

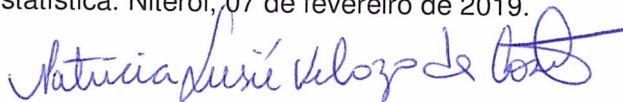
Ata da 289ª Reunião Ordinária do Departamento de Estatística

1 Aos sete dias do mês de fevereiro de dois mil e dezenove (07/02/2019) foi realizada, na sala de
2 reuniões do Instituto de Matemática e Estatística, a 289ª (ducentésima octogésima nona) reunião
3 ordinária do Departamento de Estatística (GET), que se iniciou às 10h10m horas sob a presidência da
4 professora Patrícia Lusié Velozo da Costa, chefe do GET, para deliberação sobre os seguintes itens de
5 pauta: 1) **Aprovação da ata da 288ª reunião ordinária do Departamento de Estatística;** 2)
6 **Aprovação do relatório final do projeto de iniciação à pesquisa dos professores Douglas e**
7 **Karina sob o seguinte título "Flutuações no lançamento de moedas e passeio aleatório: uma**
8 **introdução aos processos estocásticos";** 3) **Aprovação do relatório parcial e prorrogação do**
9 **projeto de pesquisa dos professores Wilson e Eduardo sob o seguinte título "Avaliando**
10 **metodologias de estimação do Valor em Risco segundo teste baseado na Regressão Quantílica";**
11 **4) Aprovação do relatório parcial do projeto de pesquisa dos professores Mariana e Wilson**
12 **intitulado "Estimação do Número de Grupos em contexto de Dados de Ranking através do**
13 **Modelo Plackett-Luce: uma Abordagem Hierárquica Bayesiana";** 5) **Aprovação do relatório final**
14 **do projeto de iniciação à docência do professor Wilson intitulado "Estudo Dirigido de Análise**
15 **Real";** 6) **Aprovação do quadro de horários para 2019-1;** 7) **Assuntos Gerais.** Estavam presentes os
16 seguintes professores: Ana Maria Lima de Farias, Eduardo Ferioli Gomes, Jony Arrais Pinto Junior,
17 José Rodrigo de Moraes, Keila Mara Cassiano, Luciane Ferreira Alcoforado, Ludmilla S. Viana
18 Jacobson, Luis Guillermo Coca Velarde, Maria Cristina Bessa Moreira, Mariana Albi de Oliveira Souza,
19 Moisés Lima de Menezes, Patrícia Lusié Velozo da Costa, Valentin Sisko e Wilson Calmon Almeida dos
20 Santos. Os seguintes professores estavam de férias: Adrian Heringer Pizzinga, Ana Beatriz Monteiro
21 Fonseca, Hugo Henrique Kegler dos Santos, Luz Amanda Melgar Santander, Márcia Marques Carvalho
22 e Núbia Karla de Oliveira Almeida. **Item 1)** A Chefia enviou dia vinte e sete de dezembro de dois mil e
23 dezoito uma versão da ata por correio eletrônico e pediu para os docentes sugerirem alterações, se
24 achassem necessário. Além disso, perguntou-se nessa Plenária se alguém queria sugerir alguma
25 alteração e ninguém se manifestou. Em seguida, submeteu-se à votação a aprovação da ata da 288ª
26 reunião ordinária, que foi aprovada por unanimidade. **Item 2)** A Chefia enviou por correio eletrônico dia
27 quatro de fevereiro deste ano o relatório final do projeto de iniciação à pesquisa sob o seguinte título
28 "Flutuações no lançamento de moedas e passeio aleatório: uma introdução aos processos estocásticos"
29 e o parecer da Comissão de Pesquisa sendo favorável ao projeto. Também disponibilizou na reunião
30 uma cópia destes documentos. O professor Wilson comentou que a Comissão de Pesquisa debateu
31 sobre o fato de que o relatório mencionava a produção de materiais, mas que parte desses não foram
32 disponibilizados [mais especificamente, scripts do R]. Como essa disponibilização não está bem
33 definida nas regras do Projeto de Pesquisa, a Comissão decidiu dar o parecer favorável
34 incondicionalmente e mencionou sobre a necessidade de discutir os assuntos (i) exigência de materiais
35 ao final dos projetos e (ii) arquivamento/acessibilidade dos materiais disponibilizados. O professor
36 Guillermo mencionou ter tido dúvidas se o projeto foi de iniciação à pesquisa ou de iniciação à docência.
37 Em seguida, submeteu-se à votação a aprovação do projeto, que foi aprovado por unanimidade. **Item 3)**
38 A Chefia enviou por email dia quatro de fevereiro deste ano o relatório parcial do projeto de pesquisa
39 sob o seguinte título "Avaliando metodologias de estimação do Valor em Risco segundo teste baseado
40 na Regressão Quantílica", o pedido de prorrogação e o parecer da Comissão de Pesquisa sendo
41 favorável ao projeto. Também disponibilizou na reunião uma cópia destes documentos. A Chefia relatou
42 que ao cadastrar os projetos de Pesquisa e Ensino no Relatório Anual de Docentes da Universidade
43 Federal Fluminense percebeu que alguns projetos aprovados pelo Departamento já atingiram o prazo
44 final de conclusão, mas os professores envolvidos nesses projetos não enviaram nem o relatório final e
45 nem um pedido de prorrogação. Sendo assim, a Chefia entrou em contato com alguns professores
46 solicitando essa regularização e aproveitou essa reunião para pedir que todos os docentes com
47 pendências desse tipo regularizem sua situação. Atendendo a esse pedido, os professores Eduardo e
48 Wilson submeteram o relatório parcial do projeto citado acima e pediram renovação do prazo. Em



seguida, submeteu-se à votação a aprovação do projeto, que foi aprovado por unanimidade. **Item 4)** A Chefia enviou por correio eletrônico no dia quatro de fevereiro deste ano o relatório parcial do projeto de pesquisa sob o seguinte título "Estimação do Número de Grupos em contexto de Dados de Ranking através do Modelo Plackett-Luce: uma Abordagem Hierárquica Bayesiana". Em seguida, submeteu-se à votação a aprovação do projeto, que foi aprovado por unanimidade. **Item 5)** A Chefia enviou por correio eletrônico no dia quatro de fevereiro deste ano o relatório final do projeto de iniciação à docência sob o seguinte título "Estudo Dirigido de Análise Real". Em seguida, submeteu-se à votação a aprovação do projeto, que foi aprovado por unanimidade. **Item 6)** A Chefia enviou por correio eletrônico no dia vinte e oito de janeiro as turmas que ainda estavam sem professor alocado e pediu que todos os docentes informassem quais as suas preferências. No dia seis de fevereiro desse ano enviou-se uma proposta de alocação dos professores nas turmas de serviço e na optativa de delineamento de experimentos. Essa proposta foi realizada com base nas seguintes considerações: Os seguintes professores já ministrarão duas disciplinas e estarão isentos da terceira (devida a planilha de pontuação): Jony, Luciane, Mariana, Maria Cristina e Nubia; Os seguintes professores não tinham escolhido qualquer disciplina anteriormente e por isso alocou-se as primeiras disciplinas escolhidas por eles, na seguinte ordem (baseada na planilha de pontuação): Ana Beatriz, Keila, Wilson, Luz Amanda, José Murilo e Adrian; Para a escolha da segunda disciplina, iniciou-se do docente com maior pontuação para o de menor pontuação, excluindo os docentes do primeiro item. Houve problemas para a alocação de alguns professores e por isso foram necessárias negociações. O candidato Marcio Watanabe, aprovado no concurso sob o Edital 165/2018, está em processo de convocação. Depois que ele tomar posse, o segundo candidato, Jaime Valdes, será convocado. Assumindo que o Marcio tomará posse antes do período letivo, finalizou-se o quadro de horários. Caso o Marcio não assuma ou algum outro problema ocorra, será necessária uma nova reunião para discutir sobre esse quadro. Pelas regras do Departamento, o Marcio deveria assumir três disciplinas, mas duas das três disciplinas que sobraram tem horário e dia da semana coincidente. Sendo assim, o professor Adrian aceitou ministrar três disciplinas. Caso o Jaime assuma no primeiro semestre desse ano, uma das disciplinas do professor Adrian passará para o Jaime. Em seguida, submeteu-se à votação a aprovação do quadro de horários de 2019-1, que foi aprovada por maioria com três abstenções. **Item 7)** Decidiu-se que a próxima reunião do GET ocorrerá dia quatorze de março (14/03), quinta-feira, às 10h30m. Em seguida, a Chefia relatou alguns problemas que estão ocorrendo para a abertura do novo concurso aprovado na 288ª Reunião Ordinária do GET. Inicialmente, a Coordenação de Pessoal Docente (CPD) questionou a formação do candidato, aprovada pelo Departamento, e sugeriu a seguinte alteração: Graduação em Estatística; Mestrado em Estatística; Doutorado em Estatística. A professora Patrícia respondeu que essas alterações reduziram muito a quantidade de possíveis candidatos. Informou que o Departamento de Estatística atualmente possui muitos professores graduados em Matemática e doutores em Estatística, por exemplo. Posteriormente, a CPD informou que o Departamento estaria impedido de abrir outro concurso para a mesma área, regime e classe do concurso sob o Edital 165/2018 pois os dois candidatos aprovados nesse concurso ainda não assumiram (processo número 23069.021597/2018-74 do Edital 165/2018). A Chefia respondeu que o Departamento possui três vagas de quarenta horas dedicação exclusiva disponíveis e que o novo concurso ocorrerá apenas em julho de 2019, quando os dois candidatos anteriores já deverão ter tomado posse. A Chefia também perguntou se poderia alterar alguma das características do novo concurso para evitar esse impedimento. A CPD então disse que não será necessário fazer qualquer alteração e que tentará acelerar o processo de posse dos dois candidatos. A professora Ana Maria informou que enviou por email o modelo da Planilha de Pontuação e pediu que todos lessem as regras antes do preenchimento dessa planilha. Qualquer erro no nome do arquivo a ser anexado, poderá produzir erro na pontuação. A Chefia estendeu o prazo máximo para o envio dessas planilhas para dia trinta e um de março (31/03). A professora Luciane solicitou que a Chefia peça ao Instituto de Matemática e Estatística a aquisição de microfone e caixa de som. Como muitas turmas de serviço possuem em torno de 60 alunos, esses itens auxiliariam muito os professores. A Chefia comunicou que excluiu duas turmas que estavam sendo abertas regularmente: A GET00177 – G1 e a GET00116 – C1. Para a exclusão dessas turmas, analisou-se o número de alunos inscritos nas disciplinas GET00177 e GET00116 (turmas B1 e C1). As turmas B1 e C1 da GET00116 excederam um pouco o número de 60 alunos em alguns semestres anteriores e por isso a Chefia mencionou que talvez seja necessária a abertura da turma C1 no período de ajuste. A professora Maria Cristina solicitou que essa divisão não fosse feita uma vez que tem havido turmas com mais de 60 alunos nos últimos

semestres e que tal divisão configuraria tratamento diferenciado entre os docentes. A Chefia informou que em 2018-2 não houve turma alguma do Departamento de Estatística com mais de 60 alunos e que sempre tenta alguma forma de contemplar as solicitações das coordenações sem prejudicar o Departamento de Estatística. Além disso, a Chefia atualmente permitirá uma turma com mais de 60 alunos caso o docente aceite e se for uma disciplina com uma única turma num dado turno. A Chefia informou que alguns professores relataram não gostarem de turmas grandes e, por isso, a Chefia informou o número de alunos inscritos nos últimos dois semestres anteriores quando enviou o quadro de horários para os professores informarem suas preferências. E que caso o número máximo de alunos seja reduzido no futuro isso implicará em um aumento no número de turmas totais e consequentemente no número de professores com três turmas. Foi discutido na reunião se haveria necessidade de cadastrar projetos de monitoria na modalidade de projeto de ensino, uma vez que alguns professores contabilizavam a carga horária do projeto apenas como orientação. O professor José Rodrigo destacou a importância do cadastramento de projetos de monitoria como "projeto de ensino" pela chefia, como vem sendo feito ao longo de algumas gestões no GET, pois desta forma seria possível contemplar aqueles projetos de monitoria que gerassem algum tipo de produto, como por exemplo, listas de exercícios, permitindo assim o registro no RAD do vínculo entre os projetos e os seus respectivos produtos de monitoria. A professora Ana Maria mencionou que antigamente era necessário esse cadastro pois a orientação estava vinculada a isso, mas que o sistema agora permite cadastrar a orientação sem ser necessário atribuir isso a um projeto. A Chefia disse que discutirá esse assunto em reuniões futuras mas pediu para os docentes que tiverem projetos de monitoria cadastrados, que tenham cuidado para não duplicar a carga horária quando forem cadastrar a orientação. Além disso, a Chefia avisou que o docente que não quiser que cadastre o projeto de monitoria em projeto de ensino basta comunicar a Chefia conforme realizado esse ano pelas docentes Ana Maria e Luz Amanda. Nada mais havendo a tratar e ninguém mais desejando fazer uso da palavra, foi encerrada a reunião às 12h09m, cuja ata vai datada e assinada por mim, Patrícia Lusié Velozo da Costa, chefe do Departamento de Estatística. Niterói, 07 de fevereiro de 2019.



Patrícia Lusié Velozo da Costa
Chefe Depto de Estatística - UFF
SIAPE 1805333

289^a REUNIÃO DO DEPARTAMENTO DE ESTATÍSTICA - UFF



Ordinária



Extraordinária

DATA: 07 / 02 / 2019

Início: 10 : 10

Término: 12 : 09

Lista de Presença

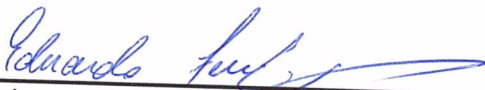
1. Adrian Heringer Pizzinga _____
2. Ana Beatriz Monteiro Fonseca _____
3. Ana Maria Lima de Farias Ana Maria Lima de Farias
4. Douglas Rodrigues Pinto _____
5. Eduardo Ferioli Gomes Eduardo Ferioli Gomes
6. Hugo Henrique Kegler dos Santos _____
7. Jessica Quintanilha Kubrusly _____
8. Jony Arrais Pinto Junior Jony Arrais Pinto Junior
9. José Murilo Ferraz Saraiva _____
10. José Rodrigo de Moraes JR
11. Keila Mara Cassiano Keila Mara Cassiano
12. Karina YurikoYaginuma _____
13. Luciane Ferreira Alcoforado Luciane F. Alcoforado
14. Ludmilla S. Viana Jacobson Ludmilla S. Viana Jacobson
15. Luis Guillermo Coca Velarde Luis Guillermo Coca Velarde
16. Luz Amanda Melgar Santander _____
17. Márcia Marques de Carvalho _____
18. Marco Aurélio dos Santos Sanfins _____
19. Mariana Albi de Oliveira Souza Mariana Albi de Oliveira Souza
20. Maria Cristina Bessa Moreira Maria Cristina B. Moreira
21. Moisés Lima de Menezes Moisés Lima de Menezes
22. Núbia Karla de Oliveira Almeida _____
23. Patrícia Lusié Vellozo da Costa Patrícia Lusié Vellozo da Costa
24. Valentin Sisko Valentin Sisko
25. Wilson Calmon Almeida dos Santos Wilson Calmon Almeida dos Santos

Assunto: Avaliação do Relatório Final associado ao Projeto de Iniciação à Pesquisa dos Professores Douglas Rodrigues Pinto (SIAPE 2283708) e Karina Yuriko Yaginuma (SIAPE 2252909).

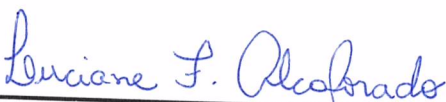
A Comissão de Pesquisa, no uso de sua atribuição, avaliou o relatório final do projeto de iniciação à pesquisa intitulado “Flutuações no lançamento de moedas e passeio aleatório: uma introdução aos processos estocásticos”, submetido pelos docentes Douglas Rodrigues Pinto (SIAPE 2283708) e Karina Yuriko Yaginuma (SIAPE 2252909), bem como o relato dos alunos envolvidos Rodolfo Hauret Spolador (matrícula UFF 117054008), Maquise de Medeiros Pinheiro (matrícula UFF 116054015) e Lucas Moura Faria e Silva (matrícula UFF 117054024). Considerando a Instrução de Serviço GET 04/2016 – e entendendo que o projeto cumpriu seu propósito, a Comissão de Pesquisa do GET considera validada a sua execução e indica o aproveitamento do trabalho realizado como atividade complementar para os alunos envolvidos.

Niterói, 20 de Dezembro de 2018

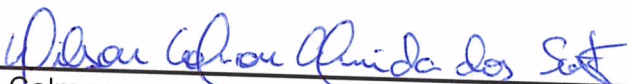
Ana Beatriz Monteiro Fonseca



Eduardo Ferioli Gomes



Luciane Alcoforado



Wilson Calmon

Relatório de Atividades 2018

Projeto de Iniciação à Pesquisa

28 de novembro de 2018

Dados do Projeto

- **Título do projeto:** Flutuações no lançamento de moedas e passeio aleatório: uma introdução aos processos estocásticos
- **Identificação do projeto:** Projeto de Iniciação à Pesquisa
- **Docente responsável 1:** Douglas Rodrigues Pinto (2283708)
- **Docente responsável 2:** Karina Yuriko Yaginuma (2252909)
- **Período global de execução:** 28/03/2018 à 14/12/2018
- **Período de avaliação:** 28/03/2018 à 14/12/2018

Participantes

Aluno 1: Rodolfo Hauret Spolador
Matrícula: 117054008
Curso: Estatística
Carga horária semanal: 12 horas

Aluno 2: Maqueise de Medeiros Pinheiro
Matrícula: 116054015
Curso: Estatística
Carga horária semanal: 12 horas

Aluno 3: Lucas Moura Faria e Silva
Matrícula: 117054024
Curso: Estatística
Carga horária semanal: 12 horas

Resumo do Plano Inicial de Pesquisa

Esse projeto foi criado com o objetivo de introduzir o aluno no estudo de processos estocásticos, utilizando o passeio aleatório nos inteiros.

O objetivo principal dos estudantes será o estudo das noções básicas e os principais resultados das flutuações do passeio aleatório simples, como quantidades de caminhos entre a origem e um ponto dado, tempos de liderança prolongada e última visita, além de estudar problemas clássicos, como o teorema da eleição (W. A. Whitworth, 1878), teste dos postos ordenados de Galton (F. Galton, 1876), testes do tipo Kolmogorov-Smirnoff (N. V. Smirnov, 1939) e o jogo teórico de lançamentos de moeda.

Paralelamente ao estudo da parte teórica do projeto, serão discutidas implementações computacionais dos modelos estudados, propiciando ao aluno o contato com a simulação de modelos estocásticos desde o início da graduação. Trabalharemos com o software R statistics.

O objeto principal de estudo é o capítulo 3 do livro *Introdução à teoria das probabilidades e suas aplicações*, de W. Feller.

Este mesmo projeto foi executado em 2017, onde foi originado um material, que poderia ser utilizado como material bibliográfico para a disciplina de Introdução aos Passeios Aleatórios. Um dos objetivos desse projeto é revisar e melhorar esse material, tornando-o acessível aos alunos do curso.

Resultados Previamente Esperados

Esperamos que os alunos compreendam os principais resultados das flutuações do passeio aleatório simples, além de serem capazes de gerar simulações estocásticas dos mesmos no programa R.

Esperamos que ao final do projeto seja gerado um material, que poderá ser utilizado como material bibliográfico para a disciplina de Introdução aos Passeios Aleatórios.

Resultados Obtidos

O projeto foi iniciado em março de 2018, e consistiu em reuniões semanais, onde realizamos discussões sobre o material da semana, e os alunos apresentavam sugestões para melhorar o entendimento da matéria em questão. Realizamos diversas discussões sobre melhoras a serem implementadas no material gerado pelo projeto de 2017. Toda discussão gerada foi redigido utilizando o software LaTeX, propiciando aos alunos o contato com essa ferramenta.

Constantemente, os alunos realizaram implementações computacionais, para simulação estocástica dos modelos discutidos, utilizando o programa R.

Os alunos mostraram um bom comprometimento com o projeto. O material gerado serviu como base para a disciplina de Introdução aos Passeios Aleatórios, ministrada no segundo semestre de 2018. Além disso, embora não previsto no projeto inicial, os mesmos prestaram auxílio aos colegas da disciplina de Introdução aos Passeios Aleatórios, apresentando aos mesmos exercícios resolvidos. Foram apresentadas também diversas sugestões de exercícios práticos e teóricos, que foram implementadas, gerando um bom material de referência para a próxima turma de Introdução aos Passeios Aleatórios, que esperamos que seja oferecido em 2020.

Seguem em anexo o material gerado pelos alunos bem como os relatórios dos mesmos.

Bibliografia

1. Feller, W. Introdução à teoria das probabilidades e suas aplicações: Parte 1 - Espaços amostrais discretos. 1a Ed. São Paulo: Edgard Blücher, 1976.
2. Ross, S. Probabilidade: Um curso moderno com aplicações. 8a Ed. Porto Alegre: Bookman, 2010.

Relatório de atividade - Projeto de Iniciação à Pesquisa
Flutuações no lançamento de moedas e Passeios Aleatórios: uma introdução aos
processos estocásticos
Aluno: Rodolfo Hauret Spolador(117054008)
Orientador 1: Douglas Rodrigues Pinto
Orientador 2: Karina Yuriko Yaginuma

Iniciado em fevereiro de 2018, o projeto teve como objetivo a introdução ao estudo de passeios aleatórios. Nossa principal referência foi o livro Introdução à teoria das probabilidades e suas aplicações de William Feller. O projeto consistia em reuniões semanais para a leitura, análise, e discussão do livro com o intuito de reescrever o material, e torna-lo mais fácil de entender.

Utilizamos o software LaTeX para redigir o material. Foi utilizado também o R, este por sua vez nos permitiu, de forma empírica, garantir a veracidade dos teoremas que demonstramos. O R também possibilitou a simulação dos passeios aleatórios, o que auxiliou no entendimento dos assuntos discutidos com os nossos orientadores.

Durante o projeto pudemos relacionar diversos resultados interessantes do estudo de passeios aleatórios com casos reais, e vimos que alguns iam contra nossa intuição. Por exemplo, em uma competição entre dois indivíduos com chances muito semelhantes de vencer, a probabilidade de que haja frequentes trocas de liderança entre eles é pequena.

Auxiliamos também na produção do material didático para as aulas de Introdução aos Passeios Aleatórios, como a criação de listas de exercícios teóricos e aplicados.

O projeto acrescentou muito, não só ao nosso conhecimento sobre passeios aleatórios, como também aprofundou nosso conhecimento do LaTeX e do R. O projeto teve seu término ao fim do segundo semestre de 2018.

Relatório de Atividades - Projeto de Iniciação à Pesquisa

Flutuações no lançamento de moedas e passeio aleatório: uma introdução aos processos estocásticos

Aluna: Maquise de Medeiros Pinheiro (116054015)

Orientador: Douglas Rodrigues Pinto

O projeto de iniciação a pesquisa científica sobre Passeios Aleatórios foi iniciado em meados de março de 2018, participando três alunos com a orientação do professor Douglas. Tivemos como base o livro “Introdução à teoria das probabilidades e suas aplicações” de William Feller e o material previamente iniciado pelos alunos participantes desse mesmo projeto no ano de 2017. Nossas reuniões foram feitas semanalmente sempre com a presença do professor Douglas.

Nosso objetivo foi revisar e concluir o material que já havia sido feito para ser usado como ferramenta auxiliar no estudo do capítulo 3 do livro escrito por Feller. Para isso, o professor nos motivava a ler uma parte do capítulo por semana e discutíamos a mesma para assim ser esclarecida qualquer dúvida. Conjuntamente, revisávamos a parte do material correspondente ao livro.

Para ser feito isso, foi necessário o uso do software LaTeX e conhecimentos de R, visto que procurávamos demonstrar de forma empírica os teoremas estudados para melhor compreensão dos mesmos.

Ao final da revisão, desenvolvemos a parte final do texto da mesma forma que os alunos anteriores, procurando deixar de forma clara o passo a passo de cada demonstração.

Paralelo ao projeto, foi criada na segunda metade de 2018 a disciplina de Passeios Aleatórios ministrada pelo professor Douglas. Essa disciplina tinha como base também o capítulo 3 do livro do Feller e o material preparado por nós. Contribuímos auxiliando o professor na elaboração de exercícios práticos e teórico e servimos como monitores aos alunos inscritos na disciplina.

Relatório de Atividades - Projeto de Iniciação à Pesquisa

Flutuações no lançamento de moedas e passeio aleatório: uma introdução aos processos estocásticos

Aluno: Lucas Moura Faria e Silva (117054024)

Orientador: Douglas Rodrigues Pinto

O projeto foi iniciado em março de 2018, junto ao começo do primeiro semestre de 2017, e tínhamos como objetivo, continuar o estudo de passeios aleatórios que já havia sido iniciado em 2017, porém com outro grupo de alunos. Utilizamos como base para nosso estudo o capítulo 3 do livro “Introdução à teoria das probabilidades e suas aplicações: Parte 1 – Espaços amostrais discretos” de William Feller e o texto que havia sido iniciado pelos alunos do mesmo projeto do ano anterior.

As reuniões com o professor Douglas eram semanais. Inicialmente, nossa principal tarefa era a de revisar o conteúdo escrito no texto. Para isso, estudávamos uma parte do livro determinada pelo professor e a discutíamos na reunião da semana. Muitas vezes era preciso usar o texto dos alunos para estudar, uma vez que a escrita do livro era por vezes obscura.

Para realizar as modificações que julgávamos necessárias no texto, de forma a melhorar seu conteúdo, foi necessário ter conhecimento no software LaTeX.

Conforme fomos avançando, começamos a fazer simulações no R com o objetivo de verificar empiricamente os resultados teóricos que obtivemos no nosso estudo. As simulações foram bastante úteis para ajudar a entender a teoria.

