden gekeken naar de huidige prijs van het aandeel. Voor ons model en onze doelstelling is dit markov-proces goed toepasbaar en wordt het meegenomen in de Methodologie.

Een markov-proces gaat uit van een overgangsmatrix A . Hierbij wordt gekeken naar verschillende overgangen en hoe groot de kans is dat deze overgangen daadwerkelijk plaatsvinden. Bijvoorbeeld, een overgang van “stijgende koers” naar “dalende koers” met een kans van 80% kan worden weergegeven als $A_{12} = 0.8$. Een overgangsmatrix ziet er typisch uit als volgt:

$$A = \begin{pmatrix} P_{11} & P_{12} & \dots & P_{1n} \\ P_{21} & P_{22} & \dots & P_{2n} \\ \vdots & \vdots & \ddots & \vdots \\ P_{n1} & P_{n2} & \dots & P_{nn} \end{pmatrix},$$

waarbij P_{ij} de kans is van de overgang van toestand i naar toestand j . Bij een Markov-keten spreken we van een *stationaire fase* als de verdeling over de toestanden op lange termijn niet meer verandert. In het begin hangt de verdeling nog sterk af van de starttoestand, maar na veel stappen “vergeet” de keten waar hij begonnen is. Onder milde voorwaarden bestaat er dan een stationaire verdeling π waarvoor geldt dat

$$\pi = \pi P,$$

waarbij P de overgangsmatrix is. Dit betekent dat als de keten eenmaal volgens π verdeeld is, één extra stap van de Markov-keten de verdeling niet meer verandert. In de tijd groeit de verdeling van toestanden geleidelijk naar deze stationaire verdeling toe: dat is wat we bedoelen met het bereiken van de stationaire fase. Voor ons model kan het handig zijn om zo’n up-to-date overgangsmatrix en verdeling π te hebben om kleine aanpassingen nog te kunnen doen aan onze allocatie waar nodig is.

2.5.2 Verborgene regimes

In veel tijdreeksmodellen, bijvoorbeeld in de econometrie en finance, gaat men uit van het idee dat er verschillende *regimes* of toestanden bestaan, zoals een rustige markt met lage volatiliteit en een turbulente markt met hoge volatiliteit (XXX). Deze regimes zijn vaak niet direct observeerbaar. Men spreekt dan van *verborgene regimes*. We nemen aan dat er een verborgen toestandsvariabele S_t bestaat (bijvoorbeeld $S_t \in \{1, 2\}$ voor een rustig en een turbulent regime), terwijl we alleen een geobserveerde variabele Y_t zien, zoals een rendement of een koers.

Het verborgen toestandsproces S_t wordt meestal gemodelleerd als een Markov-keten met bepaalde overgangskansen. Dit betekent dat de kans om van het ene regime naar het andere te gaan alleen afhangt van het huidige regime, en niet van de volledige voorgeschiedenis. Gegeven de toestand S_t heeft de geobserveerde variabele Y_t vervolgens een regime-afhankelijke verdeling. Zo kan in regime 1 een lager variantieniveau en eventueel een ander verwachtingswaarde gelden dan in regime 2. Op die manier kan het model verschillende vormen van gedrag in de data opvangen, bijvoorbeeld periodes met lage volatiliteit afgewisseld met periodes met hoge volatiliteit.

Omdat de regimes verborgen zijn, weten we niet precies in welk regime we ons op een bepaald moment bevinden. Wel kunnen we, op basis van de geobserveerde gegevens $(Y_1, Y_2, \dots, Y_t)$, voor elk tijdstip een kans toekennen aan elk regime. Na een periode met relatief kleine schommelingen zal de kans op het “rustige” regime bijvoorbeeld groot zijn, terwijl na een reeks grote uitslagen de kans op het “turbulente” regime toeneemt. In de praktijk gebruikt men speciale filter- en smoother-methoden (zoals het bekende Hamilton-filter) om deze regimekansen door de tijd heen te schatten. Verborgene-regimemodellen zijn nuttig omdat ze expliciet rekening houden met structurele omslagen in gedrag. In plaats van één enkel model met vaste parameters te gebruiken, gaat men uit van een mix van toestanden met elk hun eigen dynamiek. Dit is vooral relevant in financiële markten, waar men duidelijk onderscheid kan maken tussen normale periodes en crisisperiodes. Op basis van de geschatte regimekansen kan men het risicobeheer aanpassen: in een regime met hoge geschatte volatiliteit kan men bijvoorbeeld de risicoblootstelling verlagen, terwijl in een kalm regime meer

risico genomen kan worden. Zo bieden verborgen-regimemodellen een flexibele en realistische manier om de tijdsafhankelijke structuur van risico en rendement te modelleren.