Quando terminamos a revisão, começamos a escrever a parte final do texto, tomando o cuidado de desenvolver de forma clara o passo a passo das demonstrações do livro, que, como já citado antes, não são simples de entender.

Foi criada no segundo semestre de 2018, paralelamente ao projeto, a disciplina de Introdução aos Passeios Aleatórios, ministrada pelo professor Douglas. O conteúdo da disciplina teve como base o livro do Feller e o nosso texto. Trabalhamos como monitores para ajudar os alunos inscritos e auxiliamos o professor na criação de exercícios.

Tarefa Aplicada 1 - Introdução aos Passeios Aleatórios

1. Para gerar no R um lançamento de moeda, um método elementar é gerar aleatoriamente um número entre 0 e 1, com o comando `runif(1)`. Se o número gerado for menor que 0.5, consideramos cara, senão, o resultado foi coroa.

Construa uma função que simule um passeio aleatório de tamanho n , onde n é um argumento da função. A função deve retornar o gráfico do passeio aleatório simulado e o valor de S_n .

2. Construa uma função que simule r passeios aleatórios de tamanho n , onde n e r devem ser argumentos da função. A função deve retornar a média dos S_n e um histograma do estado onde terminou o processo.
3. (a) Faça uma função que receba dois números inteiros, $n > 0$ e r e retorne o número total de caminhos que vão da origem até o ponto (n, r) , ou seja, $N_{n,r}$. Não se esqueça de verificar se n e r têm a mesma paridade, caso contrário a função deve retornar uma mensagem de erro.
(b) Usando a função anterior, faça uma função que retorne a probabilidade de num instante $n > 0$ o caminho estar na altura r , ou seja, $p_{n,r}$.
4. Utilizando o exercício 2, construa 1000 replicações de um passeio aleatório de 100 passos. Calcule a proporção de passeios tais que $S_{100} = 10$. Compare com o resultado obtido no exercício 3.
5. (a) Faça uma função que receba um número inteiro positivo n e retorne a probabilidade de ocorrer uma volta a origem no instante n (u_n). Lembre-se de verificar se n é um número par, caso contrário a função deve retornar uma mensagem de erro.
(b) Melhore a função anterior de forma que ela retorne também essa probabilidade aproximada pela fórmula de Stirling, se o usuário preferir.
(c) Construa uma tabela que apresenta os erros absolutos e relativos quando utilizamos a aproximação de Stirling, para diversos valores de n ($n=2,4,6, 8, 10, 50, 100, 500, 1000$).
6. Utilizando o exercício 2, construa 1000 replicações de um passeio aleatório de 100 passos. Calcule a proporção de passeios tais que $S_{100} = 0$. Compare com o resultado obtido no exercício 5).
7. Faça uma função que receba um número inteiro $n > 0$ e retorne a probabilidade de que no instante n ocorra o primeiro retorno a origem (f_n). Use as funções criadas anteriormente.
8. Implemente o teorema da lei do arco seno para últimas visitas.
9. Construa 1000 replicações de um passeio aleatório de 100 passos, e guarda o instante da última visita ao estado zero. Construa um histograma dos instantes de última visita ao zero. Faça uma tabela, comparando os resultados obtidos empiricamente com os resultados teóricos do exercício 8.

Tarefa Aplicada 2 - Introdução aos Passeios Aleatórios

1. Implemente uma função que calcula a probabilidade de haver r mudanças de sinal em $2n + 1$ unidades de tempo.
2. Implemente uma função que calcula a probabilidade do máximo de um passeio aleatório de n passos ser r .
3. Construa 10.000 replicações de um passeio aleatório de 199 passos, e guarde, para cada replicação, o número de trocas de sinal, o máximo atingido por cada passeio e o número de **lados** que o passeio ficou acima do eixo.
 - (a) Construa um histograma com o número de trocas de sinal das 10.000 replicações.
 - (b) Construa uma tabela comparando os resultados empíricos com os resultados teóricos do exercício 1, para $r=1, 2, 3, 4, 5, 10, 20, 30, 40$.
 - (c) Construa um histograma com o ponto máximo das 10.000 replicações.
 - (d) Construa uma tabela comparando os resultados empíricos com os resultados teóricos do exercício 2, para $r=0, 1, 2, 3, 4, 5, 10, 15, 20, 50, 100$.
 - (e) Construa um histograma com o número de lados que o processo ficou acima do eixo das 10.000 replicações. O que é possível observar?
4. Implemente uma função que constrói o passeio aleatório em $\mathbb{Z}^2$, com n passos. A função deve retornar o desenho que a partícula realizou em $\mathbb{Z}^2$, sem o eixo temporal, e o estado final.
5. Implemente uma função que constrói r replicações do passeio aleatório de n passos em $\mathbb{Z}^2$. A função deve retornar um histograma com a distância de S_n em relação à origem. Utilize $d(x, y) = |x| + |y|$, conhecida como a *métrica de Manhattan*.
6. Implemente uma função que constrói o passeio aleatório no grafo estrela, em função de r vizinhos, tal que, ao visitar um vizinho repetido, o processo para. A função deve retornar o número de vizinhos visitados.
 - (a) Realize 1.000 replicações do processo, e retorne um histograma com o número de vizinhos visitados, para $r=3, 5, 10, 20, 100$.
7. Implemente uma função que constrói o passeio aleatório no grafo estrela, em função de r vizinhos, tal que, ao visitar todos os vizinhos, o processo para. A função deve retornar o número passos realizados.
 - (a) Realize 1.000 replicações do processo, e retorne um histograma com o número de passos realizados, para $r=3, 5, 10, 20, 100$.
8. Implemente uma função que constrói o passeio aleatório circular, em função de r vizinhos, n passos e p , onde p é a probabilidade da partícula avançar uma unidade no sentido horário. A função deve retornar o número de vizinhos visitados.
9. Implemente uma função para o problema da ruína do jogador. Os parâmetros da função devem ser: o montante inicial i do jogador; o valor r que o jogar, ao receber, para de jogar; e a probabilidade p do jogador ganhar. A função deve retornar o gráfico da evolução temporal da fortuna do jogador, e o estado final.

10. Realize 1.000 replicações do problema da ruína do jogador, para $i = 3$, $r = 6$ e 3 valores de $p : 0.3, 0.5$ e 0.7 . construa uma tabela que mostre, para cada valor de p , a proporção de processos que terminaram em 0 ou 6.
11. Implemente uma função que constrói o *frog model*, em n passos, e probabilidade p de avançar uma casa.

Lista de Exercícios

17 de setembro de 2018

1. Em 10 lançamentos de moeda responda:
 - a) Qual o número de caminhos que saem da origem e termina em 2.
 - b) Qual o número de caminhos que saem da origem e ficam acima de -7.
 - c) Qual a probabilidade de terminar em um número ímpar.
2. Qual o número de caminhos, em 10 lançamentos, que sai de 3 e chega em 1? Quantos caminhos não tocam o eixo zero?
3. Duas plantas são submetidas a um experimento (planta A e planta B), para a criação de uma substância que acelera o crescimento. Todo o dia é observado se elas crescem, se cresceram no dia é somado +1, caso contrário é somado -1. Sabendo que elas tem probabilidade de 0.5 de crescer responda:
 - a) Em 15 dias observados, qual a probabilidade da planta A ter ficado 9 medidas positivas.
 - b) Em quantos caminhos a planta B cresceu mais que a planta A, se a planta A cresceu 6 medidas positivas, em 15 dias observados.
4. Um aluno vai fazer uma prova ,de verdadeiro e falso, com dez questões. A nota da prova é calculada da seguinte forma, cada acerto representa 1 ponto e cada erro -1 ponto. Sabendo que o aluno chutou todas as questões, então:
 - a) Qual a probabilidade do aluno ficar com nota ≥ 6 .

- b) Qual a probabilidade do aluno zerar a prova.
- c) Supondo possível ficar com nota negativa, qual a probabilidade do aluno ficar com nota não negativa dado que errou a primeira questão.
5. Demonstre que $\binom{p+q}{p} = \binom{n}{(n+r)/2}$, onde p é o número de subidas, q de descidas, n é o tamanho do caminho e $r = p-q$.
6. Quantos caminhos existem num passeio que parte da origem até o ponto $(14, 6)$ sem tocar o eixo -1 ?
7. Quantos caminhos um passeio aleatório que parte da origem até o ponto $(9, 1)$ sem tocar o eixo 2 , possui?
8. Generalize os exercícios anteriores.
9. Carlos gosta de jogos de azar, mas segue determinadas regras, ele joga 8 vezes ou para quando zera o dinheiro. Sabendo que ele tem probabilidade 0.5 de vencer cada jogo, responda:
- a) Qual a probabilidade dele jogar menos de 8 vezes.
- b) Qual a probabilidade dele terminar com saldo negativo.
10. Em 20 lançamentos de um moeda honesta, qual a probabilidade do ultimo retorno ao zero ocorrerem no instante $2n = 8$? Qual a probabilidade do último retorno ocorrer no instante $2n = 12$ dado que o primeiro retorno ocorreu no instante $2n = 6$

Lista 2 - Introdução aos Passeios Aleatórios

1. No passeio aleatório circular, com estados 1, 2, 3 e 4, e probabilidade p de avançar um estado em sentido horário, seja $X_0 = 1$.

(a) Calcule $P(X_3 = 2 | X_0 = 1)$.

(b) Calcule $P(X_4 = 1 | X_0 = 1)$.

(c) Seja $T_1 = \min \{n > 0 : X_n = 1\}$. Prove que

$$P(T_1 = n) \neq 0 \Leftrightarrow n \text{ é par.}$$

2. Supondo um passeio aleatório partindo da origem, calcule:

(a) A probabilidade desse passeio, que termina no ponto $(7, 3)$, ter um máximo igual a 5.

(b) Quantos caminhos satisfazem o item anterior.

(c) A probabilidade de um passeio de tamanho 10 ter um máximo igual à 5.

3. Dado um passeio aleatório, quantos caminhos tocam o eixo 5 pela primeira vez somente no instante 20 ?

4. Determine a probabilidade de um passeio de tamanho 12 ter 6 lados acima do eixo, dado que no instante 12, ele está no estado 0.

5. Considere um passeio aleatório em $\mathbb{Z}^2$, qual a probabilidade de que em quatro passos uma partícula termine exatamente onde começou ?

6. Suponha que exista uma partícula que segue a seguinte regra: ela pode seguir para seis caminhos diferentes que são determinados de acordo com o resultado do lançamento de um dado honesto, após ela seguir por um caminho, a partícula retorna a origem e o processo continua. Supondo um processo que acaba quando a partícula visita um caminho que já foi visitado antes, calcule:

(a) A probabilidade desse processo ter exatamente 3 passos.

(b) A probabilidade da partícula conseguir visitar os seis caminhos possíveis antes de terminar o processo.

7. Elizabeth chega em um cassino com \$500 e decide que vai jogar até duplicar seu dinheiro inicial. Suponha que ela apostará \$100 por jogada e que se ganhar terá um lucro de \$100, suponha também que a probabilidade de ganhar é igual a p , se em algum momento ela ficar sem dinheiro, pedirá emprestado para poder continuar jogando.

(a) Qual a probabilidade de Elizabeth jogar exatamente 15 rodadas?

(b) Suponha que Elizabeth pode parar de jogar a qualquer momento com probabilidade r , determine a probabilidade dela jogar por 10 rodadas e sair com a mesma quantidade de dinheiro que chegou.

(c) Supondo que ela não poderá pedir dinheiro emprestado, qual a probabilidade dela jogar 10 rodadas ?

Da COMISSÃO DE PESQUISA DO GET

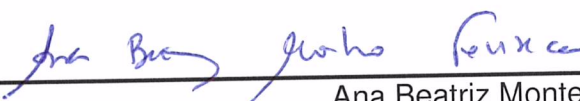
Para Chefia do Departamento de Estatística (GET)

Assunto: Avaliação do Relatório Parcial e Prorrogação do Projeto de Pesquisa dos Professores Wilson Calmon Almeida dos Santos (SIAPE: 2197625) e Eduardo Ferioli Gomes (SIAPE: 2222291).

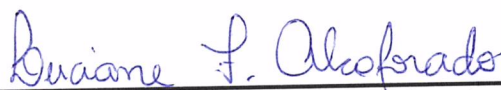
A Comissão de Pesquisa, no uso de sua atribuição, avaliou o relatório parcial do projeto de pesquisa intitulado "Avaliando metodologias de estimação do Valor em Risco segundo teste baseado na Regressão Quantílica", submetido pelos professores Wilson Calmon Almeida dos Santos (SIAPE: 2197625) e Eduardo Ferioli Gomes (SIAPE: 2222291).

Considerando a Instrução de Serviço GET – e entendendo que o projeto vem cumprindo o cronograma e seu propósito, a Comissão de Pesquisa do GET considera aprovado o relatório parcial de sua execução, bem como a prorrogação do projeto conforme solicitado.

Niterói, 19 de dezembro de 2018



Ana Beatriz Monteiro Fonseca



Luciane Alcoforado



Relatório Parcial do Projeto de Pesquisa: *Avaliando metodologias de estimação do Valor em Risco segundo teste baseado na Regressão Quantílica*

Identificação do Projeto: Projeto de Pesquisa

Nome e matrícula SIAPE dos docentes responsáveis: Wilson Calmon Almeida dos Santos (SIAPE: 2197625) e Eduardo Ferioli Gomes (SIAPE: 2222291).

Resumo do plano inicial de pesquisa: O projeto de pesquisa original tinha por objetivo comparar as principais metodologias de estimação do Valor em Risco ou VaR [*Value at Risk*] com base em testes clássicos [*backtests*] e em um teste baseado proposto mais recentemente, que é baseado na regressão quantílica. Duas formas de comparação dos métodos foram previstas: (i) a comparação empírica, em que diferentes métodos são aplicados para estimação do VaR de um conjunto de séries temporais de retornos de ativos [como índices de bolsas de valores, ações de empresas, preços de commodities, etc] e/ou (ii) a comparação, via simulação, em que diferentes métodos são aplicados a séries temporais geradas artificialmente [simuladas]. A execução do projeto deveria envolver, portanto, as tarefas: (a) levantamento [com base na literatura] das principais metodologias de estimação do VaR e de teste para validar as mesmas [*backtests*], (b) escolha dos ativos mais interessantes e/ou definição de uma estratégia de simulação; (c) implementação das metodologias de estimação do VaR e dos testes; (d) eventualmente, a elaboração de uma rotina de simulação; (e) execução da análise das séries reais e/ou simuladas e (f) resumo dos resultados.

Resultados previamente esperados: Esperávamos que, ao comparar metodologias para estimação do VaR baseando-se no teste de regressão quantílica (e outros mais tradicionais), pudessemos identificar um conjunto de melhores metodologias ou uma melhor metodologia para a estimação do VaR que sirva com referência [benchmark] para trabalhos futuros. No que diz respeito, em particular, à análise da performance com séries reais, esperávamos observar uma hierarquia das metodologias, ao menos para algumas classes de ativos [por exemplo, ao menos para índices de bolsas de valores ou taxas de câmbio].

Resumo do projeto executado e resultados efetivamente obtidos: O projeto original, enviado em Agosto de 2016, contemplava a realização das tarefas que constam no cronograma exibido na sequência e previa a conclusão do trabalho no início do ano de 2018 [o cronograma também aparece no projeto submetido ao departamento]. São 14 tarefas no plano inicial. Apesar do atraso na execução do trabalho em relação ao plano inicial, as 11 primeiras tarefas já foram cumpridas. Foi feita uma revisão da literatura, identificamos os principais métodos de estimação do VaR e os principais procedimentos de teste [*backtests*]. Tanto os métodos [6, no total], quanto os testes [4 procedimentos], foram implementados no R. De maneira similar, escolhemos os ativos [80, no total] para realizar a análise empírica. Implementamos um mecanismo de simulação e geramos 8000 séries temporais sintéticas. Boa parte do atraso na execução se deve, inclusive, à demora na geração dos dados simulados e à avaliação das séries simuladas. Optamos pela utilização de um método não usual de simulação em que os dados foram gerados a partir de distribuições condicionais estimadas para os mesmos ativos utilizados na análise empírica. Tais distribuições foram estimadas através da adoção de um método não paramétrico [estimadores de núcleo ou kernel]. Também já obtivemos todos os resultados da pesquisa. Os da análise empírica foram compilados e descritos. Os resultados do estudo de simulação estão sendo compilados e ainda precisam ser descritos.

[cronograma que consta no projeto original]

CRONOGRAMA	2016							2017												2018	
	JUN	JUL	AGO	SET	OUT	NOV	DEZ	JAN	FEV	MAR	ABR	MAI	JUN	JUL	AGO	SET	OUT	NOV	DEZ	JAN	FEV
Busca Sistemática de Artigos com Metodologias de Estimação VaR		X	X																		
Leitura de Material Básico sobre o tema do Trabalho		X	X	X																	
Leitura e Resumo das Principais Metodologias de Estimação VaR				X																	
Busca Sistemática de Artigos comparando Metodologias				X	X	X															
Leitura e Resumo das Comparações de Metodologias							X	X	X												
Programação no R do Backtesting baseado na Regressão Quantílica								X	X												
Escolha das Metodologias a Analisar e Séries para Exemplos							X	X	X												
Escrita, em formato de artigo, sobre as metodologias escolhidas										X	X										
Programação [Implementação] das metodologias escolhidas										X	X	X	X	X							
Programação [Implementação] da Análise de Simulação											X	X	X	X	X						
Programação [Implementação] da Análise Empírica												X	X	X	X	X					
Compilação dos Resultados															X	X					
Escrevendo, em formato de artigo, sobre Metodologia e Resultados																X	X	X			
Revisão do artigo																		X	X	X	

Um dos resultados do trabalho foi a execução de um projeto de iniciação científica, com apoio da FAPERJ, de um dos alunos envolvidos. Submetemos um relatório final neste ano de 2018 e o mesmo foi aprovado [parecer da FAPERJ encontra-se em anexo]. O projeto encontra-se concluído. O relatório apresentado à FAPERJ apresenta parte dos resultados empíricos obtidos.

Como um segundo e, mais importante, produto, esperamos concluir um artigo para apresentar todos os resultados da análise empírica e da análise de simulação. Associadas às tarefas acima, organizamos o artigo com as seguintes seções:

1. Introdução
2. Métodos de Estimação do VaR
3. Backtests [procedimentos para teste do VaR estimado]
4. Análise Empírica
5. Estudo de Simulação
6. Conclusão

Em relação ao artigo, boa parte já foi escrita [versão parcial também disponível em anexo]. Grosso modo, resta a tarefa de escrever os resultados da seção 5 [em particular, já descrevemos a estratégia de simulação]. Os resultados em questão já foram obtidos e estão sendo resumidos para efeitos da escrita. Além disso, será feita uma revisão do texto, formatação para envio a alguma revista.

Os arquivos referentes aos resultados da análise empírica e do estudo de simulação estão disponíveis para consulta, assim como todos os códigos de todos os programas gerados neste trabalho no site <https://sites.google.com/site/filesrgsmuff/home/var?pli=1>. Basta fazer download das pastas, clicando nos links “OutputSimulationEmpirical” e “scriptsgerados”, respectivamente.

Solicitação de extensão:

Em virtude do atraso, pedimos desculpas à comissão e ao departamento e solicitamos extensão do prazo [até Abril de 2019]. Esperamos concluir a seção 5.1 até Fevereiro de 2019 e a revisão nos últimos dois meses.

Resumo dos alunos envolvidos [transcritos dos e-mails anexados]:

Avaliação do Projeto pelo aluno Davi Saad de Almeida Lettieri*:

O período de atividades no projeto permitiu o aprendizado de conteúdos diversos. Desde a preparação e montagem de uma pesquisa em nível superior à assimilação de conteúdos específicos. A elaboração da pesquisa possibilitou o contato com conteúdos de artigos relacionados ao projeto, resultando numa familiarização com esse tipo de conteúdo; introduziu uma nova noção de análise e organização de dados; aprofundamento em tópicos específicos da estatística; utilização de uma nova linguagem de programação (R).

* O Aluno Davi Saadi Davi Saadi de Almeida Lettieri substituiu a aluna [Karolina Nascimento Melo, 214054093, discente de graduação em Estatística da UFF, no projeto – e passou a receber oficialmente a bolsa IC da FAPERJ no segundo semestre de 2017].

Avaliação do Projeto pelo aluno Johann Rodrigues de Souza Soares:

"O projeto foi e vem sendo muito importante para mim por motivos diversos, entre eles se destacam: Programação em R, no sentido de que elaboro os scripts com os procedimentos importantes para algumas tarefas realizadas na pesquisa; conhecimento acerca da literatura de VaR, que é importante pelo fato de fazer parte de uma área do conhecimento (finanças) que eu não compreendia, mas que é muito relevante para a minha carreira futura, de modo que a importância de entrar em contato com essa literatura na pesquisa é inestimável para mim; e o contato com algumas metodologias estatísticas que vão contribuir com a minha bagagem de conhecimento de forma ferramental, seja lidando com regressões quantílicas, seja na parte de estimação não paramétrica via núcleos. Além de todos esses aspectos é importante mencionar a experiência adquirida com a tarefa de pesquisar. Para quem quer seguir na área acadêmica é muito importante entrar logo cedo com metodologia científica de pesquisa, e com certeza esse projeto está me proporcionando isso."

Niterói, 14 de Dezembro de 2017.

Eduardo Ferioli Gomes

Wilson Calmon Almeida dos Santos

Chefe do Departamento



Wilson Calmon <calmonwilson@gmail.com>

Pesquisa VaR

Davi Saadi <davilettieri@gmail.com>

8 de novembro de 2017 11:06

Para: Wilson Calmon <calmonwilson@gmail.com>

Bom dia, Professor! Segue a minha avaliação, se precisar de qualquer ajuste estou disponível.

Avaliação do Projeto:

O período de atividades no projeto, permitiu o aprendizado de conteúdos diversos.

Desde a preparação e montagem de uma pesquisa em nível superior à assimilação de conteúdos específicos.

A elaboração da pesquisa possibilitou o contato com conteúdos de artigos relacionados ao projeto, resultando numa familiarização com esse tipo de conteúdo; introduziu uma nova noção de análise e organização de dados; aprofundamento

em tópicos específicos da estatística; utilização de uma nova linguagem de programação(R).

Em 7 de novembro de 2017 17:24, Wilson Calmon <calmonwilson@gmail.com> escreveu:

Caros Johann e Davi, boa tarde!

Queria pedir um grande favor a vocês. Vocês podem me enviar [por e-mail mesmo] um resumo sobre a avaliação de como o projeto os ajudou até o momento? Tenho que fazer um relatório de prestação de contas do projeto de pesquisa que está cadastrado no GET e um dos itens refere-se a tal resumo. Pode ser pequeno [1 ou 2 parágrafos].

Envio, em anexo, o relatório que pretendo submeter [após incluir vosso resumo].

Desde já agradeço enormemente!

Abraço!

Wilson Calmon Almeida dos Santos
Professor

Departamento de Estatística
Instituto de Matemática e Estatística
Universidade Federal Fluminense



Wilson Calmon <calmonwilson@gmail.com>

Script com Box plots

Johann Soares <johann.soares@hotmail.com>
Para: Wilson Calmon <calmonwilson@gmail.com>

16 de novembro de 2017 16:08

Boa tarde Wilson,

desculpe a demora. Segue o meu relato:

"O projeto foi e vem sendo muito importante para mim por motivos diversos, entre eles se destacam: Programação em R, no sentido de que elaboro os scripts com os procedimentos importantes para algumas tarefas realizadas na pesquisa; conhecimento acerca da literatura de VaR, que é importante pelo fato de fazer parte de uma área do conhecimento (finanças) que eu não compreendia, mas que é muito relevante para a minha carreira futura, de modo que a importância de entrar em contato com essa literatura na pesquisa é inestimável para mim; e o contato com algumas metodologias estatísticas que vão contribuir com a minha bagagem de conhecimento de forma ferramental, seja lidando com regressões quantílicas, seja na parte de estimação não paramétrica via núcleos.

Além de todos esses aspectos é importante mencionar a experiência adquirida com a tarefa de pesquisar. Para quem quer seguir na área acadêmica é muito importante entrar logo cedo com metodologia científica de pesquisa, e com certeza esse projeto está me proporcionando isso."

Está bom Wilson? Se não pode me falar que eu corrijo,

Grato,
Johann Soares.

De: Wilson Calmon <calmonwilson@gmail.com>

Enviado: sexta-feira, 10 de novembro de 2017 05:13

[Texto das mensagens anteriores oculto]

[Texto das mensagens anteriores oculto]

Rio de Janeiro, 14/11/2018,

Ilmo.(a) Sr. (a) Prof. (a) Dr. (a)
Wilson Calmon Almeida dos Santos

Processo nº. : E-26/ 201.731/2017 - BOLSA

Solicitante : Davi Saadi de Almeida Lettieri

Matrícula : 2017.02915.1

Assunto : Resultado da Avaliação do Relatório Final

Prezado (a) Senhor (a),

Comunicamos a V. Sa. que o Diretor Científico, com base no parecer da Assessoria Técnico-Científica desta Fundação, considerou **aprovado** o seu relatório sobre o processo acima referenciado.

Esclarecimentos adicionais (se houver):

A Assessoria Técnico-Científica da FAPERJ emitiu o seguinte parecer:

Relatório Final aprovado.

Cordialmente,

A Diretoria

An extensive comparison between the most successful methods for VaR

Calmon, Wilson[†]; Ferioli, Eduardo[†]; Lettieri, Davi^ℓ; Soares, Johann[¥]

Departamento de Estatística [†]

Departamento de Engenharia Mecânica ^ℓ

Faculdade de Economia [¥]

Universidade Federal Fluminense, Niterói

December 14, 2018

Abstract

Since 1990s, several methods for estimating Value at Risk [VaR] were proposed. In the early 2000s, the most popular VaR methods [eg, Historical Simulation, Risk Metrics, GARCH] were simplest ones. Since then, relatively more complex approaches have been adopted both in practice as in academic research, as for example, Extreme Value Theory [EVT], High Order Moments [MOM], Filtered Historic Simulation [FHS] and Conditional Autoregressive Value at Risk [CAViaR]. These four alternatives are some of the most important VaR methods nowadays. The objective of this paper is to compare the performances of EVT, MOM, FHS and CAViaR. The analysis comprises both an empirical investigation as a simulation study. Our results indicate that the EVT method outperforms its competitors. In addition, FHS and HOM are the second best methods, but they perform better in cases where EVT works. On the other hand, for assets whose VaR can not be adequately estimated by the EVT method, the CAViaR outperforms both the FHS and the HOM. Our study also indicates the need to consider other approaches, since the rejection frequencies of each method are generally greater than the level of significance.

keywords: *Value-at-Risk, Backtest, CAViaR, High Order Moments, Extreme Value Theory, Filtered Historical Simulation*

1 Introduction

Since the 1990s, several methods [VaR methods, therefore] for estimate Value at Risk [VaR] were proposed as highlighted in [Jor2006]. The two most popular

VaR methods are the normal parametric method and the historical simulation [HS] method. In both cases, it is assumed that the returns of the assets, obtained in different time instants $[\dots, R_{t-1}, R_t, R_{t+1}, \dots]$, are independent and identically distributed [i.i.d.]. However, these hypotheses were frequently questioned, inspiring other more sophisticated alternatives.

As explained in [AbaAl214]Hull (2015), a very popular alternative is the use of the GARCH model [[AbaAl214] Engle (1982) and [AbaAl214] Bollerslev (1986)] that incorporates the dependence between returns over time - especially between their second order moments. Although GARCH usually is associated with the normality hypothesis, as highlighted in [Roc216], such hypothesis can also be relaxed with the adoption of other distributions such as t-Student, GED or its asymmetric versions. More still, there are several extensions for GARCH, as the variants EGARCH, APARCH, TGARCH, for example [see [AbaAl214]Bollerslev et al (2010)].

GARCH and its variants are parametric estimation methods for VaR. Typically, these methods are focused on estimating time-varying volatility [standard deviation]. As an alternative to GARCH and its variants, several methods have been proposed in the literature. Many of them can be classified as semi-parametric as, for example, those based on the Extreme Value Theory [EVT], the Conditional Autoregressive Value at Risk approach [CAViaR] or the Filtered Historic Simulation [FHS]. Different papers have compared these methods with GARCH [or some variant] or another alternatives.

In [AbaAl214] an extensive literature review until 2013 about VaR estimation is presented. They identify almost all methods proposed in the literature and the main criteria used to compare them. Another interesting aspect is that [AbaAl214] also identified 24 articles comparing the most important methods between 2000 and 2013. In Table 1, we summarize the results obtained by Abad et al (2014) regarding the review of those articles. Both the VaR methods and the comparison criteria employed in each paper are reported. The the best methods identified by them are also indicated. In addition to these informations, available at the survey made in [AbaAl214], we report the assets considered. The set of assets and VaR methods used varies substantially.

In addition to the aforementioned VaR methods, the following alternatives were employed: i) HOM [High Order Moment time-varying]; ii) MC [Monte Carlo Simulation]; iii) NP [non-parametric estimation of the density function] and iv) APA [another basic parametric approaches]. GARCH and its variants are associated with APA group. Such approaches, as well as the traditional

non-parametric historical method [HS] are overcome in practically all comparisons. On the other hand, EVT, FHS, CAViaR and HOM have shown better performance in the VaR estimation.¹ These methods are actually better when we consider the relative frequency at which each method outperformed its competitors.

Table 1: Overview of papers that compare VaR methodologies [those considered in Abad et al 2014]

Papers that compare VaR methods considered in Abad et al 2014	What methodologies compare? (VaR Methods)		How do they compare? (Backtests)	What they compare? (Assets)
	(Best)	(Competitors)		
Abad and Benito 2013	APA	HS, EVT, MC	UC, IND, CC, BT, DQ	IBEX, CAC40, DAX 30, FTSE 100, Dow Jones, S&P 500, Nikkei 225, HANG SENG
Gerlach et al. 2011	CAViaR	APA	UC, CC, DQ	S&P 500, FTSE 100, CAC40, DAX 30, FTSEMIB.MI, TSX, Nikkei 225, HANG SENG, AORD ALL, TWII or TSEC
Sener et al. 2012	CAViaR/APA	HS, EVT, MC	DQ	IBOV, IPSA, COLCAP, PSE, BUX, MEXBOL, WIG, RTS, XU100 or ISE, JALSH, Merval, FTSE 100, Dow Jones, CAC 40, IBEX, DAX 30, Nikkei 225, AEX
Ergun and Jun 2010	HOM/EVT	APA	UC, IND, CC	S&P 500
Nozari et al. 2010	EVT	FHS	UC	BUX, PSE, NSEI, RTS, TWII or TSEC, SSE, JKSE
Polanski and Stoja 2010	HOM	APA, CAViaR	UC, IND	S&P 500, Dow Jones, Nikkei 225
Brownlees and Gallo 2010	NP	HS, APA	UC, IND, CC, DQ	BA, GE, JNJ
Yu et al. 2010	CAViaR	APA	%	S&P500, HANG SENG, Nasdaq, Nikkei 225
Ozun et al. 2010	EVT	APA	UC, IND, CC	ISE
Huang 2009	NP	HS, APA, MC	UC	AAPL, AXP, BA, CAT, CSCO, CVX, DD, DIS, GE, GS, HD, IBM, INTC, JNJ, JPM, KO, MCD, MMM, MRK, MSFT, NKE, PFE, PG, TRV, UNH, UTX, V, VZ, WMT, XOM, Dow Jones, S&P 500, Nasdaq
Marimoutou et al. 2009	FHS/EVT	HS, APA	UC, CC	WTI, Brent
Zikovic and Aktan 2009	FHS/EVT	HS, APA	UC, IND, CC	XU100 or ISE, CROBEX
Giamouridis and Ntoulia 2009	FHS/APA/EVT	HS	UC, IND, CC	HFRX-EH, HFRX-Macro, HFRX-RVA, HFRXED-DR, HFRXRV-FI, HFRX-DR, HFRXEH-EMN, HFRXED-MA
Angelidis et al. 2007	FHS	APA, EVT	UC	DJ Euro Stoxx 50
Tolikas et al. 2007	EVT	HS, APA, MC	CC	DAX 30
Alonso and Arcos 2006	APA	HS, FHS	BT	COLCAP
Bao et al. 2006	FHS/EVT	HS, APA, CAViaR, MC, NP	%	JKSE, KOSPI, Malaysia, TWII or TSEC, SET
Bhattacharyya and Ritolia 2008	EVT	HS, APA	UC	NSEI
Kuester et al. 2006	FHS/EVT	HS, APA	UC, IND, CC, DQ	Nasdaq
Bekiros and Georgoutsos 2005	EVT	HS, APA	UC, CC	Dow Jones, Cyprus Stock Exchange ²
Gençay and Selçuk 2004	EVT	HS, APA	%	Merval, IBOV, HANG SENG, JKSE, KOSPI, MEXBOL, STI, TWII or TSEC, XU100 or ISE
Gençay et al. 2003	EVT	HS, APA	%	XU100 or ISE, S&P 500
Darbha 2001	EVT	HS, APA	%	NSEI
Danielsson and de Vries 2000	EVT	APA	%	S&P 500, HANG SENG, Word Logic Company ² , DAX 30, FTSE 100, Gold Bullion, JPM, MMM, MCD, INTC, IBM, XRX, XOM, USD/FRF ¹ , USD/DEM ¹ , USD/YEN, GBP/USD

Source:

The above information was obtained mainly from Abad et al 2014 [except in the case of the last column] and in the mentioned references themselves.

Notes:

- (i) In addition to the methods used in this work, another methods appeared in the papers considered by Abad et al 2014: **HS** (Historical Simulation); **MC** (Monte Carlo Simulation); **NP** (Non-Parametric estimation of the density function); **APA** (Another basic Parametric Approaches estimated under different distributions, including the normal distributions, t-Student distributions, Cornish-Fisher expansion and Riskmetrics).
- (ii) In addition to the backtests used in his work, another procedures were considered in the papers considered by Abad et al 2014: **IND** (test based only on the statistic for the serial independence); **BT** (backtesting criterion, a backtesting criterion, where a normal approximation is used for the standardized exceptions); **%** (comparisons between the observed and expected percentages of exceptions without a formal statistical test).
- (iii) In the last column, almost all assets exposed were considered in our empirical study. We discarded two assets marked with **(1)** and replace the assets marked with **(2)** by the USD/EUR exchange rate.

(1)

¹ The NP method is best in 2 of 3 comparisons. We did not consider NP because it was the best only when compared to the worst methods. In addition, the method was compared few times.

With respect to the comparison criteria [backtests], we observe that only 6 alternatives were considered: i) Unconditional Coverage test [UC]; ii) Conditional Coverage test [CC]; iii) Independence Test [IND]; iv) Backtesting Criterion [BT]; v) Dynamic Quantile Test [DQ] and vi) comparison of Percentages of Exceptions [%]. UC, BT and percentage of exceptions attempt to evaluate the same question: the adjustment between the observed and expected exception frequencies. Of these 3 alternatives only 2 [UC and BT] are statistical tests. UC is more common than BT. The CC procedure combine test statistics of UC and IND procedures. Understandably, CC is employed in more papers than IND. Finally, the DQ procedure addresses the issue of exception dependence in relation to its outdated terms. Although its purpose is similar to that of the IND, it is not a particular case of CC.

Considering all papers, a total of 82 different assets was considered.² The set of assets comprises market indices, stock exchange indices, quotes of companies, hedge fund indices, commodities prices and exchange rates. The number of assets considered in each work varies from 1 to 33 [6, on average]. The S&P 500 is the most frequently used asset [8 times]. Second are the Japanese, German and Chinese stock indexes, as well as the Dow Jones index [4 times].

In almost cases, the EVT method is the best one. However, EVT is usually compared to the worst methods or only with a small number of best methods. In almost all cases, the EVT method is the best. However, EVT is usually compared to the worst methods or only with few methods among those considered to be the best. In particular, no study simultaneously compares EVT with CAViaR, FHS and HOM [all the best methods].

In the present work, we compare all the best methods [CAViaR, EVT, FHS and HOM], both through a large empirical analysis, and through an extensive simulation study. In these comparisons we employ the main backtests procedures [UC, CC, DQ] and a relatively recent proposal [VQR]. The VaR methods chosen are briefly described in Section 2, whereas the adopted backtests are succinctly presented in Section 3. In section 4 we report the main results obtained in the empirical analysis, after briefly describing the dataset used. In section 5 we explain the simulation strategy considered in this work and the most important results obtained in the simulation study. Section 6 concludes.

² A short description of each asset is present at Tables 10-14.

2 VaR Methods

As mentioned before, [AbaAl214] presented an extensive literature review about VaR estimation. In particular, they identified the best methodologies for producing VaR estimates [VaR Methods, for short]. Their investigation was carried out by the analysis of another works that compare the performance of distinct VaR Methods, as we intend to do here. We have reviewed the works mentioned there and have identified that the main classes of VaR methods are those associated with: **i) EVT** [Extreme Value Theory]; **ii) FHS** [Filtered Historical Simulation]; **iii) CaViaR** [Conditional Autoregressive Value at Risk]; **iv) HOM** [High-Order Moment Time-Varying Distribution]. Basically, each of these methods performed better than its competitors in some comparative study.³ Therefore, we intend to compare the main methods according to the empirical literature.⁴ In this section we will briefly describe the four methods used in our comparison [the same ones mentioned above]. Conventionally, assume that $\{r_t\}_{t=1}^n$ is a time series of stocks [or potfolio] returns. Each return r_t is seen as a realization of a random variable [r.v.] R_t and $\{R_t\}_t$ is the stochastic process associated - strictly stationary, by assumption. We could define VaR in terms of the Loss $L_t \equiv -R_t$ [$l_t \equiv -r_t$ stands for observed loss]. Denoting VaR at time τ by V_τ and following [Jor206], we could define it implicitly by the expression $\alpha = P(L_\tau \geq V_\tau | \Omega_{\tau-1})$, where $1 - \alpha$ is the VaR confidence level and $\Omega_{\tau-1}$ is the information set available at time $\tau - 1$. So, conditionally to the information set, the probability of observing a loss greater than VaR at time τ is α . If R_τ has conditional α -quantile $\mathcal{R}_{\alpha,\tau}$, then, $V_\tau = -\mathcal{R}_{\alpha,\tau}$ and $V_\tau = F_{R_\tau|\Omega_{\tau-1}}^{-1}(\alpha)$ if $F_{R_\tau|\Omega_{\tau-1}}$ is the conditional cumulative distribution function [c.d.f.] of R_t with respect to $\Omega_{\tau-1}$.⁵ Let $\mu_t \equiv E(R_t | \Omega_{t-1})$ and $\sigma_t \equiv [Var(R_t | \Omega_{t-1})]^{1/2}$ be, respectively, the conditional mean and the conditional standard deviation [conditional volatility]. Some methods assumes that for a white-noise process $\{Z_t\}$, the stochastic process $\{R_t\}_t$ satisfies

$$R_t = \mu_t + \sigma_t Z_t, \text{ where} \quad (2)$$

$$\mu_t = \phi R_{t-1} \text{ and } \sigma_t^2 = \alpha_0 + \alpha_1 Z_{t-1}^2 + \beta \sigma_{t-1}^2. \quad (3)$$

³ We discarded the parametric methods based on the Gaussian and t-Student distributions because they were the best methods in very few situations [in most cases, they were overcome by some alternative method considered here].

⁴ It is important to note that so far we have not found another paper comparing the same methods.

⁵ Let us assume that $F_{R_\tau|\Omega_{\tau-1}}$ is strictly increasing.

The two equations above describes a volatility model with varying mean.⁶ If the α -quantile of Z_t is denoted by z_α , then $V_\tau = -(\mu_\tau + \sigma_\tau z_\alpha)$. As discussed in [LutKra204], the structure presented describes an AR(1)-GARCH(1,1) model. As pointed by them, if $\hat{\phi}$, $\hat{\alpha}_0$, $\hat{\alpha}_1$ and $\hat{\beta}$ denote the corresponding parameter estimates, the fitted mean $\hat{\mu}_\tau$ and the fitted volatility $\hat{\sigma}_\tau$ could be calculated by replacing the unknown parameter value by its estimates and using some recursive update scheme. As we will see, the main difference between the GARCH approaches and the methods considered here is the treatment of $\{Z_t\}$. However, in both cases [GARCH or alternative methods EVT, FHS and HOM], $\hat{\mu}_\tau$ and $\hat{\sigma}_\tau$ could be obtained in a similar way and VaR is estimated by: (i) $\hat{V}_\tau = -\hat{\mu}_\tau - \hat{\sigma}_\tau z_\alpha$ if z_α is known; (ii) $\hat{V}_\tau = -\hat{\mu}_\tau - \hat{\sigma}_\tau \hat{z}_\alpha$, otherwise [$\hat{z}_\alpha$ is an estimate of z_α].

EVT [Extreme Value Theory]: The approach based on Extreme Value Theory [EVT] seems to be one of the best methods to estimate VaR, according to the review presented in [AbaAl214]. Among 18 comparative studies, these methods were not the best in only 3 cases. In the EVT approach, the focus is on the limit distribution of the extreme returns. There are two main alternative approaches to EVT: (1) Peaks Over Threshold [POT] and (2) Block Maxima Models [BMM]. However, analyzing the articles considered by [AbaAl214], we find that most of them [mainly the most recent ones] adopt the POT approach. The POT approach provide better VaR estimates, for example, in [ErgJun210], [Noz210], [MarAl209], [ZikAkt209], [BhaRit208], [KueAl206], [BekGeo205] [they also consider BMM and both methods perform similarly]; [GenSel204]. In contrast, POT approach was considered in [AbaBen213] and [AngAl207], but in both cases the POT method was outperformed by its competitors [alternatives based on t-Student and normal distributions, respectively]. In [AbaBen213], [BhaRit208] and [BaoAl206], POT outperforms BMM⁷. Thus, we will consider only the POT approach here [in particular, EVT will refer to the Extreme Value Theory in the Peaks Over Threshold approach]. More specifically, we will follow the conditional EVT approach method presented in [McNFre200]. First, they propose to adopt a volatility model for returns. Then apply EVT-based analysis [POT] for standardized returns. Their method can be summarized in two steps:

- (1) Firstly, adopting a gaussian pseudo-maximum-likelihood [PML], estimate the parameters in 2. In another words, only for estimation purposes, let us suppose that $Z_t \sim N(0; 1)$. It allows to obtain the mean and volatility estimates $[\hat{\mu}_\tau$ and $\hat{\sigma}_\tau]$. Define the **negative** residuals: $\hat{z}_t = -\left(\frac{r_t - \hat{\mu}_t}{\hat{\sigma}_t}\right)$.

⁶ Obviously, the Conditional Mean and Variance equations could be generalized to incorporate another lagged terms. We opt by not to adopt another lagged terms.

⁷ According to [BaoAl206], POT method tends to utilizes data more efficiently than in the BMM approach.

- (2) In step 2, we must deal with the task of estimating the $(1 - \alpha)$ -quantile of $\tilde{Z}_t$, denoted by $\tilde{z}_{1-\alpha}$.⁸ An EVT⁹ approach is proposed for modelling the distribution of $\tilde{Z}_t$. Let $\tilde{z}_{[1]} \geq \tilde{z}_{[2]} \geq \dots \geq \tilde{z}_{[n]}$ be the *decreasing* ordered values of $\{\tilde{z}_t\}_{t=1}^n$. $\tilde{z}_{1-\alpha}$ could be estimated by $\hat{\tilde{z}}_{1-\alpha}$ satisfying

$$1 - \alpha = \hat{F}_{\tilde{Z}_t}(\hat{\tilde{z}}_{1-\alpha}), \text{ where } \hat{F}_{\tilde{Z}_t}(z) = 1 + \frac{k}{n} [G_{\hat{\xi}, \hat{\lambda}}(z - \tilde{z}_{[k+1]}) - 1]. \quad (4)$$

Note that $\hat{F}_{\tilde{Z}_t}$ is the estimated cumulative distribution function of $\tilde{Z}_t$ under EVT. It depends on both an integer k and the Generalized Pareto Distribution¹⁰ [GDP] with shape parameter $\hat{\xi}$ and scale parameter $\hat{\lambda}$. If $\hat{\xi} \neq 0$, letting $\hat{z}_\alpha = -\hat{\tilde{z}}_{1-\alpha}$, then

$$\hat{V}_\tau = -\hat{\mu}_\tau - \hat{\sigma}_\tau \hat{z}_\alpha = -\hat{\mu}_\tau + \hat{\sigma}_\tau \left(\tilde{z}_{[k+1]} + \frac{\hat{\lambda}}{\hat{\xi}} \left[\left(\frac{\alpha n}{k} \right)^{-\hat{\xi}} - 1 \right] \right). \quad (5)$$

While k is a parameter that can be chosen, $\hat{\xi}$ and $\hat{\lambda}$ are the maximum likelihood estimates obtained when the data are the largest k excess returns $\tilde{z}_{[1]} - \tilde{z}_{[k+1]}, \dots, \tilde{z}_{[k]} - \tilde{z}_{[k+1]}$. $\tilde{z}_{[k+1]}$ is the threshold, automatically determined by choice of k .¹¹ Under POT assumption, the largest k excess returns are a data sample drawn from a GPD population with unknown shape parameter ξ and scale parameter λ .

FHS [Filtered Historical Simulation]: Except¹² for [Noz210], in all papers considered by [AbaAl214] the method FHS worked as well as EVT. In [MarAl209], [ZikAkt209] and [BaoAl206], for example, both were chosen the best methods. In [AngAl207] FHS outperformed EVT. The FHS method was proposed in [BarAl199]. Similar to the EVT approach, it is a semi-parametric method for estimating VaR by combining: i) a parametric model for volatility

⁸ Note that it will be an estimate of $-z_\alpha$, where z_α is the α -quantile of Z_t .

⁹ Specifically, a POT approach.

¹⁰ The Generalized Pareto Distribution has cumulative distribution function $G_{\xi, \lambda}$:

$$G_{\xi, \lambda}(y) = \begin{cases} 1 - (1 + \xi y / \lambda)^{-1/\xi}, & \text{if } \xi \neq 0 \\ 1 - \exp(-y/\lambda)^{-1/\xi}, & \text{if } \xi = 0 \end{cases}.$$

¹¹ [McNFre200] addresses the issue of choose an appropriate value for k . In the context of GPD distribution they comment that "as long as k is greater than about 50 the method is robust". We chose $k = 67$ which corresponds to approximately 15% of the sample size.

¹² In [Noz210] FHS was outperformed by EVT. In [AloArc206] FHS was outperformed by some gaussian parametric approach. These were the only two cases where the method was not chosen as best.

[for example, a GARCH model] and ii) a free-distribution approach for the standardised returns. Essentially, if $\{\hat{\mu}_t\}_t$ are the estimated conditional means and $\{\hat{\sigma}_t\}_t$ are the estimated conditional volatilities obtained from the return data $\{r_t\}_t$, it is assumed that the standardised returns $\{z_t\}_t$ [$z_t = (r_t - \hat{\mu}_t) / \hat{\sigma}_t$] are realizations of i.i.d. random variables $\{Z_t\}_t$. Let $\hat{z}_\alpha$ be the empirical α -quantile of $\{z_t\}_t$. Then the estimated VaR is¹³ $\hat{V}_\tau = -\hat{\mu}_\tau - \hat{\sigma}_\tau \hat{z}_\alpha$.

CAViaR [Conditional Autoregressive Value at Risk]: Similar to EVT and FHS methods, the Conditional Autoregressive Value at Risk [CAViaR] is also a semi-parametric approach. The investigation conducted in [AbaAl214] reveal that CAViaR is a potentially powerful method, although the literature has made relatively few comparisons using such a method. In [SenAl212] CAViaR outperforms EVT. CAViaR is also the best method in [GerAl211] and [YuAl210] comparisons. [BaoAl206] study points that CAViaR is not good enough as FHS and EVT. In [PolSto210] is outperformed by HOM. The insufficient frequency of comparisons with the CAViaR method, as well as the variability of the results found, make the classification of this method inconclusive and indicate the need for further investigations. The CAViaR method was proposed in [EngMan204]. Their approach directly explores the fact that $-V_t$ is the α -quantile of the conditional distribution of R_t . More generally, they assume that $V_t = f_t(\beta) \equiv f_t(\mathbf{x}_{t-1}, \beta)$ depends both on a unknown vector $\beta = (\beta_0, \beta_1, \dots, \beta_d)$ and $\mathbf{x}_{t-1}$, a vector of time $t - 1$ observable variables. A generic CaViaR is specified by $f_t(\beta) = \beta_0 + \sum_{i=1}^q \beta_i f_{t-i}(\beta) + \sum_{j=q+1}^d \beta_j \ell(\mathbf{x}_{t-j})$, where ℓ is a known function expressing the VaR dependency with respect to the observable variables. [EngMan204] presented the following four alternative specifications for $f_t(\beta)$:

(i) **Adaptive [ADP]:**

$$f_t(\beta) = f_{t-1}(\beta) + \beta_1 \left\{ [1 + \exp(G[r_{t-1} - f_{t-1}(\beta)])]^{-1} - \alpha \right\}$$

(ii) **Asymmetric Slope [ASY]:**

$$f_t(\beta) = \beta_1 + \beta_2 f_{t-1}(\beta) + \beta_3 [\max(r_{t-1}, 0)] + \beta_4 [\max(-r_{t-1}, 0)]$$

(iii) **Indirect GARCH(1,1) [GAR]:**

$$f_t(\beta) = \left(\beta_1 + \beta_2 f_{t-1}^2(\beta) + \beta_3 r_{t-1}^2 \right)^{1/2}$$

¹³ Equivalently, it is common to generate VaR estimates by sampling [with replacement] the standardised returns $\{z_t\}_t$. If $\{z_j^*\}_{j=1}^J$ denotes any particular sample, define $r_{j,\tau}^* = \hat{\mu}_\tau + \hat{\sigma}_\tau z_{j,\tau}^*$. Thus $\hat{V}_\tau$ could be defined as the empirical α -quantile of $-r_{1,\tau}^*, \dots, -r_{J,\tau}^*$.

(iv) **Symmetric Absolute Value [SAV]:**

$$f_t(\beta) = \beta_1 + \beta_2 f_{t-1}(\beta) + \beta_3 |r_{t-1}|.$$

In the CAViaR approach is adopted a nonlinear quantile regression for modelling VaR. Therefore, as usual for quantile regression models [see, for example, [Koe205]] the parameter vector β could be estimated by minimizing the asymmetric absolute loss: $\min_{\beta} \frac{1}{T} \sum_{t \leq T} [\alpha - \mathbb{I}(r_t < -f_t(\beta))] (r_t + f_t(\beta))$. The solution, denoted by $\hat{\beta}$, allows to define the VaR estimate $\hat{V}_t = f_t(\hat{\beta})$. To solve this problem it is necessary to deal with a nonlinear programming problem. Practical details can be found in [EngMan204]. Finally, for comparison purposes, it should be noted that only ASY, GAR and SAV were considered. The main justification is that ADP was not used in comparisons in which CAViaR was chosen as the best method.¹⁴

HOM [High-Order Moment Time-Varying Distribution]: Traditionally, in parametric approaches, only the first two moments of the return distribution can vary over time. However, at least since [Han194], alternative approaches have been proposed to model dependencies of high order moments. See, for example, [ErgJun210] and [PolSto210]. In those cited works, VaR methods based on High-Order Moment Time-Varying Distribution [HOM] outperformed some traditional parametric approaches¹⁵. We must stress that in review conducted in [AbaAl214], HOM method were only compared in these two mentioned work. In spite of the few cases, this method has not been defeated at any time. Here, we adopt an approach similar to that appearing in [BalAl208] and [ErgJun210], which are intrinsically related to the original Hansen's proposal [see [Han194]]. As before, μ_t and σ_t denotes, respectively, the returns conditional mean and standard deviation. Let η_t be the parameter [high order moments] affecting the function f_{η_t} , the density of the standardized returns $Z_t \equiv \frac{R_t - \mu_t}{\sigma_t}$. We take f_{η_t} to be the Skewed Normal Distribution with skew parameter η_t .¹⁶ Then, the density of the returns is $\frac{1}{\sigma_t} f_{\eta_t} \left(\frac{r_t - \mu_t}{\sigma_t} \right)$.¹⁷ Assume, as before, that R_t, μ_t, σ_t evolve according to equations 2 and 3. Here, we take η_t to be the skew parameter¹⁸ and assume that

$$\lambda_t = a_0 + a_1 Z_{t-1} + a_2 Z_{t-1}^2 + a_3 \lambda_{t-1}, \quad (6)$$

¹⁴ In addition, in ADP case an arbitrary constant G must be specified for each analysed series as mentioned in [EngMan204]. Computing routines to choose an appropriate value can be time consuming.