2.5.3 HMM voor financiële tijdreeksen

Een Hidden Markov Model (HMM) is een model dat heel goed past bij financiële tijdreeksen waarin verschillende regimes een rol spelen, maar waarin die regimes niet direct zichtbaar zijn (XXX). Het idee lijkt sterk op dat van verborgen regimes: er is een niet-geobserveerd proces dat aangeeft in welk toestand of regime de markt zich bevindt, en een geobserveerd proces dat we daadwerkelijk meten, zoals dagelijkse rendementen.

In een HMM gaan we uit van twee processen. Ten eerste is er de verborgen toestandsvaariabele S_t , die aangeeft in welk regime de markt op tijdstip t zit. Dit kan bijvoorbeeld een rustig regime zijn met lage volatiliteit en een onrustig regime met hoge volatiliteit. Men schrijft dan vaak $S_t \in \{1, 2, \dots, K\}$, waarbij elk getal voor een ander regime staat. Het proces (S_t) volgt een Markov-keten: de kans op het volgende regime hangt alleen af van het huidige regime. Deze overgangskansen worden samengevat in een overgangsmatrix P , zoals omschreven in hoofdstuk 2.6.1. Deze matrix geeft bondig de kans op elke overgang weer.

Ten tweede is er de geobserveerde grootheid Y_t , bijvoorbeeld het rendement van een aandelen-index op dag t . De verdeling van Y_t hangt af van het regime S_t waarin we ons op dat moment bevinden. In een rustig regime kan het rendement bijvoorbeeld worden gemodelleerd met een lagere variantie, terwijl in een turbulente fase de variantie hoger is. Op die manier krijg je voor elk regime een eigen set van parameters, zoals een gemiddelde en een variantie.

In de praktijk observeer je alleen de Y_t 's en niet de S_t 's. Het centrale probleem bij een HMM is daarom om op basis van de data schattingen te maken voor de overgangskansen tussen de regimes én voor de kans dat de markt zich op een bepaald moment in een bepaald regime bevindt. Hiervoor bestaan speciale algoritmen. Een bekend voorbeeld is het zogenoemde voorwaarts-achterwaarts-algoritme, waarmee de kansverdeling van S_t gegeven de hele reeks waarnemingen $(Y_1, \dots, Y_T)$ kan worden berekend (XXX). Voor financiële tijdreeksen is een HMM aantrekkelijk omdat het expliciet rekening houdt met verschillende marktfases, zonder dat je vooraf precies hoeft te weten wanneer zo'n fase begint of eindigt. Denk aan periodes van lage volatiliteit die worden afgewisseld met crisisperiodes waarin de volatiliteit sterk toeneemt. Met een HMM kun je op elk tijdstip een kans toekennen aan elk regime, bijvoorbeeld: "er is nu 80% kans op een rustig regime en 20% kans op een turbulent regime". Dit maakt het model geschikt voor toepassingen in risicobeheer, zoals het aanpassen van de risicoblootstelling wanneer de kans op een ongunstig regime oploopt. Op die manier biedt een HMM een flexibele manier om de structuur van financiële tijdreeksen te beschrijven, waarbij zowel het geobserveerde gedrag als de verborgen regimes een rol spelen.

Op deze manier kan een HMM gebruikt worden om in een rustig regime de σ_t variant te pakken voor risk parity in plaats van σ_t^2 voor onrustige periodes. Ook kunnen markov-processen gebruikt worden met in de vorm van een overgangsmatrix P om de kans op stijgen of dalen te berekenen. Op basis van deze informatie kunnen kleine aanpassingen worden gedaan aan de asset-allocatie. Deze allocatie was voorheen al herverdeeld op basis van risk parity met een GARCH-model.

2.6 Combinatie HMM–GARCH: regime-switching volatiliteitsmodellen

Naast het gebruik van een standaard GARCH- of EGARCH-model en een los Hidden Markov Model (HMM), combineren we in dit onderzoek beide benaderingen in model. Het idee achter een HMM–GARCH-model is dat de tijdreeks niet alleen tijdsvariërende volatiliteit kent, maar dat deze volatiliteit ook afhangt van een verborgen regime (toestand) waarin de markt zich op een bepaald moment bevindt. Elk regime heeft daarbij zijn eigen GARCH-achtige dynamiek.

We veronderstellen opnieuw een reeks rendementen $\{r_t\}_{t=1}^T$ en een bijbehorende, niet-geobserveerde toestandsvaariabele S_t die aangeeft in welk regime het proces zich op tijdstip t bevindt. Net als bij het HMM veronderstellen we dat S_t een Markov-keten vormt met overgangsmatrix

$$P = (p_{ij})_{i,j=1,\dots,K},$$

waarbij K het aantal regimes is en p_{ij} de kans voorstelt om van regime i naar regime j over te gaan. In veel toepassingen wordt gewerkt met $K = 2$, wat kan worden geïnterpreteerd als bijvoorbeeld een laag-volatiel en een hoog-volatiel regime.