¹⁵ In [ErgJun210], HOM and EVT methods performs similarly.

¹⁶ Unlike [Han194] that considered Skewed t-Student Distribution.

¹⁷ In GARCH, EVT and FHS $f_{\eta_t} = f$ [η_t is time invariant].

¹⁸ It is possible, additionally or not, to consider a shape parameter.

where $\eta_t = L + \frac{U-L}{1-e^{-\lambda t}}$.¹⁹ Once estimates $\hat{\eta}_t$, $\hat{\mu}_t$ and $\hat{\sigma}_t$ are obtained [using similar methods to those mentioned in [Han194]], defining $\hat{z}_\alpha$ as the α -quantile associated with the distribution represented by the density $f_{\hat{\eta}_t}$, we can obtain a VaR estimate by $\hat{V}_\tau = -\hat{\mu}_\tau - \hat{\sigma}_\tau \hat{z}_\alpha$.

3 Backtests

As pointed by [Roc216], a backtesting procedure allows us to assess the VaR Methods. Typically, it is done by comparison *ex ante* VaR predictions with the observed returns. [Jor206] highlights that backtesting is a model validation strategy, a kind of diagnosis tool. Usually those evaluations mentioned are made by implementation of a statistical hypothesis test, that we call a backtest. According to [AbaAl214], VaR methods can be compared or evaluated by means of (i) their accuracy or (ii) reported losses. We will consider only the accuracy aspects. Considering the statistical test approach, a specific backtest, used to evaluate whether a specific VaR method is [or not] appropriate for a specific time series of return, will result in a statistical statement as "we should reject [or not to reject] the specific VaR Method adopted for that specific return time series at the significance level adopted". The most popular backtests are the Unconditional Coverage Test [UC], proposed by [Kup195], and the Conditional Coverage Test [CC], proposed by [Chr198]. They were considered in several papers comparing VaR Methods, as [AbaBen213], [GerAl211], and [BroGal210], for example. [EngMan204] proposed the Dynamic Quantile Test [DQ], also considered in all papers mentioned before and used as the unique accuracy evaluation method in [SenAl212]. In most recent years, [GagAl211] presented a new backtest, based on the quantile regression, that we will identify as the Quantile Regression Test [VQR]. This test is not as popular as the others, but according with simulation results presented in [GagAl211], it seems to be most powerful than the traditional alternatives. [Chr198]. We opt here for using only the four backtests mentioned before. Excepting to VQR, they are, in fact, the most popular backtests. In the review made in [AbaBen213] are exhibited papers that compare accuracy of VaR Methods by comparing the expected and observed frequencies of exceptions [an exceptions occurs if the observed return is smaller than VaR] - e.g., [YuAl210]. Some works [for example, [AbaBen213]] considered a test for independency of exceptions. As we will discuss, the conditional test incorporate the analysis of exceptions dependency. Below we present a brief review of the four backtests chosen: UC, CC, DQ and VQR. Before start, however, for convenience, let $r_1, \dots, r_n, r_{n+1}, \dots, r_{n+T}$ denote the observed time series of returns. Consider that the one-day-ahead

¹⁹ Following [Han194], we assume that $\eta_t \in [L, U]$ for known constants L and U .

VaR forecasts $\hat{V}_{n+1}, \dots, \hat{V}_{n+T}$ are obtained by use the initial n data points $r_1, \dots, r_n$ and some general VaR Method. If R_τ is the r.v. associated with r_τ , then, let $I_\tau = \mathbb{I}(R_\tau < -\hat{V}_\tau)$ denote an random exception indicator [$I_\tau = 1$, if $R_\tau < -\hat{V}_\tau$ and $I_\tau = 0$, otherwise]. Then, we observe exceptions $i_{n+1}, \dots, i_{n+T}$ [realizations of $I_{n+1}, \dots, I_{n+T}$].

Unconditional Coverage Test [UC]: The exceptions $I_{n+1}, \dots, I_{n+T}$ are Bernoulli random variables. [Kup195] presented the Unconditional Coverage [UC] Test, assuming that $I_{n+1}, \dots, I_{n+T}$ are independent with a common success probability p , *i.e.*, $I_{n+1}, \dots, I_{n+T}$ are i.i.d. and $I_{n+i} \sim \text{Ber}(p)$. The UC procedure consists in test the null hypothesis $H_0 : p = \alpha$ against the alternative $H_0 : p \neq \alpha$ where $1 - \alpha$ is the VaR confidence level. In this approach a simple Likelihood Ratio Test is employed, whose test statistic is the Likelihood Ratio

$$LR_{uc} = 2 \log \left(\frac{[1 - \bar{I}]^{T(1 - \bar{I})} [\bar{I}]^{T\bar{I}}}{[1 - \alpha]^{T(1 - \bar{I})} \alpha^{T\bar{I}}} \right), \quad (7)$$

where $\bar{I} = \frac{1}{T} \sum_{i=1}^T I_{n+i}$ is the observed frequency of exceptions. Under the null hypothesis, LR_{uc} is asymptotically chi-square distributed with 1 degree of freedom: $LR_{uc} \stackrel{a}{\sim} \chi_1^2$ [[BicDok215]]. Note that reject the null hypothesis $H_0 : p = \alpha$ means that the data provide evidences that the estimated *VaR* produces exceptions incompatible with the associated confidence level, indicating that we must reject the VaR method adopted for that observed time series of return.

Conditional Coverage Test [CC]: While the unconditional approach assumes independence between exceptions and evaluates only the gap between observed and expected number of exceptions, in the conditional approach, both adjustment and independence are tested simultaneously. The Conditional Coverage Test [CC Test], proposed by [Chr198], is based on statistics LR_{cc} , defined as

$$LR_{cc} \equiv LR_{uc} + LR_{ind}, \quad (8)$$

where LR_{uc} refers to the same statistics defined in equation 7 and LR_{ind} corresponds to a likelihood ratio statistic to test the null hypothesis of independence against first-order Markov dependence, as we shall see below. The null hypothesis of the CC test postulates simultaneously both that $p = \alpha$ and the exceptions are independent. Under the null, LR_{cc} is asymptotically chi-square distributed with 2 degree of freedom: $LR_{cc} \stackrel{a}{\sim} \chi_2^2$ [[Jor206]]. Intuitively, LR_{cc} is the sum of two independent chi-square random variables. To understand LR_{ind} , let π_1 be the probability of observe an exception in time $\tau + 1$ if an exception has occurred in time τ and, analagously, let π_0 be the probability of observe an exception in time $\tau + 1$ if an exception has not occurred in time τ .

Moreover, for $i, j \in \{0, 1\}$, let X_{ij} denotes the the total of pairs $(\tau, \tau + 1)$ satisfying $I_{\tau+1} = j$ and $I_\tau = i$ ²⁰. LR_{ind} is the Likelihood Ratio Statistic adopted when testing $H_0 : \pi_1 = \pi_0$ vs $H_0 : \pi_1 \neq \pi_0$:²¹

$$LR_{ind} = 2 \ln \left(\frac{\pi_1^{X_{11}} (1 - \hat{\pi}_1)^{X_{10}} \hat{\pi}_0^{X_{01}} (1 - \hat{\pi}_0)^{X_{00}}}{\hat{\pi}_{R1}^{X_{11}} (1 - \hat{\pi}_{R1})^{X_{10}} \hat{\pi}_{R2}^{X_{01}} (1 - \hat{\pi}_{R2})^{X_{00}}} \right),$$

where $\hat{\pi}_0 = \frac{X_{01}}{X_{00} + X_{01}}$ and $\hat{\pi}_1 = \frac{X_{11}}{X_{10} + X_{11}}$ are the unrestricted maximum likelihood estimators and $\hat{\pi}_{R1} = \hat{\pi}_{R2} = \frac{X_{01} + X_{11}}{X_{00} + X_{01} + X_{10} + X_{11}}$ are the restricted maximum likelihood estimators.

Dynamic Quantile Test [DQ]: In [EngMan204] a Dynamic Quantile [DQ] test is presented in which a possible relation between the exception I_τ and variables in the information set $\Omega_{\tau-1}$ is investigated. The test is based on *Hit* variable defined as:

$$Hit_t = I_t - \alpha. \quad (9)$$

For each t , Hit_t assumes value $(1 - \alpha)$ if an exception occurs in time t and $-\alpha$, otherwise. Hit_t has null expectation. Their proposal is to adopt a Wald test in which the null hypothesis $H_0 : \theta = 0$ and $\theta = (\theta_1, \dots, \theta_d)$ is the vector of slopes in the following linear regression model:

$$Hit_t = \sum_{i=1}^d \theta_i X_{t,i} + \varepsilon_t.$$

The variables $X_{t,1}, \dots, X_{t,d}$ are explanatory variables contained in the information set Ω_{t-1} . In particular, lagged hit terms $Hit_{t-1}, \dots, Hit_{t-p}$ could be adopted as regressors. If $\hat{\theta} = (\hat{\theta}_1, \dots, \hat{\theta}_d)$ is the OLS estimate of θ , the Wald test statistics is

$$\Delta = \hat{\theta}^\top Var(\hat{\theta})^{-1} \hat{\theta},$$

where $Var(\hat{\theta})^{-1}$ is the variance matrix of $\hat{\theta}$ and $Var(\varepsilon_t)$ is set to $\alpha(1 - \alpha)$ ²². Under the null, Δ is asymptotically chi-squared distributed with d degrees of freedom.

Quantile Regression Test [VQR]: [GagAl211] propose a test for Value-at-Risk model based on Quantile Regressions [VQR]. In their approach, they suggest to adopt a linear quantile regression model using R_t as the dependent variable and the estimated VaR $\hat{V}_t$ as regressor:

$$R_t = b_0 + b_1 \hat{V}_t + \varepsilon_t, \text{ where } P(\varepsilon_t \leq 0 | \Omega_{t-1}) = \alpha.$$

²⁰ Using the exceptions indicators, $X_{00} = \sum_{t=1}^T (1 - I_t)(1 - I_{t-1})$, $X_{01} = \sum_{t=1}^T I_t(1 - I_{t-1})$, $X_{10} = \sum_{t=1}^T (1 - I_t)I_{t-1}$ and $X_{11} = \sum_{t=1}^T I_t I_{t-1}$.

²¹ For simplicity, suppose I_0 is known.

²² Repair that $Var(Hit_t) = Var(\varepsilon_t) = \alpha(1 - \alpha)$.

Under this model $b_0 + b_1 \widehat{V}_t$ is the α -quantile of R_t . Then, if $\widehat{V}_t$ is the correct VaR measurement, $b_0 = 0$ and $b_1 = -1$ [$b_0 + b_1 \widehat{V}_t = -\widehat{V}_t$]. The true values of b_0 and b_1 are unknown, but by minimizing the asymmetric absolute loss²³ [which can be done by solving a linear programming problem as discussed in [Koe205]] we can obtain their estimates $\widetilde{b}_0$ and $\widetilde{b}_1$, respectively. By setting $\theta = (b_0, b_1 + 1)$, the null hypothesis is $H_0 : \theta = 0$. As in the DQ test, they suggest the adoption of a Wald test. Following, for example, [BicDok215], if $\widetilde{\theta} = (\widetilde{b}_0, \widetilde{b}_1)$ and $\widetilde{Var}(\widetilde{\theta})$ is an estimate of the variance matrix of $\widehat{\theta}$ the Test Statistic for such Wald Test has the form

$$\zeta = \widehat{\theta}^\top \widetilde{Var}(\widetilde{\theta})^{-1} \widehat{\theta}.$$

Under the null hypothesis, the test statistic ζ is asymptotically chi-squared distributed with 2 degrees of freedom.²⁴

4 Empirical Analysis

In our empirical study, we compared the performance of the main VaR methods [EVT, FHS, CAViaR and HOM] when applied to exactly 80 daily return time series concerning: (i) 35 market or stock exchange indices; (ii) stock quotes of 31 different companies; (iii) 8 hedge fund indices; (iv) 3 commodities and (v) 3 exchange rates. Regarding to stock exchanges or market indices, were contemplated both emerging markets [*e.g.*, Brazil, Colombia, Mexico, Turkey, South Africa] and developed markets [*e.g.*, USA, England, France, Germany, Japan]. In relation to companies, were included 31 enterprises, many of them considered in the Dow Jones Industrial Average Index. Details about the collected data are shown at Tables 10-14. *SOURCE* informs you where we collect the data [**Y** for "yahoo finance", **G** for "google finance", **H** for download in "www.hedgefundresearch.com", **O** for other sources²⁵]. *SIZE* indicates the sample size and *CODE* displays the abbreviated download name in the associated data source [google finance or yahoo finance]. Finally, **W1** and **W2**

²³ Similar to the CAViaR approach to estimate VaR [EngMan204].

²⁴ Consistent estimates of the variance matrix of $\widehat{\theta}$ can be obtained using the methods discussed in [KoeMac199].

²⁵ Séries retiradas de outras fontes (8):

COLCAP (<http://www.banrep.gov.co>), CROBEX (<http://zse.hr>), PSE (<https://www.pse.cz>), RTS (<http://www.moex.com>), SET (<https://stooq.com>), Brent (<https://www.eia.gov>), WTI (<https://fred.stlouisfed.org>) e GoldBullion (<https://www.quandl.com>).

reflect two distinct sets of weights that we have adopted here, representing the relative importances of each asset.

Table 10: Market or Stock Exchange Indices

ASSET NAME	DESCRIPTION	SOURCE	CODE	SIZE	START	END	W1	W2
AEX	NETHERLANDS	Y	^AEX	2745	01/02/2007	22/08/2017	0,007	0,009
AORD ALL	AUSTRALIA	Y	^AORD	2713	02/01/2007	21/09/2017	0,007	0,009
BUX	HUNGARY	Y	^BUX	2700	23/01/2007	01/09/2017	0,014	0,017
CAC40	FRANCE	Y	^FCHI	2743	17/01/2007	06/09/2017	0,022	0,026
COLCAP	COLOMBIA	O	COLCAP	2351	24/01/2008	24/08/2017	0,014	0,017
CROBEX	CROATIA	O	CROBEX	2433	05/12/2007	16/08/2017	0,007	0,009
DAX 30	GERMANY	Y	^GDAXI	2724	18/01/2007	05/09/2017	0,036	0,043
DJ Euro Stoxx 50	EURO ZONE	Y	^STOXX50E	2700	17/01/2007	07/09/2017	0,007	0,009
Dow Jones	USA	Y	^DJI	2700	20/02/2007	07/08/2017	0,036	0,043
FTSE 100	UK	Y	^FTSE	2711	14/02/2007	08/08/2017	0,029	0,034
FTSEMIB.MI	ITALY	Y	FTSEMIB.MI	2724	25/01/2007	29/08/2017	0,007	0,009
HANG SENG	CHINA	Y	^HSI	2651	09/01/2007	14/09/2017	0,036	0,043
IBEX	SPAIN	Y	^IBEX	2740	08/02/2007	15/08/2017	0,014	0,017
IBOV	BRAZIL	Y	^BVSP	2657	04/01/2007	19/09/2017	0,014	0,017
IPSA	CHILE	Y	^IPSA	2694	08/01/2007	13/09/2017	0,007	0,009
ISE	IRELAND	Y	^ISEQ	2721	24/01/2007	30/08/2017	0,007	0,009
JALSH	SOUH AFRICA	G	JSE:JSE	2680	07/02/2007	16/08/2017	0,007	0,009
JKSE	INDONESIA	Y	^JKSE	2629	23/01/2007	30/08/2017	0,022	0,026
KOSPI	KOREA	G	KRX:KOSPI	2659	29/01/2007	25/08/2017	0,014	0,017
Malaysia	MALAYSIA [Kuala LumpuR]	Y	^KLSE	2654	31/01/2007	24/08/2017	0,007	0,009
MERVAL	ARGENTINA	Y	^MERV	2643	02/01/2007	21/09/2017	0,014	0,017
MEXBOL	MEXICO	Y	^MXX	2694	31/01/2007	23/08/2017	0,014	0,017
Nasdaq	USA	Y	NDAQ	2700	21/02/2007	04/08/2017	0,022	0,026
Nikkei 225	JAPAN	Y	^N225	2630	31/01/2007	29/08/2017	0,036	0,043
NSEI	INDIA	Y	^NSEI	2473	08/10/2007	01/09/2017	0,022	0,026
PSE	CZECH REPUBLIC	O	PSE	4538	11/08/1999	16/08/2017	0,014	0,017
RTS	RUSSIA	O	RTS	5498	05/10/1995	01/08/2017	0,014	0,017
S&P 500	USA	Y	^GSPC	2700	22/02/2007	03/08/2017	0,058	0,069
SET	THAILAND	O	SET	7383	14/08/1987	21/07/2017	0,007	0,009
SSE	CHINA	Y	000001.SS	2613	12/01/2007	13/09/2017	0,007	0,009
STI	SINGAPORE	Y	^STI	2690	07/02/2007	16/08/2017	0,007	0,009
TSX	CANADA	G	INDEXTSI: OSPTX	2692	05/01/2007	18/09/2017	0,007	0,009
TWII or TSEC	TAIWAN	Y	^TWII	2651	09/02/2007	14/08/2017	0,029	0,034
WIG	POLAND	G	WSE:WIG	2682	02/02/2007	21/08/2017	0,007	0,009
XU100 or ISE	TURKEY	Y	^XU100	2771	19/02/2007	10/08/2017	0,029	0,034

(10)

The assets we analysis practically form the union of the assets considered in those 24 papers considered by [AbaAl214] - *cf.* Table 1.²⁶ In the weighting scheme **W1**, the weight of each asset is the relative frequency of times it was analyzed in the 24 empirical studies. In the weighting scheme **W2**, we take the product between the relative frequency of the asset group by the relative frequency of the asset itself within the group.

²⁶ Only four assets were discarded. Both "USD / FRF" and "USD / DEM" exchange rates were replaced by "USD/EUR" exchange rate. We also do not use the two assets associated with "Cyprus Stock Exchange" and "Word Logic Company" due to difficulties in obtaining the data.

Table 11: Companies Stocks

ASSET NAME	DESCRIPTION	SOURCE	CODE	SIZE	START	END	W	W2
AAPL	Apple	Y	AAPL	2700	27/02/2007	31/07/2017	0,007	0,003
AXP	American Express	Y	AXP	2700	26/02/2007	01/08/2017	0,007	0,003
BA	Boeing	Y	BA	2700	28/02/2007	28/07/2017	0,014	0,005
CAT	Caterpillar	Y	CAT	2700	01/03/2007	27/07/2017	0,007	0,003
CSCO	Cisco	Y	CSCO	2700	05/03/2007	25/07/2017	0,007	0,003
CVX	Chevron	Y	CVX	2700	02/03/2007	26/07/2017	0,007	0,003
DD	Eldu Pont de Nemours and Co	Y	DD	2700	08/03/2007	20/07/2017	0,007	0,003
DIS	Disney	Y	DIS	2700	07/03/2007	21/07/2017	0,007	0,003
GE	General Electric	Y	GE	2700	12/03/2007	18/07/2017	0,014	0,005
GS	Goldman Sachs	Y	GS	2700	13/03/2007	17/07/2017	0,007	0,003
HD	Home Depot	Y	HD	2700	14/03/2007	14/07/2017	0,007	0,003
IBM	IBM	Y	IBM	2700	15/03/2007	13/07/2017	0,014	0,005
INTC	Intel	Y	INTC	2700	16/03/2007	12/07/2017	0,014	0,005
JNJ	Johnson & Johnson	Y	JNJ	2700	19/03/2007	11/07/2017	0,014	0,005
JPM	JPMorgan Chase	Y	JPM	2700	20/03/2007	10/07/2017	0,014	0,005
KO	Coca-Cola	Y	KO	2700	06/03/2007	24/07/2017	0,007	0,003
MCD	McDonald's	Y	MCD	2700	21/03/2007	07/07/2017	0,014	0,005
MMM	3 M	Y	MMM	2700	23/02/2007	02/08/2017	0,014	0,005
MRK	Merck	Y	MRK	2700	22/03/2007	06/07/2017	0,007	0,003
MSFT	Microsoft	Y	MSFT	2700	23/03/2007	05/07/2017	0,007	0,003
NKE	Nike	Y	NKE	2700	26/03/2007	03/07/2017	0,007	0,003
PFE	Pfizer	Y	PFE	2700	27/03/2007	30/06/2017	0,007	0,003
PG	Procter & Gamble	Y	PG	2700	28/03/2007	29/06/2017	0,007	0,003
TRV	Travelers Companies Inc	Y	TRV	2700	29/03/2007	28/06/2017	0,007	0,003
UNH	UnitedHealth	Y	UNH	2700	02/04/2007	26/06/2017	0,007	0,003
UTX	United Technologies	Y	UTX	2700	30/03/2007	27/06/2017	0,007	0,003
V	Visa	Y	V	2396	18/06/2008	22/06/2017	0,007	0,003
VZ	Verizon	Y	VZ	2700	03/04/2007	23/06/2017	0,007	0,003
WMT	Wal-Mart	Y	WMT	2700	05/04/2007	21/06/2017	0,007	0,003
XOM	Exxon Mobil	Y	XOM	2700	09/03/2007	19/07/2017	0,014	0,005
XRX	Xerox	Y	XRX	2700	09/04/2007	20/06/2017	0,007	0,003

(11)

Table 12: Hedge Fund

ASSET NAME	DESCRIPTION	SOURCE	CODE	SIZE	START	END	W1	W2
HFRX-ED	HFRX Event Driven Index	H	HFRX-ED	3623	10/07/2003	05/05/2017	0,007	0,004
HFRXED-DR	HFRX ED: Distressed Restructuring Index	H	HFRXED-DR	3623	03/07/2003	11/05/2017	0,007	0,004
HFRXED-MA	HFRX ED: Merger Arbitrage Index	H	HFRXED-MA	3623	07/07/2003	10/05/2017	0,007	0,004
HFRX-EH	HFRX Equity Hedge Index	H	HFRX-EH	3623	09/07/2003	08/05/2017	0,007	0,004
HFRXEHEM-N	HFRX EH: Equity Market Neutral Index	H	HFRXEHEM-N	3623	08/07/2003	09/05/2017	0,007	0,004
HFRX-Macro	HFRX Macro/CTA Index	H	HFRX-Macro	3623	11/07/2003	04/05/2017	0,007	0,004
HFRX-RVA	HFRX Relative Value Arbitrage Index	H	HFRX-RVA	3623	14/07/2003	03/05/2017	0,007	0,004
HFRXRV-FI	HFRX RV: FI-Convertible Arbitrage Index	H	HFRXRV-FI	3623	15/07/2003	02/05/2017	0,007	0,004

(12)

Table 13: Commodities

ASSET NAME	DESCRIPTION	SOURCE	CODE	SIZE	START	END	W1	W2
Brent	Brent Crude Oil Price USD/bbl.	O	BRENT	7685	02/09/1987	16/05/2017	0,007	0,023
Gold Bullion	Gold	O	GOLD	12554	03/05/1968	19/05/2017	0,007	0,023
WTI	WTI Crude Oil Price USD/bbl.	O	WTI	8256	18/04/1986	12/05/2017	0,007	0,023

(13)

Table 14: Exchange Rates

ASSET NAME	DESCRIPTION	SOURCE	CODE	SIZE	START	END	W1	W2
GBP/USD	UK-USA	Y	GBP=X	2800	19/04/2007	04/06/2017	0,007	0,023
USD/EURO	GERMANY-USA	Y	EURUSD=X	2800	17/04/2007	06/06/2017	0,007	0,023
USD/YEN	JAPAN-USA	Y	JPY=X	2800	18/04/2007	05/06/2017	0,007	0,023

(14)

In the empirical analysis, we consider only the last 2000 observations for each asset. The first half [1000 observations] was used for estimation purposes [in sample]. The last 1000 observations were used for testing purposes [out-of-sample]. For each of the 80 assets considered we employ VaR by the following six methods: i) **EVT**, ii) **FHS**, iii) **HOM**, iv) **ADP**, v) **ASY** and vi) **GAR**. The latter 3 correspond to the CAViaR alternatives. One-step ahead forecasts were obtained for the out-of-sample returns. We estimated VaR considering both $\alpha = 5\%$ [95% confidence level] and $\alpha = 1\%$ [99% confidence level]. Then, for each pair "*Asset-VaR Method*" we employ the following four backtests: i) **UC**, ii) **CC**, iii) **DQ** and iv) **VQR**. We also analyse the difference between the observed and expected frequencies of exceptions. Summarizing results are exhibited in the next subsection.

4.1 Results of the Empirical Investigation

The p-values obtained for each triple "*Asset-VaR Method-Backtest*" are shown both in Tables 30-32 [$\alpha = 5\%$] and Tables 31-33 [$\alpha = 1\%$]. For each VaR method, a success is a situation in which the null hypothesis is non rejected for some specific backtest. Obviously, success will depend simultaneously on: i) the VaR method adopted, ii) the specific backtest employed, iii) the specific asset considered and iv) the significance level chosen. At Tables 15-16 we present the success rate for each configuration. The success rate indicate the relative frequency of assets for which the method in question has not been rejected by the specific backtest [UC, CC, DQ or VQR], by some of these backtests [**ANY**] or jointly for all four backtests [**ALL**].