In een HMM-GARCH-model wordt de verdeling van r_t niet alleen bepaald door de toestand S_t , maar heeft elk regime zijn eigen volatilitiedynamiek. In een eenvoudige specificatie schrijven we:

$$r_t = \mu_{S_t} + \varepsilon_t,$$

waarbij μ_{S_t} het gemiddelde rendement is in het actuele regime en ε_t de foutterm. De foutterm wordt weer geschreven als

$$\varepsilon_t = \sigma_{S_t,t} z_t,$$

waar $\sigma_{S_t,t}$ de voorwaardelijke standaarddeviatie is op tijdstip t , behorend bij het regime S_t , en z_t een gestandaardiseerde willekeurige term met gemiddelde nul en variantie één. Voor elk regime i wordt de volatiliteit $\sigma_{i,t}^2$ vervolgens gemodelleerd met een eigen GARCH- of EGARCH-structuur. In het geval van een regime-afhankelijk EGARCH(1,1)-model kan dit bijvoorbeeld worden geschreven als

$$\ln(\sigma_{i,t}^2) = \omega_i + \beta_i \ln(\sigma_{i,t-1}^2) + \alpha_i (|z_{t-1}| - E|z_{t-1}|) + \gamma_i z_{t-1},$$

voor $i = 1, \dots, K$. Hieruit blijkt dat elk regime zijn eigen parameterset $(\omega_i, \beta_i, \alpha_i, \gamma_i)$ heeft. Dit maakt het mogelijk dat de volatiliteit niet alleen verandert over de tijd, maar ook een andere dynamiek vertoont afhankelijk van de onderliggende marktsituatie.

Conditioneel op het regime $S_t = i$ en de informatie tot en met tijdstip $t - 1$ veronderstellen we doorgaans dat z_t gestandaardiseerd normaal verdeeld is. De rendementen hebben dan een regime-afhankelijke conditionele verdeling met gemiddelde μ_i en variantie $\sigma_{i,t}^2$. In woorden: als we zouden weten in welk regime de markt zich bevindt, dan weten we ook welke GARCH-structuur voor de volatiliteit geldt en hoe de verdeling van r_t eruitziet.

De schatting van een HMM-GARCH-model verloopt via maximum likelihood, maar is complexer dan bij een standaard GARCH-model. Omdat de toestanden S_t niet geobserveerd zijn, moet bij de constructie van de likelihood rekening worden gehouden met alle mogelijke toestandsreeksen. In de praktijk wordt hiervoor een efficiënte recursieve methode gebruikt (bijvoorbeeld het Baum-Welch-algoritme of een Hamilton-filter), die tegelijkertijd de toestandskansen en de modelparameters bijwerkt. Het resultaat is een set geschatte overgangskansen tussen regimes, samen met per regime een parameterset voor de gemiddelde structuur en de volatilitiedynamiek.

In dit onderzoek gebruiken we het HMM-GARCH-fundament om te onderzoeken of de volatiliteit van de rendementen beter beschreven kan worden door verschillende regimes met elk hun eigen GARCH-dynamiek. Dit biedt enerzijds een flexibeler model voor de volatiliteit, en anderzijds een directe interpretatie in termen van regimewisselingen, bijvoorbeeld tussen normale marktomstandigheden en periodes van financiële stress.

2.7 Integratie van GARCH-volatiliteit in de risk parity-aanpak

In dit onderzoek combineren we de inzichten uit de GARCH-modellering met een zogeheten *risk parity*-benadering voor portefeuillesamenstelling. Het basisidee van risk parity is dat niet de gewichten in termen van kapitaal leidend zijn, maar de bijdrage van elke belegging aan het totale risico van de portefeuille (XXX). In een eenvoudige variant betekent dit dat posities met een hoge volatiliteit een kleiner gewicht krijgen dan posities met een lage volatiliteit, zodat de risico's meer gelijkmatig over de portefeuille worden verdeeld.

In de klassieke benadering wordt vaak uitgegaan van historische volatiliteit, bijvoorbeeld op basis van een voortschrijdende standaarddeviatie van rendementen. Dit impliceert dat de geschatte risico's relatief traag reageren op veranderingen in de markt. Door in plaats daarvan een GARCH-model te gebruiken, maken we de risicometing expliciet tijdsafhankelijk en laten we de volatiliteit dynamisch meebewegen met recente informatie. Een GARCH- of EGARCH-model levert voor elk tijdstip t een schatting van de voorwaardelijke variantie $\sigma_{i,t}^2$ voor activum i . De bijbehorende voorwaardelijke standaarddeviatie $\sigma_{i,t}$ kunnen we vervolgens gebruiken als input in de risk parity-regel.