When $\alpha = 5\%$, EVT method exhibits the best success rates in almost cases [EVT is the best method in 14 different situations]. When CC and DQ backtests are employed, its success rates are higher than those observed for alternative methods, regardless of the significance level. When the significance level is 10%, the EVT outperform the alternative approaches for all backtests considered. In particular, the frequency of assets for which the EVT method

was not rejected according to the four backtests simultaneously is 35%, the largest one. Similar results occur when the significance level is 1% or 5%.

Table 15: Success Rates (%) by significance level [SL] when $\alpha = 5\%$

SL = 1%	ASY	GAR	SAV	EVT	FHS	HOM
UN	85.00	87.50	87.50	88.75	91.25	85.00
CC	81.25	85.00	83.75	85.00	85.00	81.25
DQ	82.50	83.75	77.50	87.50	87.50	85.00
VQR	73.75	68.75	80.00	75.00	71.25	70.00
ALL	63.75	60.00	67.50	67.50	63.75	65.00
ANY	92.50	92.50	92.50	95.00	96.25	91.25
SL = 5%	ASY	GAR	SAV	EVT	FHS	HOM
UN	61.25	71.25	65.00	72.50	68.75	70.00
CC	60.00	60.00	62.50	66.25	65.00	65.00
DQ	67.50	60.00	61.25	71.25	71.25	71.25
VQR	56.25	47.50	62.50	60.00	56.25	53.75
ALL	36.25	27.50	40.00	41.25	41.25	41.25
ANY	81.25	81.25	80.00	87.50	85.00	81.25
SL = 10%	ASY	GAR	SAV	EVT	FHS	HOM
UN	51.25	56.25	57.50	65.00	57.50	60.00
CC	43.75	53.75	48.75	57.50	55.00	56.25
DQ	57.50	46.25	46.25	61.25	56.25	60.00
VQR	38.75	38.75	45.00	53.75	51.25	46.25
ALL	25.00	21.25	27.50	35.00	33.75	26.25
ANY	71.25	73.75	68.75	80.00	73.75	77.50

(15)

Table 16: Success Rates (%) by significance level [SL] when $\alpha = 1\%$

SL = 1%	ASY	GAR	SAV	EVT	FHS	HOM
UN	91.25	96.25	88.75	97.50	97.50	96.25
CC	93.75	95.00	90.00	95.00	97.50	93.75
DQ	88.75	91.25	72.50	87.50	83.75	77.50
VQR	81.25	83.75	82.50	92.50	95.00	91.25
ALL	72.50	76.25	61.25	80.00	78.75	73.75
ANY	98.75	98.75	96.25	100.00	100.00	100.00
SL = 5%	ASY	GAR	SAV	EVT	FHS	HOM
UN	78.75	81.25	78.75	87.50	87.50	77.50
CC	83.75	83.75	75.00	87.50	85.00	75.00
DQ	78.75	75.00	62.50	75.00	76.25	62.50
VQR	71.25	67.50	72.50	77.50	82.50	81.25
ALL	56.25	50.00	47.50	56.25	60.00	51.25
ANY	97.50	96.25	96.25	98.75	98.75	95.00
SL = 10%	ASY	GAR	SAV	EVT	FHS	HOM
UN	68.75	72.50	62.50	81.25	83.75	71.25
CC	77.50	73.75	71.25	81.25	81.25	66.25
DQ	72.50	70.00	56.25	73.75	72.50	55.00
VQR	61.25	58.75	67.50	66.25	72.50	73.75
ALL	43.75	38.75	36.25	47.50	51.25	41.25
ANY	95.00	92.50	93.75	96.25	96.25	91.25

(16)

When VQR backtest is employed and the significance level is 1% or 5%, SAV outperforms EVT. However, the difference between their success rates is, at most, 5% [4 assets]. SAV is the best method in 3 different situations. Other methods are also best in some contexts: i) GAR [1 time], ii) HOM [2 times]

and, more importantly, iii) FHS [6 times]. FHS seems to be the second most important method. Its success rates are closer to those associated with the EVT method - the differences between them are less than 7.5% [6 assets]. However, excepting when the significance level is 1% and the UC backtest is employed, FHS success rates are less than or equal to those associated with EVT. It seems reasonable to state that EVT is the best method when $\alpha = 5\%$ regarding to the success rates.

The result is slightly different when $\alpha = 1\%$. If UC backtest is employed, the highest success rates are reached by the FHS method regardless of the significance level. At the same time, EVT method is the second best and the differences between their success rates are smaller than 2.5% [2 assets]. Relative to CC backtest, again the two best methods are EVT and FHS: their success rates are the largest and the difference between them is at most 2.5%. If the adopted backtest is DQ, then EVT, ASY and GAR are the best one methods when the significance levels are 10%, 5% and 1%, respectively. The differences between the success rates of EVT and the best method are smaller than 3.75% [3 assets]. The performance of EVT is worse when VQR backtest is employed. In this case, FHS and HOM outperforms their competitors²⁷. The difference between their success rates to those associated with EVT are, at most, 7.5% [6 assets]. EVT and FHS methods seems to be the best ones when $\alpha = 1\%$. When the significance level is 1%, for every asset there is a backtest for which those methods can not be rejected [it is also true when the method is HOM]. These rates decreases to 98.75% [96.25%] when the significance level is 5% [10%]. Taking into consideration the context in which all four backtests are simultaneously considered ["ALL" case], EVT and FHS also presents the largest success rates. If the significance level is 10% or 5%, then FHS outperforms EVT - the difference between their success rates is 3.75%. The reverse occurs if the significance level is 1%.

We also compared the performance of the alternative VaR methods with respect to the exception frequencies. For each asset [generically denoted by a], let i_a denote the observed frequency of exceptions. Then, we evaluated the Root Mean Square Error [RMSE] and Mean Absolute Error [MAE] both in weighted and unweighted version. In the **weighted approach** we use

$$RMSE = \sqrt{\sum_{a=1}^{n_a} w_a (i_a - \alpha)^2} \text{ and } MAE = \sum_{a=1}^{n_a} w_a |i_a - \alpha|, \quad (17)$$

where $w_1, \dots, w_{80}$ denote the weights of the n_a assets considered. For weighted metrics we consider the two distinct systems of weights [**W1** and **W2**] present

²⁷ Except when the significance level is 1% [EVT outperforms HOM].

at Tables 10-14. By taking $w_a \equiv \frac{1}{80}$ we obtain the **unweighted metrics** in a similar way:

$$RMSE = \sqrt{\frac{1}{n_a} \sum_{a=1}^{n_a} (i_a - \alpha)^2} \text{ and } MAE = \frac{1}{n_a} \sum_{a=1}^{n_a} |i_a - \alpha|. \quad (18)$$

The results are presented at Table 19.

In unweighted approach the smallest errors are associated with EVT method. If $\alpha = 5\%$, EVT method does not present the smallest errors only when we use MAE with weights system W2 [HOM is the best method]. Similarly, if $\alpha = 1\%$, EVT method does not present the smallest errors only when we use MAE with weights system W1 [FHS is the best method]. We must stress the fact that the, exception for $\alpha = 5\%$ in unweighted versions, the differences between errors presented by EVT and FHS methods are smaller than 2×10^{-4} . Once again, in almost all cases they outperform their competitors [and EVT is subtly better than FHS].

Table 19: RMSE and MAE by VaR method

alfa = 5%	ASY	GAR	SAV	EVT	FHS	HOM
MAE_0	0.0125	0.0117	0.0110	0.0101	0.0106	0.0116
RMSE_0	0.0155	0.0148	0.0132	0.0121	0.0126	0.0143
MAE_1	0.0125	0.0122	0.0114	0.0106	0.0108	0.0111
RMSE_1	0.0150	0.0145	0.0133	0.0124	0.0126	0.0137
MAE_2	0.0135	0.0123	0.0120	0.0107	0.0108	0.0102
RMSE_2	0.0165	0.0141	0.0138	0.0125	0.0126	0.0130
alfa = 1%	ASY	GAR	SAV	EVT	FHS	HOM
MAE_0	0.0042	0.0041	0.0047	0.0030	0.0030	0.0044
RMSE_0	0.0053	0.0049	0.0060	0.0038	0.0039	0.0054
MAE_1	0.0047	0.0044	0.0053	0.0032	0.0031	0.0046
RMSE_1	0.0057	0.0052	0.0068	0.0041	0.0041	0.0056
MAE_2	0.0053	0.0045	0.0056	0.0030	0.0032	0.0048
RMSE_2	0.0065	0.0054	0.0072	0.0040	0.0041	0.0057

W0 indicates unweighted errors while W1 and W2 indicate the weighted errors where the weight systems are respectively W1 and W2.

(19)

Table 20 shows the frequencies of assets for which at least one method was not rejected [acceptance percentage]. As before, the column labeled "ANY" is associated with assets for which there is some method not rejected according to at least one backtest. Similarly, the column labeled "ALL" is associated with assets for which there is some method not rejected simultaneously when we use **all** four backtests [UC, CC, DQ and VQR].

Table 20: Acceptance Percentage (%) by backtest

alfa = 5%	UN	CC	DQ	VQR	ALL	ANY
SL = 1%	96.25	93.75	96.25	96.25	88.75	98.75
SL = 5%	92.50	86.25	88.75	85.00	72.50	96.25
SL = 10%	85.00	81.25	83.75	77.50	60.00	95.00
alfa = 1%	UN	CC	DQ	VQR	ALL	ANY
SL = 1%	98.75	98.75	96.25	98.75	98.75	100.00
SL = 5%	93.75	96.25	90.00	97.50	97.50	100.00
SL = 10%	93.75	91.25	88.75	93.75	93.75	100.00

(20)

If $\alpha = 1\%$, for every asset considered there is a VaR method [ASY, GAR, SAV, EVT, FHS or HOM] that can not be rejected by some test [acceptance percentages are 100% in column "ANY"] independently from the significance level adopted. Otherwise, when $\alpha = 5\%$, if the significance level adopted is 10%, there are 4 assets [5%] for which there is no satisfactory VaR method regardless of the backtest employed. Analysing the acceptance percentages associated with column "ALL" we observe that when $\alpha = 1\%$ the group of VaR methods mentioned is a satisfactory one in a sense that for almost all assets there is at least one method that can not be rejected. The group of VaR methods does not works only for 5 [respectively, 2 and 1] assets when the significance level is 10% [respectively, 5% and 1%]. Otherwise, when $\alpha = 5\%$, at least when the significance level is 10%, the group of assets for which any VaR method is satisfactory has a considerably smaller size - no method works for 32 [40%] assets. Other VaR methods other than those considered herein must be employed for these assets.

Comparing the acceptance percentages at Table 20 with the success rates at Tables 16-15, we conclude that the set of assets for which each method works vary substantially. In particular, if the significance level is 10% and the four backtests are considered simultaneously, while the most successful methods presents success rates of 35% and 51.25% [for $\alpha = 5\%$ and $\alpha = 1\%$, respectively], the acceptance percentages are 60% and 93.75% in similar contexts. We report at Tables 21 and 22 the success rates for all pairs of VaR methods. The percentage indicated in any cell indicates the frequency of assets for which at least one of the two methods associated [related to the cell column and row] can not be rejected, regardless of the backtest used. The values that appear on the diagonals serve as benchmark [on diagonals, success does not refer to pairs of methods]²⁸.

It is interesting to note that the most successful pairs of VaR methods vary substantially. Both pairs (HOM,SAV) and (ASY,FHS) are the best twice each.

²⁸ Instead, these values correspond to success rates for individual methods that appear in the "ALL" row in Tables 15-16].

Both pairs (EVT,ASY) and (GAR,FHS) are the best only once each. Repair that EVT appears only once. However, by keeping one of the methods fixed as EVT itself, the best pairs have relatively high success rates. Except for the case where $\alpha = 5\%$ and the significance level is 5%, the difference of their success rates with respect to the larger ones are smaller than 2.5% [2 assets].

Table 21: Success Rates (%) for pairs of VaR methods when $\alpha = 5\%$

SL = 1%	ASY	GAR	SAV	EVT	FHS	HOM
ASY	63.75					
GAR	76.25	60.00				
SAV	78.75	73.75	67.50			
EVT	78.75	77.50	82.50	67.50		
FHS	76.25	73.75	80.00	70.00	63.75	
HOM	80.00	78.75	83.75	73.75	72.50	65.00
SL = 5%	ASY	GAR	SAV	EVT	FHS	HOM
ASY	36.25					
GAR	46.25	27.50				
SAV	53.75	43.75	40.00			
EVT	55.00	47.50	56.25	41.25		
FHS	53.75	48.75	57.50	43.75	41.25	
HOM	57.50	55.00	61.25	56.25	56.25	41.25
SL = 10%	ASY	GAR	SAV	EVT	FHS	HOM
ASY	25.00					
GAR	36.25	21.25				
SAV	40.00	32.50	27.50			
EVT	47.50	42.50	45.00	35.00		
FHS	46.25	42.50	43.75	35.00	33.75	
HOM	42.50	40.00	45.00	45.00	45.00	26.25

(21)

Table 22: Success Rates (%) for pairs of VaR methods when $\alpha = 1\%$

SL = 1%	ASY	GAR	SAV	EVT	FHS	HOM
ASY	72.50					
GAR	86.25	76.25				
SAV	83.75	82.50	61.25			
EVT	86.25	87.50	83.75	80.00		
FHS	88.75	90.00	85.00	82.50	78.75	
HOM	87.50	86.25	81.25	87.50	85.00	73.75
SL = 5%	ASY	GAR	SAV	EVT	FHS	HOM
ASY	56.25					
GAR	71.25	50.00				
SAV	68.75	61.25	47.50			
EVT	76.25	63.75	65.00	56.25		
FHS	77.50	70.00	68.75	65.00	60.00	
HOM	75.00	66.25	65.00	65.00	66.25	51.25
SL = 10%	ASY	GAR	SAV	EVT	FHS	HOM
ASY	43.75					
GAR	60.00	38.75				
SAV	56.25	47.50	36.25			
EVT	65.00	56.25	56.25	47.50		
FHS	67.50	60.00	58.75	56.25	51.25	
HOM	63.75	56.25	52.50	58.75	58.75	41.25

(22)

Another important result is the fact that in all situations considered there is

a CAViaR method such that the pair²⁹ (CAViaR,EVT) is the best between those pairs where one of the methods is EVT. This suggests that although individual FHS and HOM success rates are, in some cases, higher than those associated with CAViaR approaches, CAViaR performs best when EVT does not work well. In another words, CAViaR serve as a best complementary alternative to EVT, while FHS and HOM usually are best alternatives for assets in which EVT is not rejected.

5 *Simulation Study*

By alternating the set of assets or the period of time analyzed it would be possible to draw different conclusions in empirical comparison. Obviously, a natural strategy to evaluate the robustness of the conclusion obtained in a comparative study is to consider a wide variety of time series and different time periods. Considering that in the empirical analysis the set of assets adopted is large relatively to the empirical literature, we decide also to analyse simulated data to deal with question concerning dataset variability. This strategy is not uncommon. [LimNer207] compared the performances of classical VaR methods using artificial data. [GagAl211] compared backtest procedures using simulated data. Simulations are constructed by adopting some specific Data Generating Process [DGP]. At aforementioned works, as usual, different DGPs are considered, each of them corresponding to parametric models. However, here we adopt a different approach. Similar to [ShaAl197], our simulations are generated from a distribution estimated by nonparametric methods as we explain further. The two main reasons for this choice are: (i) the generated synthetic time-series should to have similar characteristic to those observed for real data and (ii) we did not want impose any specific patterns [specific parametric distributions or specific dynamic equations for time-varying parameters] in our simulation process.

The objective of this section is to compare the aforementioned *VaR* methods [EVT, FHS, HOM and CAViaR (ASY, GAR and SAV)] with respect to same backtests considered in empirical analysis [UC, CC, DQ or VQR], in a similar way to what was done in section 5. However, the main difference is the fact that in this section we consider simulated [artificial] returns data sets. For each of the 80 assets considered in the empirical analysis [see Tables 10-14], we associate a DGP identified by its own name. The DGP generically labeled "*a*"

²⁹ In the pair (CAViaR,EVT), CAViaR indicates a generic CAViaR method: ASY, GAR or SAV.

is identified by the [estimated] conditional distribution of the instantaneous return of the asset "a" with respect to its "p" most recent lagged returns. The nonparametric kernel method [see [WanJon195] and [Sil186]] was adopted for estimate the associated conditional cumulative distribution function [CCDF]. Then, using the inversion method, similar to that mentioned in [Dev186], it is possible simulate from the DGP. It is sufficient to draw a number u of a standard uniform distribution and obtain the u -quantile of the estimated distribution. For each of 80 DGPs, we simulate 100 artificial returns time series of size 2000. In the sequence we briefly describes our simulation strategy.

Consider, generically, a DGP labeled "a". Suppose $\{R_t\}_{t \in \mathbb{Z}}$ describe the stochastic process associated with the returns time series $\{r_t\}_{t=1}^T$ observed for the asset "a". As pointed by [ShaAl197], when simulating a time series, it is common assume that $\{R_t\}_{t \in \mathbb{Z}}$ is both stationary and markovian. In this sense, to simulate the corresponding instantaneous return R_t it is only necessary to take account its dependency with respect to the p most recent lagged terms $R_{t-1}, \dots, R_{t-p}$. Assuming that the process is strictly stationary, we could simulate values for the related time series at t time using the CCDF, given [for each $r^* \in \mathbb{R}$ and $\mathbf{r}^* = (r_1^*, \dots, r_p^*) \in \mathbb{R}^p$] by

$$F_C(r^*|\mathbf{r}^*) \equiv F_{R_t|R_{t-1}, \dots, R_{t-p}}(r^*|r_1^*, \dots, r_p^*). \quad (23)$$

Obviously, we rarely have access to the true function F_C . However, we can replace the unknown CCDF with some estimate of it. Similar to [LiRac208], we opt to estimate F_C by Nonparametric Kernel method. The estimated CCDF, $\hat{F}_C$, is defined, for each $(r^*, \mathbf{r}^*) \in \mathbb{R}^{p+1}$, by

$$\hat{F}_C(r^*|\mathbf{r}^*) \equiv \frac{T^{-1} \sum_{i=p+1}^T \Psi\left(\frac{r^* - r_i}{h_0}\right) M_\psi(\mathbf{r}^*; r_{i-1}, \dots, r_{i-p})}{T^{-1} \sum_{\ell=p+1}^T M_\psi(\mathbf{r}^*; r_{\ell-1}, \dots, r_{\ell-p})}, \quad (24)$$

where: (i) $M_\psi(\mathbf{r}^*; r_{i-1}, \dots, r_{i-p}) = \prod_{j=1}^p \frac{1}{h_j} \psi\left(\frac{r_j^* - r_{i-j}}{h_j}\right)$; (ii) ψ is an univariate kernel whose primitive is Ψ ; (iii) constants $h_0, h_1, \dots$ and h_p are the bandwidth parameters.

According to [Sil186], an univariate kernel function ψ is any function satisfying $\int_{-\infty}^{+\infty} \psi(s) ds = 1$ and, in particular, any density function is a kernel.³⁰ We use the gaussian kernel: $\psi = \phi$.³¹ There are several methods and rules for choosing the values of the bandwidth parameters. Most of them are based on

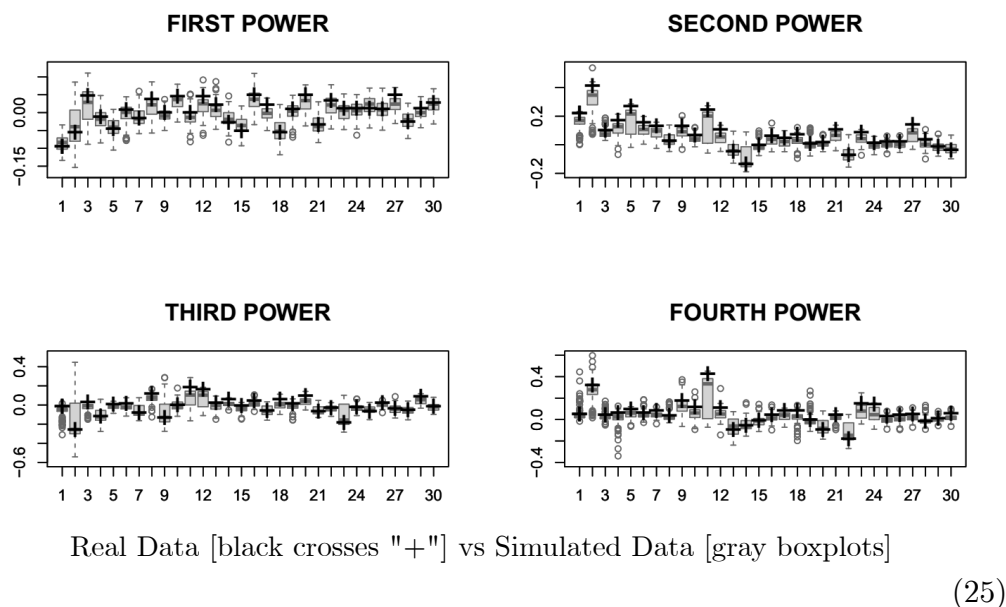
³⁰ [Sil186] presents some popular kernels like *Epanechnikov*, *Triangular*, *Uniform*, among others.

³¹ ϕ stand for the density function of the standard normal distribution

cross-validation concept. [LiAl213], for example, presents an selection strategy focused on the case which data are realizations of independent random variables. [CaiXu208], on the other hand, considered the case of weakly dependent data. In our exercises, we adopt the normal reference rule - see [Sil186]. Finally, although different ideal values for p may be associated with different assets, we chose $p = 10$ in all cases.

To illustrate the fact that the generated scenarios seem reasonable, we compared estimates of some statistics evaluated in the original dataset with those obtained for the artificial data. For each estimate associated with the actual data, we obtained 100 estimates associated with the 100 artificial data sets. The DGP chosen for illustrative comparisons was Dow Jones. In Figure 25 we display the PACFs for the first four powers of the Dow Jones return time series data [real data]. The estimates can be compared with the distributions of the estimates obtained for the artificial data.

Figure 25: PACF for the first four powers - Dow Jones index returns

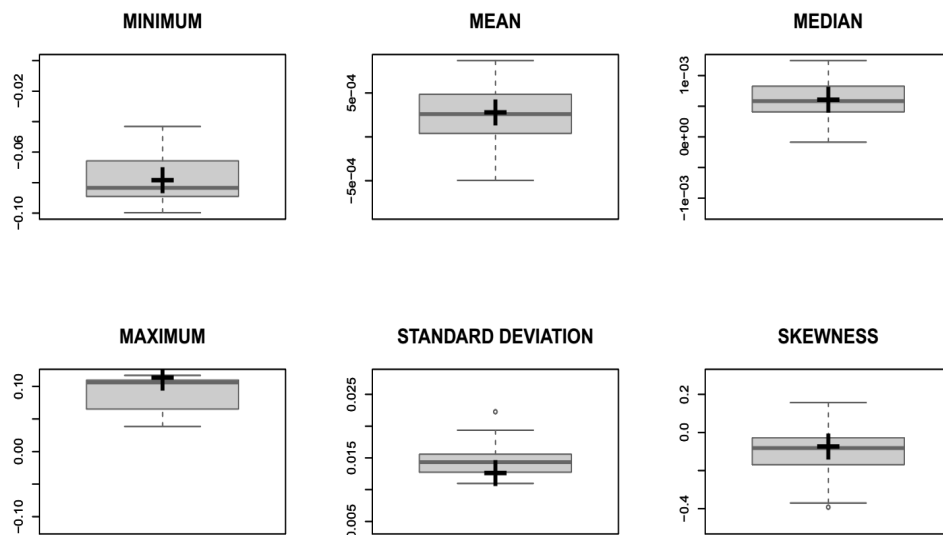


We observed that the estimates of the PACF coefficients obtained for the real time series [in its first four powers] are matching with the distribution of those estimates obtained for the associated coefficients when the artificial data are employed. The evolution of PACF coefficients estimates obtained for the real data when we move from lag 1 to lag 30 is similar to that observed for the distribution's median of the estimates obtained for the artificial data. In particular, the values of aforementioned medians and the values of the

coefficients associated with the real data are very close to each other [in almost cases, they are practically identical].

Similarly, we observe a good adjustment between estimates of important unconditional statistics in Figure 26. Both central tendencies [mean and median] and extreme values [maximum and minimum] are extremely compatible: the median of the distribution of estimates obtained for the artificial data are practically equal to the estimates obtained for the real data. The same occurs with the skewness coefficient. The largest difference is observed for the standard deviation. The estimate based on the real data [0.012] is closer to the first quartile [0.013] than the median [0.014] [with respect to the distribution of the estimates associated with the artificial data sets]. However, the differences for the 5th and 95th percentiles are less than 10^{-2} , indicating that the values of the standard deviation obtained for the simulated data are reasonable.