In een eenvoudige setting met N activa, waarbij geen rekening wordt gehouden met correlaties tussen activa, kan een *naïeve* risk parity-gewichtsvorming worden gedefinieerd als

$$w_{i,t} \propto \frac{1}{\hat{\sigma}_{i,t}},$$

waar $w_{i,t}$ het gewicht van activum i in de portefeuille op tijdstip t voorstelt en $\hat{\sigma}_{i,t}$ de uit het (E)GARCH-model geschatte voorwaardelijke standaarddeviatie is. De gewichten worden vervolgens geschaald zodat ze sommeren tot één:

$$w_{i,t} = \frac{\frac{1}{\hat{\sigma}_{i,t}}}{\sum_{j=1}^N \frac{1}{\hat{\sigma}_{j,t}}}.$$

In woorden: activa met een hoge geschatte volatiliteit krijgen een relatief klein gewicht, terwijl minder volatiele activa juist een groter gewicht krijgen (XXX). Doordat de volatiliteitsschattingen uit het GARCH-model van dag tot dag (of periode tot periode) kunnen veranderen, past de portefeuille zich continu aan de marktomstandigheden aan.

Het gebruik van GARCH-volatiliteiten binnen een risk parity portefeuille heeft twee belangrijke voordelen. Ten eerste houdt de gewichtsverdeling expliciet rekening met het feit dat volatiliteit in de tijd clustert en niet constant is. Periodes met verhoogde onzekerheid leiden daardoor automatisch tot lagere gewichten voor de betreffende activa. Ten tweede kan asymmetrische informatie uit een model als EGARCH (bijvoorbeeld een sterkere reactie op negatieve rendementsschokken dan op positieve) indirect worden meegenomen in de risicometing: een plotselinge negatieve schok vergroot de geschatte volatiliteit, wat zich vertaalt in een neerwaartse bijstelling van het gewicht.

In de praktische implementatie schatten we voor elk activum afzonderlijk een (E)GARCH-model op de historische rendementen (XXX). Na schatting gebruiken we de uit het model voortkomende conditionele volatiliteitsschattingen $\hat{\sigma}_{i,t}$ om de portefeuillegewichten volgens bovenstaande risk parity-regel te bepalen. De vergelijking met een traditionele, op historische standaarddeviaties gebaseerde aanpak maakt het mogelijk om te onderzoeken of een dynamische, door GARCH gedreven risicometing leidt tot stabielere risicobijdragen en mogelijk gunstigere risico-rendementsverhoudingen.

2.7.1 Aanpassing van risk parity-gewichten per regime

Naast het gebruik van GARCH-modellen voor de schatting van tijdsafhankelijke volatiliteit kan een Markov-proces worden ingezet om inzicht te krijgen in de verwachte richting van een activum. Het uitgangspunt hiervan is dat rendementen niet volledig willekeurig door de tijd bewegen, maar dat de richting van de verandering mogelijk afhangt van recente marktdynamiek. Door de rendementen in eenvoudige categorieën op te delen (bijvoorbeeld “stijging”, “daling” of eventueel “neutraal”), kan een Markov-transitiemodel worden geschat dat laat zien hoe waarschijnlijk het is dat de markt van de ene categorie naar de andere overgaat.

In de eenvoudigste vorm worden de rendementen geclassificeerd als positief of negatief, waarna voor elk tijdstip wordt bepaald of de asset in een “stijgtoestand” of “daaltoestand” zit. Op basis hiervan kan een 2×2 -overgangsmatrix worden geschat:

$$P = \begin{pmatrix} p_{\text{up,up}} & p_{\text{up,down}} \\ p_{\text{down,up}} & p_{\text{down,down}} \end{pmatrix}.$$

Hier geeft $p_{\text{up,up}}$ bijvoorbeeld de kans aan dat een activum, na een positieve rendementssignaal in periode $t-1$, ook in periode t positief beweegt. Een hoge waarde van $p_{\text{up,up}}$ duidt op positieve momentum- of trendpatronen, terwijl een hoge $p_{\text{down,down}}$ wijst op neerwaartse persistentie.

Wanneer de overgangsmatrix eenmaal is geschat, kunnen we voor elk tijdstip de verwachte kans op een stijging of daling bepalen, conditioneel op de meest recente categorie waarin het activum zich bevindt. Deze informatie kan worden gebruikt als aanvullende signalering bovenop de volatiliteitsmeting. In plaats van uitsluitend te heralloceren op basis van risico (bijvoorbeeld via de risk parity-regel), kan de portefeuille een lichte aanpassing krijgen op basis van de geschatte stijg- of daalkansen.

Een mogelijke manier om dit te implementeren is als volgt. Stel dat $\pi_{i,t}^{\text{up}}$ de kans voorstelt dat activum i in periode t een positief rendement zal laten zien, zoals afgeleid uit het Markov-proces. Laat $\pi_{i,t}^{\text{down}}$ de neerwaartse kans zijn. De risk parity-gewichten worden eerst bepaald op basis van de geschatte GARCH-volatiliteit:

$$w_{i,t}^{\text{RP}} = \frac{1/\hat{\sigma}_{i,t}}{\sum_{j=1}^N 1/\hat{\sigma}_{j,t}}.$$

Vervolgens kan een kleine, richtinggevoelige correctie worden toegepast. Een eenvoudige vorm hiervan is:

$$w_{i,t} = w_{i,t}^{\text{RP}} \left(1 + \lambda (\pi_{i,t}^{\text{up}} - \pi_{i,t}^{\text{down}}) \right),$$

waar λ een parameter is die bepaalt hoe sterk het richtingseffect in de gewichten wordt meegenomen. De term $\pi_{i,t}^{\text{up}} - \pi_{i,t}^{\text{down}}$ vertegenwoordigt de netto verwachte richting. Een activum dat volgens het Markov-model een grotere kans op stijging heeft, krijgt dan een iets hoger gewicht, terwijl een activum met een overwegend negatieve verwachting iets wordt verlaagd. Tenslotte worden de gewichten opnieuw geschaald zodat ze optellen tot één.

Deze aanpak zorgt ervoor dat de allocatie primair wordt gedreven door het risico (via GARCH), maar daarnaast subtiel rekening houdt met eenvoudige momentum- of trendpatronen die zichtbaar worden via het Markov-proces. Hiermee wordt de risicogebaseerde portefeuillestrategie uitgebreid met een lichte voorspellende component, zonder dat er een complexe voorspellings- of timingstrategie hoeft te worden ingevoerd.

3 Methodologie

Dit is het deel van het onderzoek waar het om draait: het opstellen van het model. Het doel van het model is om het risico zo klein mogelijk te maken tegen een gunstig rendement. Het is daarbij niet heel erg als er rendement wordt ‘ingeleverd’ voor een kleiner risico. Het doel is dan ook dat de backtesting van het model op historische data leidt tot een plot met weinig variantie en een stijgende lijn. Een sharpe ratio van $S \geq 0.7$ wordt daarbij als doel beschouwd. De optimale situatie zou zijn om een GARCH-HMM-risk-parity model te maken dat dit risico binnen het portefeuille minimaliseert. Op die manier zal de probleemstelling worden behandeld.

3.1 Onderzoekopzet

In dit hoofdstuk wordt behandeld hoe de hoofdvraag beantwoord wordt. De probleemstelling is dat markten dynamisch zijn, volatiliteit en regimes wisselen, en klassieke methoden houden daar weinig rekening mee. De onderzoeksvraag was als volgt geformuleerd: ***In hoeverre leidt een portefeuillemodel dat gebruikmaakt van GARCH-volatiliteitsvoorspellingen in combinatie met een HMM-regimefilter tot een hogere risico-gewogen performance dan eenvoudige benchmarkstrategieën, gegeven realistische beperkingen op leverage, posities en transactiekosten?*** Het doel is om een dynamisch (E)GARCH-HMM-risk-parity model te bouwen en te kijken of het risico gewogen rendement beter is dan simpele benchmarks.

3.1.1 Onderzoeksaanpak: kwantitatieve backtesting

Dit onderzoek betreft een kwantitatieve en empirische analyse, en is daarbij voornamelijk gefocust op de backtesting. Daarbij wordt gekeken hoe het model ten opzichte van andere modellen in het verleden zou hebben gepresteerd. Er wordt vergeleken met drie andere benchmarks: buy-and-hold (1), gelijkgewogen portfolio (2), simpele $1/\sigma$ parity (3). Buy and hold is een beleggingsstrategie waarbij een belegger effecten koopt, zoals aandelen of obligaties, en deze voor een lange periode aanhoudt, ongeacht de fluctuaties op de korte termijn. Het belangrijkste doel van deze strategie is het genereren van vermogensgroei door te profiteren van de verwachte langetermijngroei van de geselecteerde activa, in plaats van te proberen te timen wanneer de markt stijgt of daalt.