Figure 25: Unconditional Summary Statistics - Dow Jones index returns



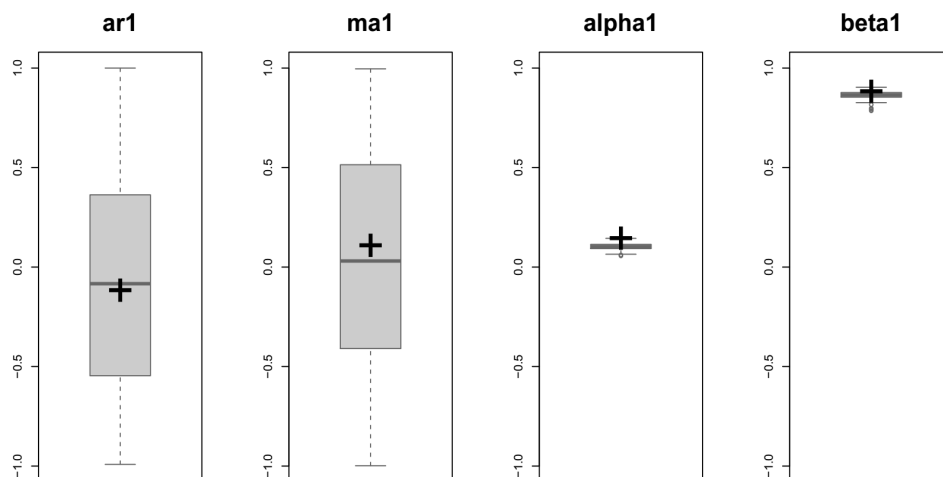
Real Data [black crosses "+"] vs Simulated Data [gray boxplots]

(26)

Finally, we also compare the estimates obtained for coefficients of models ARMA(1,1) [ar(1) and ma(1)] and GARCH(1,1) [alpha(1) and beta(1)] obtained both for real and simulated data sets - see Figure 27. Again, the estimates obtained for the real data seems to be compatible with those obtained for artificial data sets [in particular, with their median values].

The results exhibited for Dow Jones case are merely illustrative. However, for most DGPs or assets, we observed similar proximity between characteristics of both real and simulated data. We stress that our objective is only to produce reasonable scenarios - not necessarily to reproduce perfectly the DGP associated with the assets considered. An advantage of our approach is that we are not assuming parametric DGPs that could be extremely simplorious or not adequated for the real data. More still, in thesis, some parametric DGPs could induces biases in comparisons [in a sense that some methods are constructed in a such way that could performs better or worst depending the structure of the theoretical assumptions of the assumed DGPs]. In these view, we are not assuming very restrictive assumptions about DGPs.

Figure 27: ARMA(1,1) and GARCH(1,1) coefficients - Dow Jones index returns



Real Data [black crosses "+"] vs Simulated Data [gray boxplots]

(27)

In summary, for each of the 80 DGPs [corresponding to the same assets considered in the empirical analysis] we generate 100 artificial time series of return data [100 different scenarios: scen1, $\dots$, scen100]. Each artificial time series contains 2000 data points that we split in "in sample" [first half] and "out-of-sample" [second half] sets. Similar analysis to those presented in empirical analysis were done with all artificial time series. We employ the six VaR Methods considered before [EVT, FHS, HOM, ASY, GAR and SAV] to obtain one-ahead-forecast for the "out-of-sample" set using estimates based on "in sample" set. Next, the predicted and observed values for the "out-of-sample" set were used for testing purposes [using UC, CC, DQ and VQR backtests]. Once again, we adopted $\alpha = 5\%$ and $\alpha = 1\%$ and considered the adjustment between observed and expected exceptions, success rates, acceptance percentages, similar to what was done in section 5. In the next subsection, we present the main results obtained - a lot of them are summarized by DGPs.

5.1 Results of the Simulation Study [PENDENTE]

6. Conclusion [PENDENTE]

7. APPENDIX

Table 28: p-values for *companies stocks* [$\alpha = 5\%$]

ASSETS\TESTS	ASY				GAR				SAV				EVT				FHS				HOM			
	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR
AAPL	0.180	0.390	0.400	0.310	0.660	0.790	0.930	0.500	0.770	0.400	0.800	0.940	0.370	0.500	0.940	0.490	0.460	0.760	0.980	0.630	0.010	0.010	0.180	0.000
AXP	0.130	0.190	0.340	0.270	0.030	0.090	0.520	0.020	0.230	0.330	0.730	0.440	0.670	0.410	0.060	0.000	1.000	0.650	0.060	0.000	0.180	0.250	0.300	0.000
BA	0.030	0.090	0.350	0.000	0.050	0.120	0.380	0.010	0.010	0.020	0.170	0.020	0.100	0.140	0.230	0.120	0.100	0.140	0.230	0.140	0.010	0.020	0.180	0.150
CAT	0.010	0.050	0.500	0.000	0.070	0.180	0.610	0.000	0.030	0.100	0.580	0.010	0.130	0.280	0.700	0.000	0.030	0.100	0.440	0.000	0.050	0.130	0.340	0.000
CSCO	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
CVX	0.130	0.190	0.560	0.040	0.880	0.290	0.500	0.980	0.770	0.250	0.580	0.800	0.300	0.580	0.760	0.790	0.560	0.840	0.870	0.900	0.180	0.340	0.550	0.370
DD	0.560	0.170	0.250	0.620	0.230	0.150	0.290	0.350	0.560	0.170	0.270	0.800	0.050	0.120	0.220	0.030	0.050	0.120	0.220	0.020	0.010	0.010	0.060	0.000
DIS	0.010	0.050	0.210	0.000	0.300	0.450	0.220	0.290	0.180	0.340	0.300	0.070	0.010	0.020	0.120	0.000	0.010	0.020	0.120	0.000	0.000	0.000	0.010	0.000
GE	0.020	0.060	0.020	0.000	0.010	0.010	0.010	0.000	0.010	0.000	0.000	0.000	0.000	0.000	0.010	0.000	0.000	0.000	0.010	0.000	0.000	0.000	0.030	0.000
GS	0.300	0.580	0.560	0.700	0.560	0.840	0.970	0.460	0.770	0.940	0.560	0.990	0.130	0.310	0.650	0.010	0.130	0.310	0.650	0.030	0.010	0.050	0.280	0.000
HD	0.070	0.010	0.040	0.100	0.230	0.480	0.640	0.220	0.460	0.610	0.930	0.840	0.370	0.510	0.860	0.110	0.370	0.510	0.860	0.230	0.050	0.120	0.650	0.040
IBM	0.120	0.300	0.100	0.250	0.160	0.360	0.060	0.000	0.090	0.220	0.030	0.040	0.390	0.690	0.440	0.140	0.390	0.690	0.440	0.140	0.130	0.190	0.390	0.040
INTC	0.300	0.420	0.670	0.300	0.070	0.170	0.500	0.000	0.560	0.400	0.860	0.330	0.100	0.140	0.490	0.040	0.050	0.060	0.310	0.020	0.050	0.120	0.450	0.020
JNJ	0.880	0.960	0.100	0.150	0.200	0.090	0.020	0.620	0.160	0.280	0.000	0.610	0.670	0.900	0.280	0.160	0.380	0.690	0.120	0.170	0.670	0.900	0.080	0.150
JPM	0.020	0.000	0.000	0.060	0.030	0.000	0.030	0.070	0.020	0.000	0.000	0.070	0.100	0.000	0.000	0.030	0.050	0.000	0.000	0.010	0.010	0.000	0.000	0.000
KO	0.230	0.150	0.180	0.050	0.770	0.540	0.230	0.280	0.460	0.330	0.180	0.420	0.560	0.400	0.170	0.280	0.560	0.400	0.170	0.270	0.180	0.250	0.510	0.090
MCD	0.180	0.040	0.200	0.350	0.100	0.020	0.040	0.050	0.180	0.010	0.050	0.410	0.100	0.000	0.010	0.020	0.070	0.000	0.010	0.010	0.010	0.000	0.000	0.000
MMM	0.370	0.500	0.350	0.680	0.050	0.030	0.150	0.090	0.230	0.400	0.050	0.400	0.020	0.020	0.060	0.070	0.020	0.020	0.060	0.030	0.010	0.030	0.040	0.010
MRK	0.770	0.440	0.460	0.870	0.560	0.840	0.880	0.930	0.320	0.450	0.160	0.650	0.770	0.940	0.560	0.880	0.770	0.940	0.560	0.860	0.030	0.100	0.370	0.290
MSFT	0.300	0.450	0.640	0.060	0.070	0.180	0.460	0.010	0.180	0.340	0.710	0.090	0.070	0.100	0.590	0.320	0.050	0.060	0.510	0.140	0.020	0.030	0.270	0.020
NKE	0.000	0.000	0.080	0.000	0.050	0.020	0.120	0.020	0.030	0.010	0.080	0.020	0.030	0.040	0.150	0.000	0.010	0.010	0.080	0.000	0.000	0.000	0.000	0.000
PFE	0.040	0.010	0.000	0.020	0.320	0.210	0.000	0.010	0.040	0.020	0.000	0.020	0.670	0.190	0.010	0.160	0.670	0.190	0.010	0.170	0.200	0.010	0.000	0.000
PG	0.300	0.200	0.020	0.000	0.130	0.190	0.080	0.460	0.180	0.390	0.430	0.170	0.300	0.580	0.640	0.200	0.180	0.390	0.450	0.140	0.030	0.090	0.310	0.150
TRV	0.020	0.010	0.000	0.130	0.130	0.020	0.040	0.180	0.230	0.050	0.090	0.270	0.030	0.000	0.000	0.000	0.030	0.000	0.000	0.000	0.000	0.000	0.000	0.000
UNH	0.070	0.010	0.070	0.100	0.070	0.010	0.070	0.050	0.070	0.010	0.060	0.980	0.000	0.000	0.020	0.000	0.010	0.000	0.030	0.000	0.000	0.000	0.010	0.000
UTX	0.130	0.280	0.220	0.010	0.230	0.400	0.460	0.020	0.180	0.390	0.240	0.220	0.370	0.670	0.580	0.270	0.560	0.700	0.460	0.370	0.130	0.280	0.490	0.050
V	0.020	0.060	0.320	0.230	0.230	0.460	0.600	0.320	0.300	0.580	0.590	0.510	0.300	0.420	0.530	0.120	0.300	0.420	0.530	0.240	0.000	0.000	0.030	0.010
VZ	0.660	0.020	0.000	0.000	0.000	0.000	0.000	0.000	0.560	0.050	0.130	0.630	0.300	0.200	0.190	0.660	0.300	0.200	0.190	0.740	0.130	0.020	0.030	0.090
WMT	0.200	0.290	0.390	0.460	0.390	0.530	0.660	0.330	0.320	0.450	0.580	0.860	0.770	0.860	0.950	0.660	0.770	0.860	0.950	0.680	0.370	0.670	0.960	0.190
XOM	0.660	0.090	0.350	0.610	0.770	0.860	0.820	0.500	1.000	0.950	0.960	0.990	1.000	0.950	0.950	0.540	0.460	0.610	0.880	0.450	0.770	0.860	0.940	0.470
XRX	0.100	0.020	0.060	0.160	0.010	0.000	0.040	0.000	0.000	0.000	0.020	0.020	0.130	0.020	0.040	0.000	0.010	0.010	0.060	0.000	0.070	0.040	0.040	0.000

(28)

Table 29: p-values for *companies stocks* [$\alpha = 1\%$]

ASSETS\TESTS	ASY				GAR				SAV				EVT				FHS				HOM				
	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	
AAPL	0.360	0.560	0.930	0.580	0.310	0.570	0.980	0.000	0.230	0.400	0.590	0.740	1.000	0.900	0.960	0.360	1.000	0.900	0.960	0.360	1.000	0.900	0.990	0.770	
AXP	0.170	0.380	0.840	0.000	0.750	0.870	0.990	0.310	0.510	0.750	0.960	0.230	0.510	0.750	0.970	0.150	1.000	0.900	0.990	0.270	0.540	0.710	0.950	0.300	
BA	0.750	0.840	0.930	0.870	0.140	0.270	0.230	0.730	1.000	0.900	0.990	0.740	0.540	0.710	0.610	0.890	0.140	0.270	0.210	0.710	0.140	0.270	0.180	0.370	
CAT	1.000	0.900	0.970	0.420	0.360	0.560	0.910	0.860	0.360	0.560	0.900	0.780	0.230	0.400	0.690	0.810	0.230	0.400	0.690	0.810	0.140	0.270	0.410	0.240	
CSCO	0.510	0.750	0.950	0.240	0.170	0.380	0.940	0.170	0.000	0.000	0.220	0.000	0.170	0.380	0.960	0.010	0.170	0.380	0.900	0.010	0.310	0.570	0.980	0.210	
CVX	0.750	0.840	0.810	0.030	0.510	0.750	1.000	0.810	0.750	0.840	0.940	0.330	1.000	0.900	0.910	0.100	0.540	0.710	0.840	0.220	0.230	0.490	0.420	0.140	
DD	1.000	0.900	0.800	0.670	0.080	0.110	0.050	0.430	0.020	0.050	0.000	0.340	0.750	0.870	0.990	0.350	1.000	0.900	1.000	0.420	1.000	0.230	0.150	0.840	
DIS	0.310	0.570	0.980	0.300	0.310	0.570	0.980	0.890	0.170	0.380	0.940	0.390	0.310	0.570	0.950	0.540	0.540	0.710	0.840	0.580	0.750	0.870	1.000	0.940	
GE	0.010	0.030	0.510	0.000	0.030	0.090	0.700	0.000	0.010	0.030	0.500	0.000	0.030	0.090	0.670	0.000	0.030	0.090	0.670	0.030	0.310	0.070	0.010	0.110	
GS	0.310	0.570	0.980	0.300	0.170	0.380	0.980	0.010	0.510	0.750	0.970	0.020	0.030	0.090	0.690	0.020	0.170	0.380	0.940	0.290	0.170	0.030	0.010	0.450	
HD	0.310	0.570	0.980	0.470	0.170	0.380	0.940	0.060	0.170	0.380	0.950	0.160	0.310	0.570	0.920	0.160	0.080	0.210	0.820	0.060	0.750	0.840	0.150	0.490	
IBM	0.140	0.270	0.650	0.850	0.230	0.400	0.570	0.020	0.140	0.270	0.650	0.840	0.230	0.400	0.770	0.230	0.230	0.400	0.770	0.260	0.080	0.170	0.460	0.250	
INTC	0.230	0.400	0.830	0.540	0.360	0.560	0.930	0.740	0.230	0.400	0.700	0.760	0.540	0.710	0.980	0.790	0.540	0.710	0.980	0.750	0.140	0.270	0.670	0.490	
JNJ	0.000	0.000	0.000	0.110	0.020	0.050	0.020	0.000	0.010	0.010	0.000	0.120	0.000	0.000	0.000	0.000	0.010	0.010	0.030	0.000	0.010	0.000	0.000	0.010	
JPM	0.000	0.000	0.230	0.000	0.030	0.090	0.730	0.010	0.030	0.090	0.690	0.030	0.080	0.010	0.000	0.010	1.000	0.230	0.000	0.220	0.510	0.120	0.050	0.280	
KO	0.360	0.560	0.800	0.410	0.140	0.160	0.160	0.640	0.080	0.110	0.120	0.230	0.080	0.110	0.120	0.030	0.080	0.110	0.120	0.040	0.940	0.070	0.970	0.060	
MCD	0.170	0.380	0.750	0.050	0.310	0.070	0.020	0.000	0.540	0.260	0.010	0.090	0.310	0.070	0.020	0.070	0.750	0.170	0.990	0.140	0.940	0.060	0.000	0.000	
MRK	0.310	0.570	0.950	0.360	0.310	0.570	0.960	0.710	0.310	0.570	0.970	0.730	0.510	0.750	0.960	0.610	0.510	0.750	0.960	0.610	0.750	0.840	0.320	0.260	
MSFT	0.310	0.570	0.930	0.190	0.080	0.210	0.000	0.060	0.080	0.210	0.380	0.120	0.170	0.380	0.310	0.030	0.450	0.750	0.300	0.450	0.750	0.090	0.000	0.000	
NFT	0.750	0.870	0.920	0.000	0.640	0.360	0.560	0.690	0.000	0.750	0.870	1.000	0.190	0.750	0.870	0.960	0.020	1.000	0.900	0.950	0.030	0.230	0.490	0.690	0.010
OKS	0.080	0.210	0.430	0.000	0.170	0.030	0.010	0.350	0.750	0.170	0.120	0.550	0.080	0.210	0.460	0.180	0.750	0.170	0.110	0.870	0.750	0.170	0.100	0.730	
PG	0.230	0.000	0.000	0.290	0.080	0.000	0.000	0.440	0.080	0.000	0.000	0.020	0.510	0.000	0.000	0.040	0.750	0.020	0.000	0.150	0.540	0.020	0.000	0.000	
PFE	0.360	0.560	0.010	0.990	1.000	0.900	1.000	0.670	0.230	0.400	0.000	0.020	0.750	0.260	0.240	0.990	1.000	0.230	0.190	1.000	0.230	0.400	0.290	0.220	
TRV	0.510	0.120	0.060	0.450	1.000	0.230	0.010	0.360	0.510	0.750	0.940	0.490	1.000	0.230	0.000	0.110	1.000	0.230	0.000	0.100	0.750	0.260	0.010	0.140	
UNH	0.080	0.210	0.750	0.000	0.030	0.090	0.690	0.020	0.510	0.120	0.030	0.230	0.030	0.090	0.720	0.010	0.030	0.090	0.720	0.010	0.800	0.210	0.760	0.020	
UTX	0.310	0.570	0.870	0.320	0.360	0.560	0.790	0.030	0.080	0.110	0.000	0.020	0.540	0.710	0.830	0.040	0.540	0.710	0.830	0.060	0.230	0.490	0.570	0.040	
V	0.750	0.870	0.800	0.000	0.750	0.260	0.250	0.990	0.360	0.240	0.290	0.430	0.510	0.750	0.940	0.240	1.000	0.900	1.000	0.290	1.000	0.900	0.920	0.240	
VZ	1.000	0.900	0.750	0.050	0.750	0.870	0.890	0.050	1.000	0.900	0.740	0.070	1.000	0.900	0.720	0.010	1.000	0.900	0.720	0.010	0.230	0.490	0.240	0.020	
WMT	0.010	0.010	0.000	0.000	0.140	0.270	0.420	0.000	0.040	0.100	0.060	0.000	0.080	0.170	0.170	0.010	0.080	0.170	0.170	0.010	0.010	0.010	0.000	0.000	
XOM	0.540	0.710	0.260	0.020	0.080	0.210	0.760	0.190	0.540	0.710	0.210	0.020	0.170	0.380	0.910	0.290	0.170	0.380	0.910	0.330	0.540	0.710	0.970	0.490	
YRD	0.510	0.120	0.030	0.000	0.510	0.750	0.980	0.470	0.510	0.750	1.000	0.760	0.510	0.750	0.960	0.000	1.000	0.610	0.000	0.000	0.750	0.020	0.000	0.000	

Table 30: p-values for *market or stock exchange indices* [$a = 5\%$]

ASSETS\TESTS	ASY				GAR				SAV				EVT				FHS				HOM			
	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR
AEX	0.010	0.020	0.120	0.060	0.010	0.010	0.090	0.010	0.030	0.090	0.300	0.050	0.010	0.020	0.240	0.030	0.010	0.020	0.240	0.000	0.100	0.240	0.320	0.240
AORD ALL	0.370	0.510	0.470	0.040	0.130	0.280	0.260	0.000	0.130	0.310	0.320	0.010	0.230	0.480	0.210	0.010	0.230	0.480	0.210	0.000	0.580	0.700	0.210	0.000
BUX	0.030	0.010	0.040	0.010	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.050	0.020	0.030	0.080	0.070	0.030	0.070	0.250	0.010	0.020	0.140	0.000
CAC40	0.050	0.120	0.110	0.090	0.010	0.030	0.150	0.000	0.010	0.050	0.050	0.000	0.000	0.010	0.180	0.020	0.010	0.010	0.190	0.060	0.050	0.120	0.380	0.210
COLCAP	0.880	0.910	0.340	0.170	0.460	0.040	0.070	0.740	1.000	0.040	0.030	1.000	0.770	0.090	0.160	0.430	0.770	0.090	0.160	0.440	0.370	0.030	0.070	0.330
CROBEX	0.020	0.030	0.040	0.000	0.000	0.000	0.000	0.000	0.180	0.250	0.140	0.240	0.010	0.020	0.000	0.000	0.010	0.020	0.010	0.000	0.000	0.000	0.000	0.000
DAX 30	0.020	0.060	0.310	0.000	0.010	0.040	0.280	0.030	0.370	0.510	0.390	0.310	0.130	0.020	0.080	0.120	0.130	0.020	0.080	0.110	0.370	0.100	0.290	0.790
DJ Euro Stoxx 50	0.050	0.000	0.000	0.160	0.010	0.000	0.000	0.000	0.030	0.000	0.000	0.010	0.030	0.000	0.000	0.030	0.030	0.000	0.000	0.010	0.070	0.000	0.000	0.110
Dow Jones	0.010	0.040	0.050	0.050	0.070	0.170	0.080	0.030	0.130	0.310	0.150	0.030	0.130	0.310	0.150	0.120	0.130	0.310	0.150	0.110	0.230	0.330	0.170	0.270
FTSE 100	0.300	0.420	0.370	0.090	0.130	0.310	0.260	0.340	0.230	0.480	0.650	0.070	0.100	0.140	0.310	0.280	0.070	0.170	0.330	0.190	0.560	0.400	0.250	0.690
FTSEMIB.MI	0.000	0.000	0.000	0.000	0.370	0.510	0.400	0.040	0.100	0.050	0.310	0.070	0.020	0.020	0.220	0.000	0.300	0.420	0.370	0.030	0.300	0.420	0.420	0.070
HANG SENG	0.230	0.080	0.610	0.470	0.070	0.040	0.440	0.060	0.180	0.070	0.410	0.370	0.370	0.500	0.750	0.390	0.300	0.450	0.670	0.260	0.300	0.450	0.670	0.480
IBEX	0.180	0.390	0.340	0.040	0.100	0.230	0.020	0.040	0.070	0.180	0.090	0.060	0.130	0.280	0.030	0.070	0.130	0.280	0.030	0.080	0.230	0.400	0.050	0.130
IBOV	0.570	0.850	0.190	0.200	0.320	0.560	0.300	0.380	0.390	0.610	0.190	0.480	0.390	0.690	0.260	0.140	0.320	0.600	0.230	0.130	0.670	0.790	0.210	0.160
IPSA	0.070	0.170	0.110	0.010	0.070	0.100	0.130	0.000	0.020	0.010	0.080	0.060	0.230	0.480	0.780	0.470	0.230	0.480	0.770	0.340	0.130	0.310	0.520	0.160
ISE	0.180	0.040	0.140	0.530	0.460	0.130	0.310	0.390	0.370	0.100	0.270	0.380	0.230	0.050	0.160	0.270	0.370	0.030	0.030	0.370	0.130	0.080	0.410	0.140
JALSH	0.000	0.000	0.000	0.000	0.010	0.020	0.010	0.000	0.000	0.000	0.000	0.000	0.050	0.150	0.310	0.150	0.000	0.000	0.000	0.030	0.200	0.010	0.210	0.180
JKSE	0.560	0.840	0.830	0.800	0.120	0.010	0.030	0.480	0.260	0.010	0.030	0.250	0.070	0.030	0.130	0.670	0.120	0.080	0.300	0.720	0.300	0.420	0.840	0.510
KOSPI	0.000	0.000	0.010	0.000	0.010	0.020	0.170	0.000	0.020	0.010	0.060	0.000	0.070	0.040	0.080	0.000	0.050	0.020	0.060	0.000	0.100	0.050	0.090	0.000
Malaysia	0.020	0.020	0.020	0.000	0.100	0.230	0.000	0.030	0.020	0.020	0.000	0.000	0.130	0.060	0.000	0.070	0.020	0.020	0.000	0.000	0.050	0.030	0.000	0.060
MERVAL	0.070	0.030	0.020	0.080	1.000	0.040	0.000	0.750	0.880	0.030	0.020	0.910	0.570	0.420	0.120	0.570	0.670	0.410	0.100	0.610	1.000	0.530	0.090	0.990
MEXBOL	0.010	0.000	0.030	0.000	0.460	0.040	0.050	0.610	0.050	0.020	0.110	0.070	0.130	0.190	0.570	0.230	0.300	0.420	0.620	0.410	0.130	0.190	0.570	0.190
Nasdaq	0.000	0.000	0.040	0.010	0.000	0.000	0.040	0.000	0.000	0.000	0.040	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000
Nikkei 225	0.070	0.190	0.010	0.360	0.120	0.300	0.000	0.380	0.020	0.060	0.000	0.060	0.160	0.360	0.020	0.220	0.260	0.510	0.060	0.250	0.320	0.600	0.030	0.400
NSEI	0.570	0.710	0.820	0.960	0.050	0.130	0.380	0.000	0.300	0.450	0.550	0.390	0.560	0.570	0.910	0.360	0.880	0.910	0.860	0.380	0.660	0.590	0.910	0.430
PSE	0.040	0.050	0.000	0.070	0.160	0.170	0.010	0.100	0.040	0.050	0.000	0.060	0.880	0.150	0.120	0.850	0.770	0.170	0.150	0.810	0.560	0.050	0.190	0.890
RTS	0.030	0.090	0.380	0.000	0.230	0.330	0.300	0.000	0.020	0.060	0.130	0.000	0.460	0.780	0.120	0.020	0.070	0.170	0.220	0.000	0.300	0.580	0.030	0.080
S&P 500	0.010	0.010	0.090	0.030	0.010	0.020	0.040	0.020	0.010	0.020	0.030	0.000	0.010	0.040	0.050	0.010	0.030	0.090	0.110	0.020	0.070	0.170	0.050	0.020
SET	0.050	0.030	0.340	0.070	0.670	0.790	0.010	0.820	0.670	0.790	0.020	0.610	0.470	0.280	0.010	0.240	0.470	0.280	0.010	0.230	0.230	0.400	0.010	0.300
SSE	0.770	0.860	0.130	0.090	1.000	0.330	0.030	0.020	0.880	0.600	0.010	0.070	0.180	0.250	0.040	0.000	0.770	0.540	0.020	0.000	0.180	0.250	0.040	0.000
STI	0.300	0.200	0.240	0.190	0.660	0.070	0.060	0.390	0.880	0.110	0.090	0.040	0.560	0.170	0.030	0.060	0.460	0.130	0.040	0.030	0.370	0.100	0.070	0.030
TSX	0.130	0.060	0.220	0.080	0.460	0.330	0.080	0.570	0.300	0.200	0.110	0.560	0.660	0.790	0.230	0.720	0.770	0.860	0.270	0.920	0.880	0.600	0.210	0.890
TWI or TSEC	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.010	0.020
WIG	0.660	0.790	0.290	0.040	0.660	0.470	0.280	0.000	0.670	0.710	0.160	0.000	1.000	0.330	0.180	0.000	1.000	0.330	0.180	0.000	1.000	0.330	0.160	0.000
XU100 or ISE	0.660	0.090	0.060	0.030	0.230	0.480	0.010	0.010	0.460	0.590	0.060	0.020	0.370	0.670	0.040	0.160	0.370	0.670	0.040	0.190	0.230	0.480	0.340	0.150