3.1.2 Dataselectie (markten, assets, periode, frequentie)

Voor het onderzoek is het belangrijk om concreet te omschrijven welke data gebruikt wordt en hoe deze gebruikt wordt. Voor de data van de assets uit het onderzoek gelden de volgende regels:

1. De assets die voor het kwantitatieve onderzoek zijn meegenomen zijn: S&P 500, Nasdaq 100, goud en staatsobligaties (lang/kort). Deze selectie zorgt voor een mix van risicovolle en defensieve beleggingscategorieën, zodat het model getest kan worden in zowel aandelen- als rente- en grondstoffenmarkten. Indien nodig kunnen later nog aanvullende benchmarks of vergelijkbare indices worden toegevoegd.

2. De historische data wordt via Python uit de **yfinance**-library (Yahoo Finance) gehaald (XXX). Dit zorgt ervoor dat de data op een gestandaardiseerde manier wordt ingelezen, en dat er eenvoudig meerdere tijdreeksen tegelijk kunnen worden opgehaald. Bovendien maakt dit het gehele proces reproduceerbaar: andere onderzoekers kunnen in principe dezelfde code gebruiken om dezelfde datasets opnieuw te downloaden.
3. De historische data start minimaal vanaf 1970; indien het mogelijk is om deze set eerder te starten, wordt dat gedaan. Op deze manier wordt een zo lang mogelijke tijdshorizon gebruikt, zodat het model getest kan worden over meerdere marktregimes (bijvoorbeeld inflatoire periodes, crises en langdurige bull-markten). Wanneer een specifieke asset pas later beschikbaar is, wordt de analyse voor die asset uitgevoerd vanaf de vroegst beschikbare datum.
4. Er is sprake van dagelijkse data. Dagelijkse observaties maken het mogelijk om volatiliteit en regimewisselingen op relatief korte termijn te detecteren, wat belangrijk is voor GARCH- en HMM-modellen. Tegelijkertijd blijft de dataset beheersbaar, in tegenstelling tot bijvoorbeeld intraday-data, die voor een profielwerkstuk onnodig complex zou zijn.
5. Bewerkingen worden gedaan door de prijzen naar (log-)rendementen om te zetten. Log-rendementen hebben wiskundige voordelen: ze zijn beter optelbaar over de tijd en sluiten aan bij veel theoretische modellen in de financiële econometrie.

3.1.3 Afhankelijke variabelen in het model

In dit onderdeel wordt beschreven wat er precies wordt gemodelleerd en welke variabelen en modellen daarin een rol spelen. Er wordt onderscheid gemaakt tussen afhankelijke variabelen, risico-variabelen en de toegepaste modellen. De belangrijkste grootheden die in dit onderzoek worden gemodelleerd zijn de rendementen van de individuele assets en het totale portefeuillerendement. Voor elke asset i op tijdstip t wordt het (log-)rendement genoteerd als

$$r_{i,t},$$

waarbij $i = 1, \dots, N$ het aantal assets voorstelt en $t = 1, \dots, T$ de tijd voorstelt. Op basis van de portefeuillengewichten $w_{i,t}$ wordt het totale portefeuillerendement op tijdstip t geschreven als

$$r_{p,t} = \sum_{i=1}^N w_{i,t} r_{i,t}.$$

Dit portefeuillerendement $r_{p,t}$ vormt de basis voor de evaluatie van het model (bijvoorbeeld via Sharpe-ratio en drawdown). Naast de rendementen zelf spelen verschillende risicogerelateerde variabelen een centrale rol:

- De voorwaardelijke volatiliteit per asset, genoteerd als $\hat{\sigma}_{i,t}$, wordt geschat met behulp van een EGARCH-model. Deze grootheid geeft weer hoe volatiel asset i verwacht wordt te zijn op tijdstip t , gegeven de beschikbare historische informatie.
- De regime-indicator S_t geeft aan in welk marktregime de economie zich op tijdstip t bevindt. In dit onderzoek wordt aangenomen dat er twee regimes zijn, bijvoorbeeld een rustig (laag-volatiel) regime en een stress- of crisisregime. Deze regimes worden gemodelleerd met een Hidden Markov Model (HMM).
- De UP/DOWN-kansen per asset worden verkregen uit een eenvoudige Markov-keten op de voortekens van de rendementen. Hierbij wordt de kans geschat dat een positief rendement (UP) volgt op een eerder positief rendement, dan wel dat een negatief rendement (DOWN) volgt op een negatief rendement. Deze kansen worden later gebruikt als klein richtinggevend signaal in de gewichtsaanpassing.