(30)

Table 31: p-values for *market or stock exchange indices* [$a = 1\%$]

ASSETS\TESTS	ASY				GAR				SAV				EVT				FHS				HOM				
	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	
AEX	0.230	0.400	0.790	0.290	0.540	0.710	0.980	0.730	0.230	0.400	0.810	0.450	0.540	0.710	0.980	0.380	0.230	0.200	0.290	0.190	0.140	0.160	0.250	0.090	
AORD ALL	0.750	0.840	1.000	0.880	0.540	0.710	0.990	0.870	0.750	0.840	1.000	0.810	1.000	0.900	1.000	0.850	0.040	0.100	0.410	0.340	0.020	0.050	0.030	0.280	
BUX	0.030	0.090	0.710	0.010	0.170	0.380	0.870	0.120	0.310	0.580	0.700	0.000	0.010	0.030	0.550	0.000	0.000	0.010	0.370	0.000	0.310	0.580	0.010	0.100	
CAC40	0.140	0.270	0.650	0.100	0.230	0.400	0.830	0.280	0.230	0.400	0.770	0.130	0.540	0.710	0.970	0.400	0.540	0.710	0.970	0.410	0.140	0.270	0.660	0.140	
COLCAP	0.310	0.570	0.940	0.280	1.000	0.230	0.090	0.630	1.000	0.230	0.080	0.320	0.750	0.870	1.000	0.860	0.750	0.870	1.000	0.870	0.080	0.170	0.170	0.520	
CROBEX	0.310	0.570	0.960	0.610	0.080	0.210	0.860	0.080	0.080	0.210	0.860	0.310	0.170	0.380	0.950	0.070	0.170	0.380	0.950	0.070	0.170	0.380	0.930	0.200	
DAX 30	0.540	0.710	0.980	0.390	0.360	0.560	0.980	0.760	0.230	0.400	0.240	0.350	0.750	0.870	0.990	0.710	1.000	0.900	0.990	0.610	0.230	0.400	0.280	0.760	
DJ Euro Stoxx 50	0.230	0.200	0.260	0.470	0.750	0.260	0.250	0.710	0.140	0.020	0.000	0.190	1.000	0.230	0.180	0.730	0.230	0.020	0.000	0.400	0.140	0.020	0.000	0.140	
Dow Jones	0.000	0.010	0.360	0.000	0.030	0.090	0.720	0.020	0.080	0.010	0.000	0.090	0.170	0.030	0.000	0.050	0.170	0.030	0.000	0.050	0.230	0.200	0.000	0.090	
FTSE 100	0.080	0.110	0.240	0.400	0.360	0.240	0.040	0.790	0.000	0.000	0.000	0.350	0.540	0.260	0.030	0.890	0.360	0.240	0.040	0.740	0.010	0.030	0.000	0.410	
FTSEMIB MI	0.080	0.120	0.860	0.020	0.510	0.750	1.000	0.840	0.310	0.570	0.980	0.410	0.510	0.750	1.000	0.730	0.750	0.870	1.000	0.720	0.750	0.840	0.990	0.900	
HANG SENG	0.310	0.570	0.980	0.620	0.310	0.570	0.030	0.770	0.510	0.750	0.960	0.820	0.540	0.710	0.030	0.950	0.750	0.870	1.000	0.970	0.230	0.400	0.060	0.320	
IBEX	1.000	0.990	0.910	0.400	0.360	0.560	0.300	0.930	0.540	0.710	0.310	0.930	0.510	0.750	1.000	0.920	0.750	0.870	1.000	1.000	0.360	0.560	0.320	0.910	
IBOV	0.510	0.750	0.990	0.000	0.510	0.750	0.870	0.020	0.750	0.870	0.970	0.230	0.310	0.570	0.980	0.310	0.310	0.570	0.980	0.340	1.000	0.900	0.960	0.720	
IPSA	0.750	0.870	0.880	0.010	0.610	0.030	0.550	0.010	0.010	0.030	0.540	0.000	0.170	0.380	0.920	0.330	0.510	0.750	0.980	0.540	0.230	0.400	0.280	0.340	
ISE	0.230	0.200	0.290	0.540	0.230	0.400	0.810	0.400	0.140	0.160	0.210	0.560	0.360	0.240	0.240	0.440	0.230	0.200	0.230	0.170	0.020	0.050	0.040	0.120	
JALSH	0.360	0.560	0.010	0.10	0.350	0.230	0.400	0.150	0.190	0.750	0.840	0.090	0.270	0.540	0.710	0.150	0.830	0.360	0.560	0.180	0.730	0.000	0.000	0.000	0.120
JKSE	0.310	0.570	0.940	0.060	0.310	0.570	0.980	0.500	0.080	0.210	0.820	0.050	0.310	0.570	0.980	0.080	0.510	0.750	0.980	0.230	0.510	0.750	0.980	0.690	
KDPSI	0.310	0.570	0.980	0.020	0.510	0.750	1.000	0.790	0.750	0.870	1.000	0.890	0.170	0.380	0.940	0.080	0.170	0.380	0.940	0.260	0.510	0.750	0.990	0.780	
Malaysia	0.310	0.570	0.880	0.020	0.510	0.750	0.960	0.870	0.510	0.750	0.930	0.840	0.510	0.750	1.000	0.790	0.750	0.840	0.990	0.910	0.360	0.560	0.020	0.710	
Merval	0.020	0.050	0.000	0.160	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.540	0.710	0.920	0.590	0.740	0.100	0.210	0.090	0.020	0.050	0.000	0.120	0.080
MEXBOL	0.310	0.570	0.960	0.040	0.080	0.210	0.850	0.030	0.310	0.570	0.940	0.440	0.170	0.380	0.950	0.270	0.170	0.380	0.950	0.140	1.000	0.900	0.990	0.990	
Nasdaq	0.000	0.000	0.120	0.000	0.000	0.000	0.120	0.000	0.000	0.000	0.120	0.000	0.000	0.000	0.120	0.000	0.000	0.000	0.120	0.000	0.000	0.000	0.120	0.000	0.120
Nikkei 225	0.010	0.030	0.020	0.080	0.010	0.010	0.010	0.010	0.000	0.000	0.000	0.000	0.010	0.030	0.030	0.120	0.010	0.030	0.030	0.120	0.010	0.030	0.040	0.100	
NSE	0.010	0.010	0.020	0.370	0.010	0.010	0.010	0.070	0.000	0.000	0.000	0.000	0.540	0.710	0.920	0.580	0.540	0.710	0.950	0.490	0.010	0.030	0.070	0.050	
PSE	0.510	0.750	0.060	1.000	0.510	0.750	0.060	0.880	0.140	0.270	0.000	0.890	0.170	0.380	0.950	0.250	0.510	0.750	0.060	0.820	0.750	0.870	0.120	0.940	
RTS	0.540	0.260	0.230	0.480	0.230	0.200	0.000	0.490	0.750	0.260	0.020	0.790	0.750	0.260	0.020	0.710	0.750	0.260	0.020	0.590	0.140	0.160	0.070	0.530	
SAP 500	1.000	0.010	0.360	0.000	0.080	0.210	0.830	0.130	0.080	0.010	0.000	0.230	0.510	0.000	0.000	0.190	0.510	0.000	0.000	0.200	0.750	0.020	0.000	0.180	
SET	0.750	0.840	0.230	0.430	0.170	0.380	0.010	0.570	0.170	0.380	0.010	0.250	0.310	0.570	0.030	0.830	0.170	0.380	0.010	0.390	0.310	0.570	0.030	0.830	
SSE	0.080	0.170	0.000	0.10	0.360	0.560	0.000	0.090	0.140	0.270	0.000	0.240	0.360	0.560	0.000	0.910	0.360	0.560	0.000	0.010	0.010	0.100	0.000	0.000	
STI	0.540	0.710	0.300	0.730	0.510	0.750	0.070	0.120	1.000	0.230	0.000	0.950	1.000	0.900	0.190	0.980	0.900	0.190	0.990	0.360	0.560	0.310	0.330	0.330	
TSX	0.080	0.210	0.660	0.000	0.010	0.030	0.550	0.000	0.000	0.010	0.370	0.000	0.310	0.570	0.020	0.770	0.170	0.380	0.950	0.300	0.140	0.270	0.000	0.530	
TWI or TSEC	0.080	0.210	0.660	0.210	0.310	0.070	0.020	0.280	0.170	0.030	0.010	0.340	0.750	0.260	0.240	0.990	0.540	0.260	0.260	0.790	0.010	0.100	0.000	0.040	
WIG	0.040	0.100	0.010	0.450	1.000	0.000	0.840	0.320	0.230	0.400	0.090	0.260	0.510	0.750	0.990	0.680	0.510	0.750	0.990	0.640	0.080	0.110	0.040	0.670	
XU100 or ISE	0.360	0.560	0.220	0.100	0.140	0.270	0.110	0.000	0.230	0.400	0.000	0.000	0.360	0.560	0.430	0.000	0.140	0.270	0.190	0.000	0.140	0.270	0.420	0.400	

Table 32: p-values for *hedge funds, commodities and exchange rates* [$a = 5\%$]

ASSETS \ TESTS	ASY				GAR				SAV				EVT				FHS				HOM			
	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR
HFRX-ED	0.010	0.030	0.120	0.020	0.070	0.060	0.030	0.050	0.000	0.000	0.000	0.010	0.010	0.030	0.090	0.100	0.010	0.020	0.060	0.070	0.160	0.280	0.330	0.230
HFRX-EH	0.370	0.260	0.230	0.580	0.570	0.200	0.180	0.560	0.670	0.070	0.060	0.340	0.880	0.960	0.810	0.930	0.390	0.690	0.740	0.530	0.770	0.860	0.840	1.000
HFRX-Macro	0.460	0.610	0.060	0.010	0.670	0.900	0.590	0.120	0.320	0.600	0.380	0.200	0.770	0.580	0.220	0.010	0.670	0.390	0.090	0.020	0.880	0.910	0.220	0.030
HFRX-RVA	0.880	0.920	0.910	0.760	0.010	0.010	0.000	0.000	0.010	0.020	0.000	0.000	0.000	0.000	0.000	0.000	0.010	0.010	0.000	0.000	0.880	0.960	0.400	0.190
HFRXED-DR	0.370	0.500	0.860	0.350	0.660	0.790	0.470	0.050	0.460	0.540	0.570	0.220	0.660	0.790	0.380	0.190	0.560	0.700	0.290	0.170	0.460	0.610	0.090	0.530
HFRXED-MA	0.000	0.000	0.020	0.000	0.000	0.000	0.080	0.000	0.020	0.060	0.330	0.010	0.000	0.000	0.040	0.000	0.010	0.010	0.050	0.000	0.000	0.000	0.000	0.000
HFRXEHEM-N	0.230	0.150	0.200	0.080	0.020	0.010	0.060	0.000	0.070	0.010	0.030	0.010	0.020	0.010	0.060	0.000	0.020	0.010	0.060	0.000	0.130	0.080	0.100	0.020
HFRXRV-FI	0.000	0.000	0.000	0.010	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.000	0.040	0.090	0.140	0.310	0.000	0.000	0.010	0.100	0.370	0.500	0.370	0.190
Brent	0.880	0.950	0.170	0.150	0.320	0.450	0.040	0.740	0.160	0.350	0.010	0.220	0.570	0.710	0.420	0.960	0.660	0.900	0.240	0.920	0.880	0.910	0.380	0.730
Gold Bullion	0.130	0.080	0.000	0.050	0.230	0.150	0.000	0.010	0.180	0.110	0.000	0.060	0.130	0.080	0.000	0.140	0.100	0.050	0.000	0.060	0.000	0.000	0.000	0.000
WTI	0.000	0.000	0.000	0.000	0.000	0.010	0.030	0.030	0.000	0.000	0.000	0.000	0.030	0.080	0.180	0.610	0.020	0.060	0.110	0.490	0.570	0.850	0.980	0.970
GBP/USD	0.670	0.410	0.230	0.350	0.560	0.100	0.100	0.220	0.070	0.040	0.010	0.510	0.230	0.080	0.230	0.380	0.180	0.070	0.310	0.320	0.100	0.050	0.140	0.170
USD/EURO	0.000	0.000	0.050	0.000	0.000	0.000	0.020	0.000	0.000	0.000	0.040	0.000	0.000	0.000	0.060	0.000	0.000	0.000	0.040	0.000	0.010	0.020	0.310	0.020
USD/¥EN	0.000	0.000	0.000	0.000	0.660	0.070	0.200	0.350	0.460	0.010	0.000	0.100	0.770	0.860	0.720	0.820	0.370	0.670	0.800	0.740	0.770	0.580	0.440	0.010

(32)

Table 33: p-values for *hedge funds, commodities and exchange rates* [$a = 1\%$]

ASSETS \ TESTS	ASY				GAR				SAV				EVT				FHS				HOM			
	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR	UN	CC	DQ	VQR
HFRX-ED	0.540	0.710	0.840	0.950	0.360	0.560	0.900	0.630	0.360	0.560	0.920	0.420	0.360	0.560	0.840	0.490	0.360	0.560	0.840	0.510	0.020	0.050	0.040	0.130
HFRX-EH	1.000	0.900	0.980	0.310	0.750	0.870	0.990	0.640	0.310	0.570	0.940	0.330	0.510	0.750	0.060	0.820	1.000	0.010	0.000	0.930	0.040	0.010	0.000	0.570
HFRX-Macro	0.750	0.670	0.930	0.230	0.170	0.380	0.900	0.020	0.080	0.210	0.830	0.780	0.010	0.030	0.520	0.280	0.750	0.670	1.000	0.720	0.230	0.400	0.030	0.120
HFRX-RVA	0.360	0.560	0.230	0.670	0.010	0.010	0.010	0.510	0.010	0.030	0.000	0.360	1.000	0.900	0.140	0.070	0.360	0.560	0.230	0.160	0.230	0.400	0.220	0.550
HFRXED-DR	0.360	0.560	0.050	0.090	0.540	0.710	0.910	0.120	0.750	0.840	0.730	0.270	0.360	0.560	0.780	0.200	0.540	0.710	0.830	0.190	0.080	0.170	0.440	0.060
HFRXED-MA	0.080	0.210	0.860	0.060	0.310	0.570	0.970	0.000	0.080	0.210	0.820	0.000	0.170	0.380	0.930	0.000	0.170	0.380	0.930	0.000	0.310	0.570	0.970	0.000
HFRXEHEM-N	0.310	0.070	0.020	0.050	0.360	0.020	0.000	0.050	0.750	0.010	0.000	0.000	1.000	0.010	0.000	0.500	0.140	0.020	0.000	0.420	0.230	0.020	0.000	0.020
HFRXRV-FI	0.010	0.010	0.000	0.460	0.000	0.000	0.000	0.010	0.000	0.000	0.000	0.300	0.040	0.100	0.000	0.110	0.020	0.050	0.000	0.080	0.010	0.010	0.000	0.170
Brent	0.000	0.000	0.000	0.000	0.020	0.050	0.010	0.000	0.080	0.170	0.140	0.030	1.000	0.900	0.980	0.360	1.000	0.900	0.980	0.390	0.360	0.560	0.880	0.520
Gold Bullion	0.030	0.090	0.000	0.000	0.510	0.750	1.000	0.000	0.750	0.260	0.000	0.000	0.080	0.210	0.000	0.050	0.510	0.750	0.960	0.560	0.510	0.750	0.070	0.880
WTI	0.360	0.240	0.280	0.340	0.510	0.750	1.000	0.950	0.080	0.010	0.000	0.610	0.750	0.870	1.000	0.590	0.310	0.570	0.980	0.260	0.040	0.010	0.000	0.720
GBP/USD	0.000	0.010	0.010	0.340	0.010	0.010	0.010	0.120	0.000	0.010	0.010	0.310	0.020	0.050	0.010	0.080	0.010	0.030	0.010	0.060	0.020	0.050	0.010	0.040
USD/EURO	0.020	0.050	0.000	0.180	0.750	0.260	0.200	0.540	0.020	0.050	0.010	0.230	0.750	0.170	0.120	0.860	0.750	0.170	0.120	0.850	0.750	0.170	0.120	0.730
USD/¥EN	0.510	0.750	1.000	0.940	0.360	0.240	0.130	0.410	0.540	0.260	0.160	0.360	1.000	0.230	0.190	0.260	0.750	0.170	0.120	0.230	0.230	0.200	0.170	0.660

(33)

References

- [AbaBen213] Abad, P.; Benito, S. (2013): "A detailed comparison of value at risk in international stock exchanges". *Mathematics and Computers in Simulation*, Volume 94, August, 2013, Pages 258-276.
- [AbaAl214] Abad, Pilar; Benito, Sonia ; López, Carmen (2014): "A comprehensive review of Value at Risk methodologies". *The Spanish Review of Financial Economics*, 12, 15-32.
- [AloArc206] Alonso, J., Arcos, M. (2006): "Cuatro Hechos Estilizados de las Series de Rendimientos: una Ilustración para Colombia". *Estudios Gerenciales* 22, 103–123.
- [AngAl207] Angelidis, T.; Benos, A.; Degiannakis, S. (2007): "A robust VaR model under different time periods and weighting schemes". *Review of Quantitative Finance and Accounting* 28, 187–201.

- [BalAl208] Bali, T., Hengyong, M., Tang, Y., 2008. The role of autoregressive conditional skewness and kurtosis in the estimation of conditional VaR. *Journal of Banking & Finance* 32, 269–282.
- [BaoAl206] Bao, Y.; Lee, T.; Saltoglu, B. (2006): "Evaluating predictive performance of value-at-risk models in emerging markets: a reality check". *Journal of Forecasting* 25, 101–128.
- [BarAl199] Barone-Adesi, G.; Giannopoulos, K.; Vosper, L. (1999). VaR without correlations for nonlinear portfolios. *Journal of Futures Markets* 19, 583–602.
- [BekGeo205] Bekiros, S. and Georg., D. (2005): "Estimation of value at risk by extreme value and conventional methods: a comparative evaluation of their predictive performance". *Journal of International Financial Markets, Institutions & Money* 15 (3), 209–228.
- [BicDok215] Bickel, Peter J. and Doksum, Kjell A. (2015): "Mathematical Statistics, Basic Ideas and Selected Topics, Vol. 1, (2nd Edition) 2nd Edition, Chapman & Hall/CRC Texts in Statistical Science.
- [BhaRit208] Bhattacharyya, M. and Ritolia, G. (2008): "Conditional VaR using EVT. Towards a planned margin scheme". *International Review of Financial Analysis* 17, 382–395.
- [BroGal210] Brownlees, C.; and Gallo, G. (2010): "Comparison of volatility measures: a risk management perspective"; *Journal of Financial Econometrics* 8, 29–56.
- [CaiXu208] Cai, Z. and Xu, X. (2008): "Nonparametric quantile estimations for dynamic smooth coefficient models"; *Journal of the American Statistical Association* 103(484), 1595–1608.
- [Chr198] Christoffersen, P. (1998): "Evaluating Interval Forecasts". *International Economic Review*, 39, 841–862.
- [Dev186] Devroye, Luc (1986): "Non-Uniform Random Variate Generation"; Springer-Verlag New.
- [EngMan204] Engle, Robert F. & Manganelli, Simone (2004): "CAViaR: Conditional Autoregressive Value at Risk by Regression Quantiles"; *American Statistical Association, Journal of Business & Economic Statistics*, October, Vol. 22, No. 4.
- [ErgJun210] Ergun, A. and Jun, J. (2010): "Time-varying higher-order conditional moments and forecasting intraday VaR and expected shortfall"; *Quarterly Review of Economics and Finance* 50, 264–272.
- [GagAl211] Gaglianone et al (2011), "Evaluating Value-at-Risk Models via Quantile Regression". *Journal of Business & Economic Statistics*, 29:1, 150–160.

- [GenSel204] Gençay, R.; Selçuk, F. (2004): "Extreme value theory and value-at-risk: relative performance in emerging markets"; *International Journal of Forecasting* 20, 287–303.
- [GerAl211] Gerlach, R., Chen, C., Chan, N. (2011): "Bayesian time-varying quantile forecasting for value-at-risk in financial markets". *Journal of Business & Economic Statistics* 29, 481–492.
- [Han194] Hansen, B.E. (1994): "Autoregressive conditional density estimation"; *International Economic Review*, 35, 3, 705–730.
- [Jor206] Jorion, Philippe (2006). "Value at Risk: The New Benchmark for Managing Financial Risk". Third Edition, McGraw Hill Companies.
- [Koe205] Koenker, R. (2005): "Quantile Regression"; Cambridge University Press.
- [KoeMac199] Koenker, R., and Machado, J. A. F. (1999): "Goodness of Fit and Related Inference Processes for Quantile Regression" *Journal of the American Statistical Association*, 94 (448), 1296–1310.
- [KueAl206] Kuester, K.; Mittnik, S.; Paolella, M. (2006): "Value-at-risk prediction: a comparison of alternative strategies". *Journal of Financial Econometrics* 4, 53–89.
- [Kup195] Kupiec, P. (1995). "Techniques for verifying the accuracy of risk measurement models". *Journal of Derivatives*, 3:73–84.
- [LiRac208] Li, Qi & Racine, Jeffrey S. (2008): "Nonparametric Estimation of Conditional CDF and Quantile Functions with Mixed Categorical and Continuous Data"; *Journal of Business & Economic Statistics*, Vol. 26, No. 4 (Oct.), pp. 423–434.
- [LiAl213] Li, Q.; Lin, J.; Racine, J.S. (2013): "Optimal bandwidth selection for nonparametric conditional distribution and quantile functions", *Journal of Business and Economic Statistics*, 31, 57–65.
- [LimNer207] Lima, Luiz Renato and Néri, Breno Pinheiro (2007): "Comparing Value-at-Risk Methodologies"; *Brazilian Review of Econometrics* 27(1) May.
- [LutKra204] Lutkepohl, Helmut and Kratzig, Markus (2004): *Applied Time Series Econometrics*; Cambridge University Press.
- [MarAl209] Marimoutou, V.; Raggad, B.; Trabelsi, A. (2009): "Extreme value theory and value at risk: application to oil market". *Energy Economics* 31, 519–530.
- [McNFre200] McNeil, A. J. and Frey, R. (2000): "Estimation of tail-related risk measures for heteroscedastic financial time series: an extreme value approach"; *Journal of Empirical Finance*; 7, 271–300.

- [Noz210] Nozari, M.; Raei, S.; Jahanguin, P.; Bahramgiri, M. (2010): "A comparison of heavy-tailed estimates and filtered historical simulation: evidence from emerging markets". *International Review of Business Papers* 6–4, 347–359.
- [PolSto210] Polanski, A., Stoja, E. (2010): "Incorporating higher moments into value-at-risk forecasting". *Journal of Forecasting* 29, 523–535.
- [Roc216] Roccioletti, Simona (2016): "Backesting Value at Risk and Expected Shortfall"; Springer.
- [SenAl212] Sener, E.; Baronyan, S.; Meng., L. (2012): "Ranking the predictive performances of value-at-risk estimation methods". *International Journal of Forecasting* 28, 849–873.
- [ShaAl197] Sharma, Ashish; Tarboton, David G.; Lall, Upmanu (1997): "Streamflow simulation: A nonparametric approach"; *Water Resources Research*, Vol. 33, No. 2, pp 291-308, feb.
- [Sil186] Silverman, B. W. (1986): "Density Estimation for Statistics and Data Analysis"; Chapman & Hall, First edition.
- [WanJon195] Wand, M.P. & Jones, M. C. (1995): "Kernel Smoothing"; Chapman & Hall, New York, First edition.
- [YuAl210] Yu, P.L.H.; Li, W.K.; Jin, S. (2010): "On some models for value-at-risk". *Econometric Reviews* 29, 622–641.
- [ZikAkt209] Zikovic, S., Aktan, B. (2009): "Global financial crisis and VaR performance in emerging markets: a case of EU candidate states – Turkey and Croatia." *Proceedings of Rijeka faculty of economics. Journal of Economics and Business* 27, 149–170.

Da COMISSÃO DE PESQUISA DO GET

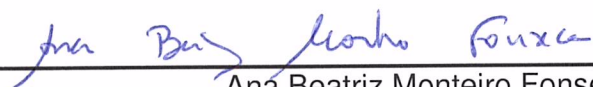
Para Chefia do Departamento de Estatística (GET)

Assunto: Avaliação do Relatório Parcial associado ao Projeto de Pesquisa dos Professores Mariana Albi de Oliveira Souza (SIAPE: 1809003) e Wilson Calmon Almeida dos Santos (SIAPE: 2197625).

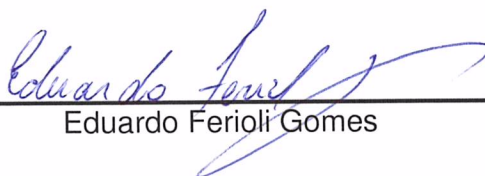
A Comissão de Pesquisa, no uso de sua atribuição, avaliou o relatório parcial do projeto de pesquisa intitulado “Estimação do Número de Grupos em contexto de Dados de Ranking através do Modelo Plackett-Luce: uma Abordagem Hierárquica Bayesiana”, submetido pelos professores .

Considerando a Instrução de Serviço GET – e entendendo que o projeto vem cumprindo o cronograma e seu propósito, a Comissão de Pesquisa do GET considera aprovado o relatório parcial de sua execução.