Het onderzoek maakt gebruik van een combinatie van modellen, die elk een specifiek onderdeel van het gedrag van de financiële tijdreeksen beschrijven:

- Voor elke afzonderlijke asset wordt een EGARCH(1,1)-model gespecificeerd op de tijdreeks van rendementen $r_{i,t}$. Dit model levert een tijdsvariërende inschatting van de volatiliteit $\hat{\sigma}_{i,t}$, waarbij onder meer rekening wordt gehouden met asymmetrische reacties op positieve en negatieve schokken (leverage effect).
- Een Hidden Markov Model met twee regimes (laag- en hoogvolatiel) wordt toegepast op een gekozen referentiereeks (bijvoorbeeld de S&P 500-index of het portefeuillerendement). Het HMM levert door de tijd heen de kansverdeling over de regimes, zodat op elk tijdstip een inschatting gemaakt kan worden van de kans op een rustig versus een stressregime.
- De voortekens van de rendementen (UP voor positief, DOWN voor negatief) worden gebruikt om een eenvoudige Markov-keten te schatten. Dit resulteert in een overgangsmatrix met kansen zoals $P(\text{UP} \rightarrow \text{UP})$ en $P(\text{DOWN} \rightarrow \text{DOWN})$. Deze informatie wordt gebruikt als een extra, lichte richtingindicator naast het puur risicogedreven onderdeel.
- De basisgewichten in de portefeuille worden bepaald volgens een risk parity-logica, waarbij assets met een hogere geschatte volatiliteit een lager gewicht krijgen. In de eenvoudigste vorm geldt dat de risk parity-gewichten evenredig zijn aan de inverse van de geschatte volatiliteit:

$$w_{i,t}^{\text{RP}} \propto \frac{1}{\hat{\sigma}_{i,t}}.$$

Na normalisatie zodat $\sum_{i=1}^N w_{i,t}^{\text{RP}} = 1$ ontstaat een portefeuille waarin de risicobijdragen meer in balans zijn. Vervolgens wordt een kleine correctie op deze basisgewichten aangebracht op basis van de UP/DOWN-kansen uit het Markov-model (bijvoorbeeld door $w_{i,t}^{\text{RP}}$ licht te verhogen als de UP-kans groter is dan de DOWN-kans). Daarna worden de gewichten opnieuw genormaliseerd.

3.1.4 Constructie van het portefeuillemodel (stappenplan)

In dit onderdeel wordt stap voor stap beschreven hoe het GARCH–HMM–risk parity-model in de tijd wordt toegepast. Het gaat om een iteratief proces, waarbij op elk tijdstip t alleen informatie wordt gebruikt die tot en met $t-1$ bekend is. Dit voorkomt het gebruik van voorkennis (look-ahead bias) en sluit aan bij hoe een belegger in de praktijk beslissingen zou nemen. Het algoritme kan in woorden als volgt worden weergegeven:

1. Gebruik historisch bekende data tot en met dag $t-1$. Voor elk nieuw beslismoment t wordt alleen de informatie gebruikt die beschikbaar is tot en met dag $t-1$. Dit omvat de historische prijzen, rendementen en eerder geschatte volatiliteiten en regimes.
2. Voor elke asset i wordt op basis van de rendementen tot en met $t-1$ een EGARCH(1,1)-model geschat of geüpdatet. Dit levert een inschatting van de voorwaardelijke volatiliteit $\hat{\sigma}_{i,t}$, die aangeeft hoe risicovol asset i wordt geacht op dag t .
3. Het Hidden Markov Model wordt toegepast op een gekozen referentiereeks (bijvoorbeeld een index of het eerdere portefeuillerendement) om de kans te schatten dat de markt zich op dag t in een rustig dan wel stressregime bevindt. Dit resulteert in een regime-indicator S_t of een kansverdeling over de regimes.
4. Op basis van de historische UP/DOWN-sequenties van de rendementen van elke asset wordt een overgangsmatrix geconstrueerd. Hieruit worden de kansen afgeleid dat een positieve (UP) of negatieve (DOWN) beweging zich voortzet. Deze kansen worden gebruikt als een lichte richtinggevende factor.
5. Uit de geschatte volatiliteiten $\hat{\sigma}_{i,t}$ worden de basisgewichten bepaald volgens een risk parity-logica. In de eenvoudigste vorm geldt dat

$$w_{i,t}^{\text{RP}} \propto \frac{1}{\hat{\sigma}_{i,t}},$$

waarna de gewichten worden geschaald zodat $\sum_{i=1}^N w_{i,t}^{\text{RP}} = 1$. Assets met een hogere volatiliteit krijgen zo automatisch een lager gewicht.

6. De basisgewichten $w_{i,t}^{\text{RP}}$ worden vervolgens licht aangepast op basis van de netto UP–DOWN-indicator per asset. Als de kans op een stijging (UP) duidelijk groter is dan de kans op een daling (DOWN), kan het gewicht van die asset iets worden verhoogd; in het omgekeerde geval wordt het gewicht verlaagd. De aanpassing blijft bewust klein, zodat de risicocomponent (volatiliteit) dominant blijft.
7. Na de richtingstilt worden de aangepaste gewichten opnieuw genormaliseerd, zodat geldt:

$$\sum_{i=1}^N w_{i,t} = 1.$$

Hiermee ontstaat de uiteindelijke gewichtsverdeling voor dag t .

8. Aan de hand van de gerealiseerde rendementen $r_{i,t}$ en de gekozen gewichten $w_{i,t}$ wordt het portefeuillerendement op dag t bepaald:

$$r_{p,t} = \sum_{i=1}^N w_{i,t} r_{i,t}.$$

9. De verandering in gewichten ten opzichte van dag $t - 1$ wordt gebruikt om de turnover te berekenen, bijvoorbeeld als

$$\text{turnover}_t = \frac{1}{2} \sum_{i=1}^N |w_{i,t} - w_{i,t-1}|.$$

Vermenigvuldiging van deze turnover met een verondersteld kostentarief per eenheid handel levert de transactiekosten op dag t op. Deze kosten worden in mindering gebracht op het bruto portefeuillerendement om een netto rendement te verkrijgen.

Door dit stappenplan gedurende de hele onderzoeksperiode herhaaldelijk toe te passen, ontstaat een tijdreeks van portefeuillerendementen die kan worden vergeleken met de gekozen benchmarkstrategieën. Hierdoor kan beoordeeld worden in hoeverre het GARCH-HMM-risk-parity model leidt tot een betere risico-gewogen performance.

4 Empirische resultaten

$$r_{t,i} = \ln \left(\frac{P_{t,i}}{P_{t-1,i}} \right)$$

$$\ln(\sigma_{t,i}^2) = \omega_i + \alpha_i (|\varepsilon_{t-1,i}| - E|\varepsilon_{t-1,i}|) + \gamma_i \varepsilon_{t-1,i} + \beta_i \ln(\sigma_{t-1,i}^2)$$

$$S_{t,i} \in \{0, 1\}$$

$$P_i(a, b) = P(S_{t,i} = b \mid S_{t-1,i} = a)$$

$$\pi_{t,i}^{\text{up}} = P_i(S_{t,i}, 1)$$

$$\pi_{t,i}^{\text{down}} = 1 - \pi_{t,i}^{\text{up}}$$

$$w_{t,i}^{\text{RP}} = \frac{\sigma_{t,i}^{-1}}{\sum_j \sigma_{t,j}^{-1}}$$

$$\Delta_{t,i} = \lambda (2\pi_{t,i}^{\text{up}} - 1)$$

$$\widehat{w}_{t,i} = \max \left(0, w_{t,i}^{\text{RP}} (1 + \Delta_{t,i}) \right)$$

$$w_{t,i}^{\text{tilt}} = \frac{\widehat{w}_{t,i}}{\sum_j \widehat{w}_{t,j}}$$

$$r_t^{\text{gross}} = \sum_i w_{t,i}^{\text{tilt}} r_{t,i}$$

$$\text{TO}_t = \frac{1}{2} \sum_i |w_{t,i}^{\text{tilt}} - w_{t-1,i}^{\text{tilt}}|$$

$$r_t^{\text{net}} = r_t^{\text{gross}} - c \cdot \text{TO}_t$$

4.1	Beschrijvende statistieken van de data
4.2	Resultaten van de GARCH-schattingen
4.3	Resultaten van de HMM / regimeclassificatie
4.4	Prestaties van de GARCH-risk parity-portefeuille
4.5	Prestaties van HMM-regime-afhankelijke risk parity
4.6	Vergelijking met Markowitz- en andere benchmarkportefeuilles
4.7	Risicoprofiel, drawdowns en stabiliteit van gewichten
5	Robuustheidsanalyses
5.1	Alternatieve GARCH- en HMM-specificaties
5.2	Verschillende herallocatiefrequenties
5.3	Impact van transactiekosten
5.4	Subperiode- en stressperiode-analyses
6	Discussie
6.1	Interpretatie van de resultaten
6.2	Relatie met Markowitz, EMH en regime-switching
6.3	Praktische implementeerbaarheid van (HMM-)GARCH-risk parity
6.4	Beperkingen van het onderzoek
7	Conclusie en aanbevelingen
7.1	Antwoord op de onderzoeksvragen
7.2	Theoretische en praktische implicaties
7.3	Aanbevelingen voor vervolgonderzoek
	Literatuurlijst
A	Extra tabellen en figuren
B	Technische details van GARCH- en HMM-schatting
C	Overzicht van gebruikte code en algoritmes