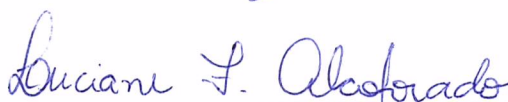
Niterói, 17 de dezembro de 2018



Ana Beatriz Monteiro Fonseca



Eduardo Ferioli Gomes



Luciane Alcoforado



Relatório Parcial do Projeto de Pesquisa:

Estimação do Número de Grupos em contexto de Dados de Ranking através do Modelo Plackett-Luce: uma Abordagem Hierárquica Bayesiana

Identificação do Projeto: Projeto de pesquisa intitulado “*Estimação do Número de Grupos em contexto de Dados de Ranking através do Modelo Plackett-Luce: uma Abordagem Hierárquica Bayesiana*”

Nome e matrícula SIAPE dos docentes responsáveis: Mariana Albi de Oliveira Souza (SIAPE: 1809003) e Wilson Calmon Almeida dos Santos (SIAPE: 2197625).

Resumo do plano inicial de pesquisa: São abundantes os dados de ranqueamento [doravante, **RD** de *Ranking Data*]. RD são essencialmente multivariados e, dessa forma, técnicas tradicionais de análise multivariada podem ser empregadas para analisá-los. Em particular, técnicas de agrupamento podem ser utilizadas para identificar grupos e empates não revelados. Tipicamente, o número de grupos é um insumo de muitos procedimentos de agrupamento e a importância deste parâmetro é frequentemente negligenciada. Todavia, nos últimos 20 anos, métodos interessantes foram propostos para estimá-lo. Contudo, nenhum deles foi desenvolvido especificamente para RD. Nosso objetivo inicial era desenvolver métodos específicos para estimação do número de grupos no contexto de RD a partir do modelo Plackett-Luce através de uma abordagem bayesiana.

Resultados previamente esperados: Neste trabalho esperamos apresentar estimadores razoáveis do número de grupos no contexto de RD. Esperamos que o método proposto apresente melhores resultados do que as alternativas encontradas na literatura [propostas, tipicamente, para outros contextos].

Resumo do projeto executado e resultados efetivamente obtidos: O cronograma inicial está sendo cumprido. De acordo com o mesmo, que se encontra disponível ao final deste relatório parcial, até o final do ano esperávamos: i) compreender e escrever sobre o modelo PL; ii) propor e implementar o modelo PL em uma abordagem bayesiana hierárquica com fins a obter uma maneira nova de estimar o número de grupos em um contexto de RD; iii) implementar a simulação de conjuntos de RD a partir do modelo PL; iv) fazer um levantamento bibliográfico das principais metodologias de estimação do número de grupos, escrever sobre elas e implementá-las; v) criar um programa para analisar os dados simulados a partir do método proposto e algumas alternativas selecionadas. Todas estas tarefas foram cumpridas. No momento, estamos revisando algumas partes do código e do texto. No que diz respeito aos métodos alternativos, consideramos 5 propostas clássicas e 4 mais recentes, totalizando 9 alternativas - um bom número quando comparado a trabalhos similares da literatura. Enviamos em anexo a este relatório dois arquivos com extensão pdf que tratam de versões preliminares que descrevem a metodologia proposta em nosso trabalho (item ii deste resumo) e os diferentes métodos alternativos estudados (item iv deste resumo).

Niterói, 7 de Dezembro de 2018

Mariana Albi de Oliveira Souza [SIAPE: 1809003]

Wilson Calmon Almeida dos Santos [SIAPE: 2197625]

1 Alternative Approaches

Let $Y^n = [Y_j(i)]_{i=1, j=1}^{k, n}$ denote the sample, a $k \times n$ matrix, where the (i, j) element $Y_j(i)$ is the j -th feature of the object i . For example, $Y_j(i)$ could be the rank $R_j(i)$ or the score $X_j(i)$ of object i in the j -th judgment. Usually, in classification or cluster analysis, $Y_j(i)$ corresponds to j -th attribute of the object i . The i -th row of Y^n will be denoted by $\tilde{Y}_i = [Y_1(i) \cdots Y_n(i)]$, which is the features vector of object i . Alternatively, the j -th column of Y^n will be denoted by $Y_j = [Y_j(1) \cdots Y_j(k)]^\top$. We can write $Y^n = [\tilde{Y}_1^\top \cdots \tilde{Y}_k^\top]^\top$ or equivalently $Y^n = [Y_1 \cdots Y_n]$. As pointed by [?], for a specific number of groups $g \leq k$, the classification problem consists in identify a non-empty cluster partition¹ $\mathcal{C}^g = \{\mathcal{C}_1^g, \dots, \mathcal{C}_g^g\}$ of $\{1, \dots, k\}$. Based on sample Y^n , several distinct clustering methods as k -means, agglomerative hierarchical or finite mixtures, *etc*, are intended to generate a cluster partition $\hat{\mathcal{C}}^g = \{\hat{\mathcal{C}}_1^g, \dots, \hat{\mathcal{C}}_g^g\}$ as discussed, for example, in [EveAl211] and [JamAl215]. In a broad sense, it is expected that two distinct objects i and i' will belong to the same cluster if, and only if, their characteristic vectors $\tilde{Y}_i$ and $\tilde{Y}_{i'}$ are close to each other. Let $d(i, i')$ denote the distance between any pair of objects i and i' [the distance between $\tilde{Y}_i$ and $\tilde{Y}_{i'}$]. Usually d is the Euclidean Distance:

$$d(i, i') = \left[\sum_{j=1}^n (Y_j(i) - Y_j(i'))^2 \right]^{1/2}. \quad (1)$$

We expect that in a well-defined cluster $\hat{\mathcal{C}}_\ell^g$ the average of the distances within the cluster $\bar{d}_\ell^2$ is relatively small with respect to the overall distances average² $\bar{d}^2$ where

$$\bar{d}_\ell^2 \equiv \frac{1}{(\#\hat{\mathcal{C}}_\ell^g)^2} \sum_{i \in \hat{\mathcal{C}}_\ell^g} \sum_{i' \in \hat{\mathcal{C}}_\ell^g} d^2(i, i') \text{ and } \bar{d}^2 \equiv \frac{1}{k^2} \sum_{i=1}^k \sum_{i'=1}^k d^2(i, i'). \quad (2)$$

Typically, the goodness of clustering partition $\hat{\mathcal{C}}^g = \{\hat{\mathcal{C}}_1^g, \dots, \hat{\mathcal{C}}_g^g\}$ is measured by comparing the Within Cluster Sum of Squares $W(g)$ with the Between Cluster Sum of Squares $B(g)$:

$$W(g) \equiv \sum_{\ell=1}^g \left[\frac{(\#\hat{\mathcal{C}}_\ell^g) \bar{d}_\ell^2}{2} \right] \text{ and } B(g) \equiv \frac{k}{2} \bar{d}^2 - W(g). \quad (3)$$

¹ The cluster partition $\mathcal{C}^g$ is a class containing g elements $\mathcal{C}_1^g, \dots, \mathcal{C}_g^g$ which are non-empty sets satisfying both (i) $\mathcal{C}_\ell^g \cap \mathcal{C}_{\ell'}^g = \emptyset$, if $\ell \neq \ell'$ and (ii) $\cup_{\ell=1}^g \mathcal{C}_\ell^g = \{1, \dots, k\}$.

² It is possible to express both $\bar{d}_\ell^2$ and $\bar{d}^2$ in terms of the distances between the characteristic vectors and the centroids - see [?]. We chose to follow the definitions in [KrzLai188].

Using the entities defined in equations (2) and (3), we can also define the Total Sum of Squares by

$$T(g) \equiv \frac{k}{2} d^2. \quad (4)$$

C.1 Silhouette [Sil]: Let $d(i, i')$ denote the distance between any objects i and i' [for example, as in equation (1), it could express the Euclidean distance between the features vectors $\tilde{Y}_i$ and $\tilde{Y}_{i'}$]. It seems reasonable evaluate the distance between any object i and a cluster $\tilde{\mathcal{C}}$ by

$$\tilde{d}(i, \tilde{\mathcal{C}}) \equiv \frac{1}{\#\tilde{\mathcal{C}}} \sum_{i' \in \tilde{\mathcal{C}}} d(i, i'). \quad (5)$$

In particular, for a given cluster partition $\hat{\mathcal{C}}^g = \{\hat{\mathcal{C}}_1^g, \dots, \hat{\mathcal{C}}_g^g\}$, suppose $i \in \tilde{\mathcal{C}}_i$ and $\tilde{\mathcal{C}}_i \in \hat{\mathcal{C}}^g$. So we can evaluate both the distance from object i to its own cluster $[OC_i^g]$ as the distance to the nearest distinct cluster $[NDC_i^g]$ by

$$OC_i^g \equiv \tilde{d}(i, \tilde{\mathcal{C}}_i) \text{ and } NDC_i^g \equiv \min_{\tilde{\mathcal{C}} \in \hat{\mathcal{C}}^g; \tilde{\mathcal{C}} \neq \tilde{\mathcal{C}}_i} \{\tilde{d}(i, \tilde{\mathcal{C}})\}. \quad (6)$$

[Rou187] proposed adopt a measure of goodness of classification $S(g)$ [Silhouette] as the average of the individual classification goodness measures $s_1^g, \dots, s_k^g$:

$$S(g) \equiv \frac{1}{k} \sum_{i=1}^k s_i^g, \text{ where } s_i^g \equiv \begin{cases} \frac{NDC_i^g - OC_i^g}{\max(OC_i^g, NDC_i^g)}, & \text{if } \#\tilde{\mathcal{C}}_i > 1 \\ 0, & \text{if } \#\tilde{\mathcal{C}}_i = 1 \end{cases} \quad (7)$$

As discussed in [KauRou190], $-1 \leq s_i^g \leq 1$ and, therefore, $-1 \leq S(g) \leq 1$. The number of clusters g could be estimated by

$$\hat{g} = \arg \max_{g \in \{2, \dots, k\}} \{S(g)\}. \quad (8)$$

C.2 Calinski & Harabasz [CH]: A simpler method was proposed in [CalHar174] and consists in evaluate the statistic [Calinski & Harabasz Index or CH, for short]:

$$CH(g) = \frac{[B(g) / (g - 1)]}{[W(g) / (k - g)]}. \quad (9)$$

The number of clusters g could be estimated by

$$\hat{g} = \arg \max_{g \in \{2, \dots, k-1\}} \{CH(g)\}. \quad (10)$$

They call the proposed index of “variance ratio criterion” and highlight the fact that this works as a F-statistic.

C.3 Krzanowski and Lai [KL]: [KrzLai188] proposed to estimate g by maximizing the statistic $KL(g)$:

$$KL(g) \equiv \left| \frac{DIFF(g)}{DIFF(g+1)} \right|, \quad (11)$$

where $DIFF(g) \equiv (g-1)^{2/n} W(g-1) - g^{2/n} W(g)$.

The estimated value of g is obtained by

$$\hat{g} = \arg \max_{g \in \{2, \dots, k-1\}} KL(g). \quad (12)$$

C.4 Gap: The Gap approach was proposed in [TibAl201]. It is based on Gap measure defined as

$$Gap_k^*(g) = E_k^* \{\log(W(g))\} - \log(W(g)), \quad (13)$$

where E_k^* denotes expectation under some reference distribution. In the original proposal, it is proposed to replace $E_k^* \{\log(W(g))\}$ with an moment estimate $\frac{1}{M} \sum_{m=1}^M \{\log(W_m^*(g))\}$, where $W_m^*(g)$ is the Within Cluster Sum of Squares [equation (3)] evaluated for a cluster partition containing g clusters when the original data set Y^n is replaced by Y_{*m}^n , an artificial simulated dataset³; resulting in the Gap Statistic

$$Gap_k(g) = \frac{1}{M} \sum_{m=1}^M \{\log(W_m^*(g))\} - \log(W(g)). \quad (14)$$

Finally, the number of clusters g could be estimated by

$$\hat{g} = \min_{g \in \{1, \dots, k-1\}} \{g; Gap_k(g) \geq Gap_k(g+1) - sd(\log(W^*(g)))\}, \quad (15)$$

where $sd(\log(W^*(g)))$ is an estimate of the standard deviation of $\log(W_m^*(g))$.⁴

C.5 Jump: The proposal of [SugJam203] is similar to the KL and Gap methods and consists of obtaining a cluster partition so that the Within-Cluster Sum of the Squares is sufficiently small in the sense that any subdivision does

³ Specifically, M datasets must be generated: $\{Y_{*m}^n\}_{m=1}^M$. A possible strategy for generate the artificial dataset $Y_{*m}^n = [Y_{*m,1} \dots Y_{*m,n}]$ is to generate each feature column $Y_{*j,n}$ by drawing random numbers from the uniform distribution defined over the range of the associated observed feature vector Y_j , where $Y^n = [Y_1 \dots Y_n]$.

⁴ The rule adopted for estimate g , made explicit in equation (15), is a kind of “1-standard-error” rule, as discussed in [BreAl184].

not substantially decrease it. Basically, g could be estimated by maximizing the Jump Statistic $Jump(g)$:

$$\hat{g} = \arg \max_{g \in \{1, \dots, k-1\}} Jump(g), \text{ where} \quad (16)$$

$$Jump(g) \equiv W(g)^{-n/2} - W(g-1)^{-n/2} \text{ and } W(0) \equiv 0. \quad (17)$$

Although there is a more general formulation for the Jump statistic⁵, we will adopt this simpler version as done in [ZhaAl217] and [KolAl216].⁶

C.6 Hartigan [Har]: [Har175] proposed an simple and iterative rule of thumb in which g could be estimated by

$$\hat{g} = \min_{g \in \{1, 2, \dots, k-1\}} \{g; H(g) \leq 10\}, \text{ where} \quad (18)$$

$$H(g) \equiv (k - g - 1)^{-1} \left[\frac{W(g)}{W(g+1)} - 1 \right]. \quad (19)$$

The Hartigan statistics H , called a mean square ratio, is a measure of the reduction of within-cluster variance. It is distributed as $F_{1, k-g-1}$ if the features are normal [see [Har175], p.136].

N.1 Slope [Slo]: Similar to [Rou187], [FujAl214] adopt as a measure of goodness of the cluster partition the silhouette statistic $S(g)$ defined in equation (7). As pointed out by them, while in the silhouette method the goal is to maximize $S(g)$, they also intend to identify “clusters that cannot be partitioned into smaller clusters” in a sense that any partition could drastically decreases $S(g)$. Their proposal to solve the trade-off between a larger silhouette and its

⁵ Let $\hat{\mathcal{C}}^g = \{\hat{\mathcal{C}}_1^g, \dots, \hat{\mathcal{C}}_g^g\}$ be the cluster partition and $\bar{Y}_\ell$ denotes the centroid of $\hat{\mathcal{C}}_\ell^g$. $W(g)$ could be replaced for

$$\hat{d}_k = n^{-1} \sum_{i=1}^k \min_{1 \leq \ell \leq g} \left(\tilde{Y}_i - \bar{Y}_\ell \right) \Gamma^{-1} \left(\tilde{Y}_i - \bar{Y}_\ell \right)^\top,$$

where Γ is the clusters covariance matrix. [see, for example, [Wan210]]. As discussed in [SugJam203], if Γ is the identity matrix, then $\hat{d}_k$ is proporcional to $W(g)$. The authors themselves suggest to adopt the simpler formulation, since it is difficult to estimate Γ .

⁶ It is also possible to generalize Jump by replacing n with another constant in the equation (17). However, [SugJam203] suggest to adopt this specific value. Again, this is typically the choice made in the literature - [ZhaAl217] and [KolAl216].

negligible variability of it [when increasing the number of clusters by one] is to consider the slope statistic $\mathcal{S}_p(g)$ defined as

$$\mathcal{S}_p(g) = -[S(g+1) - S(g)] S(g)^p, \quad (20)$$

where p is a positive integer value that controls the trade-off between silhouette level and its variation. The number of groups is given by ⁷

$$\hat{g} = \arg \max_{g \in \{2, \dots, k-1\}} \{\mathcal{S}_p(g)\}. \quad (21)$$

N.2 Instability [Ins]: [Wan210] presented a method based on the minimization of the “Cluster Instability”, an alternative measure of goodness of clustering previously discussed, for example, in [BenAl206] and [LanAl204]. Suppose $\tilde{Y}, \tilde{Y}_1, \dots, \tilde{Y}_k, \tilde{\mathcal{Y}}, \tilde{\mathcal{Y}}_1, \dots, \tilde{\mathcal{Y}}_k$ are independent and identically distributed n -dimensional random vectors and $\tilde{Y} \sim \tilde{P}$. As before, let $Y^n = [\tilde{Y}_1^\top \dots \tilde{Y}_k^\top]^\top$ and $\mathcal{Y}^n = [\tilde{\mathcal{Y}}_1^\top \dots \tilde{\mathcal{Y}}_k^\top]^\top$ denote two samples with size k draw from $\tilde{P}$. Typically, for any specific number of groups g , a clustering ψ_g is a mapping $\psi_g: \mathbb{R}^n \mapsto \{1, \dots, g\}$ that associates for any vector n -dimensional vector Y the associated group $\psi_g(Y)$, an integer between 1 and g . In [Wan210], the distance between two clustering $\psi_{1,g}$ and $\psi_{2,g}$ is defined as

$$\Delta(\psi_{1,g}, \psi_{2,g}) \equiv E \left\{ \left| \mathbb{I}[\psi_{1,g}(\tilde{Y}) = \psi_{1,g}(\tilde{\mathcal{Y}})] - \mathbb{I}[\psi_{2,g}(\tilde{Y}) = \psi_{2,g}(\tilde{\mathcal{Y}})] \right| \right\}. \quad (22)$$

Accordingly, the distance between two clustering $\psi_{1,g}$ and $\psi_{2,g}$ is the probability that the two identically distributed random vectors $\tilde{Y}$ and $\tilde{\mathcal{Y}}$ are similarly classified in only one of the two clustering ⁸. Generally, a clustering algorithm Ψ_g associates a unique cluster mapping $\Psi_g(Y^n) = \psi_{Y^n,g}$ with any training sample Y^n . ⁹ Otherwise, if the training sample is $\mathcal{Y}^n$, then, $\Psi_g(\mathcal{Y}^n) = \psi_{\mathcal{Y}^n,g}$ and it is possible that both $\psi_{\mathcal{Y}^n,g}$ and $\psi_{Y^n,g}$ are distinct clustering mappings, even if the samples were drawn from the same population. The clustering instability of Ψ_g is defined by

$$\mathcal{I}(\Psi_g) \equiv E[\Delta(\Psi_g(Y^n), \Psi_g(\mathcal{Y}^n))] = E[\Delta(\psi_{Y^n,g}, \psi_{\mathcal{Y}^n,g})]. \quad (23)$$

⁷ [FujAl214] highlight the fact that $\mathcal{S}_p(g)$ is a discrete version of the $S(g)^{p+1} / (p+1)$ [justifying the name “slope”]. They also commented that if the correlation between the silhouette and the number of groups is positive, the number of groups could be chosen as 1.

⁸ Is it straightforward to see that $\Delta(\psi_{1,g}, \psi_{2,g})$, defined in equation (7), is equal to

$$\Pr \left[\mathbb{I} \left\{ \psi_{1,g}(\tilde{Y}) = \psi_{1,g}(\tilde{\mathcal{Y}}) \right\} + \mathbb{I} \left\{ \psi_{2,g}(\tilde{Y}) = \psi_{2,g}(\tilde{\mathcal{Y}}) \right\} = 1 \right].$$

⁹ We point out the fact that we are using the same notation for the data matrix and for the sample: Y^n .

All the expectations¹⁰ should be taken with respect to the $\tilde{P}$ distribution. Since $\tilde{P}$ is unknown, $\mathcal{I}(\Psi_g)$ is also unknown. [FanWan212] propose the bootstrap method to produce an estimate $\widehat{\mathcal{I}}(\Psi_g)$ of the cluster instability $\mathcal{I}(\Psi_g)$.¹¹ Therefore, the optimal number of clusters can be estimated by

$$\hat{g} = \arg \min_{2 \leq g \leq k} \widehat{\mathcal{I}}(\Psi_g). \quad (24)$$

N.3 Utility [Uti]: [LiaAl212] considered a mixed data context. They propose to adopt a cluster validity index $[CUF(\hat{\mathcal{C}}^g)]$ based on the category utility function [[GluCor185]] to access the quality of a clustering partition $\hat{\mathcal{C}}^g = \{\hat{\mathcal{C}}_1^g, \dots, \hat{\mathcal{C}}_g^g\}$ and estimate g by maximizing¹² it:

$$\hat{g} = \arg \max_{g_{\min} \leq g \leq g_{\max}} CUF(\hat{\mathcal{C}}^g). \quad (25)$$

Suppose that the features vector $\tilde{Y}_i$ has the first n_C coordinates corresponding to categorical variables, while the last n_N are related to numerical ones. Then, $\tilde{Y}_i = [\tilde{Y}_i^C; \tilde{Y}_i^N]$, where $\tilde{Y}_i^C = [Y_1(i) \ \dots \ Y_{n_C}(i)]$ and $\tilde{Y}_{N_i} = [Y_{n_C+1}(i) \ \dots \ Y_n(i)]$. The index $CUF(\hat{\mathcal{C}}^g)$ is given by $CUF(\tilde{\mathcal{C}}^g) \equiv \frac{n_C}{n} CUC(\tilde{\mathcal{C}}^g) + \frac{n_N}{n} CUN(\tilde{\mathcal{C}}^g)$, where $CUN(\tilde{\mathcal{C}}^g)$ and $CUC(\tilde{\mathcal{C}}^g)$ denote validity indexes in numerical and categorical contexts, respectively. In the numerical case,

$$CUN(\tilde{\mathcal{C}}^g) \equiv \sum_{j=n_C+1}^n \left\{ \frac{1}{g} \left[\sum_{\ell=1}^g \left(\frac{\#\tilde{\mathcal{C}}_{\ell}^g}{k} \right) [\delta_j^2 - \delta_{\ell j}^2] \right] \right\} \quad (26)$$

¹⁰ Note that clustering instability is defined by take a double expectation as below $\mathcal{I}(\Psi_g) = E \left\{ E_{\tilde{P}} \left\{ \left| \mathbb{I} \left\{ \psi_{1,g}(\tilde{Y}) = \psi_{1,g}(\tilde{Y}) \right\} - \mathbb{I} \left\{ \psi_{2,g}(\tilde{Y}) = \psi_{2,g}(\tilde{Y}) \right\} \right| \right\} \right\}$.

¹¹ Let $Y^n = [\tilde{Y}_1^\top \ \dots \ \tilde{Y}_k^\top]^\top$ be the observed sample. The algorithm is shown below.

1. Generate B independent bootstrap sample-pairs $(Y_{b*}^n, \mathcal{Y}_{b*}^n)$, $b = 1, \dots, B$, where $Y_{b*}^n = [\tilde{Y}_{b*,1}^\top \ \dots \ \tilde{Y}_{b*,k}^\top]^\top$ and $\mathcal{Y}_{b*}^n = [\tilde{\mathcal{Y}}_{b*,1}^\top \ \dots \ \tilde{\mathcal{Y}}_{b*,k}^\top]^\top$. Each sample consists of k observations generated from empirical distribution $\hat{P}$ [sampling with replacement from $\tilde{Y}_1, \dots, \tilde{Y}_k$].

2. Construct $\psi_{Y_{b*}^n, g}$ and $\psi_{\mathcal{Y}_{b*}^n, g}$ based on $(Y_{b*}^n, \mathcal{Y}_{b*}^n)$, $b = 1, \dots, B$.

3. For each pair $(\psi_{Y_{b*}^n, g}, \psi_{\mathcal{Y}_{b*}^n, g})$, estimate $\Delta(\psi_{Y_{b*}^n, g}, \psi_{\mathcal{Y}_{b*}^n, g})$ by $\hat{\Delta}_b$, where $\hat{\Delta}_b \equiv \frac{1}{k^2} \sum_{i=1}^k \sum_{i'=1}^k \left| \mathbb{I} \left\{ \psi_{Y_{b*}^n, g}(\tilde{Y}_i) = \psi_{Y_{b*}^n, g}(\tilde{Y}_{i'}) \right\} - \mathbb{I} \left\{ \psi_{\mathcal{Y}_{b*}^n, g}(\tilde{Y}_i) = \psi_{\mathcal{Y}_{b*}^n, g}(\tilde{Y}_{i'}) \right\} \right|$.

4. The clustering instability $\mathcal{I}(\Psi_g)$ can be estimated by $\widehat{\mathcal{I}}(\Psi_g) \equiv \frac{1}{B} \sum_{b=1}^B \hat{\Delta}_b$. We must stress that [Wan210] proposed previously estimate $s(\Psi_g)$ through a “modified leave-many-out crossvalidation scheme”.

¹² Both $g_{\min}$ and $g_{\max}$ are constants to be defined previously.

where $\delta_{\ell j}^2$ and δ_j^2 are, respectively, the sample variances of the j -th feature within the cluster $\hat{\mathcal{C}}_\ell^g$ and overall. In the categorical case, for each categorical variable indexed by j [$j = 1, \dots, n_C$] suppose there are exactly d_j distinct values $Y_j^{*1}, \dots, Y_j^{*d_j}$ in the set $\{Y_j(1), \dots, Y_j(k)\}$. Let $C_{j\lambda}^g \equiv \{i; Y_j(i) = Y_j^{*\lambda}\}$ for $\lambda = 1, \dots, d_j$. Then,

$$CUC(\tilde{\mathcal{C}}^g) \equiv \sum_{j=1}^{n_C} \left\{ \frac{1}{g} \left[\sum_{\ell=1}^g \left(\frac{\#\tilde{\mathcal{C}}_\ell^g}{k} \right) Q_{\ell j} \right] \right\}, \text{ where} \quad (27)$$

$$Q_{\ell j} \equiv \sum_{\lambda=1}^{d_j} \left\{ \left(\frac{\#(C_{j\lambda}^g \cap \tilde{\mathcal{C}}_\ell^g)}{\#\tilde{\mathcal{C}}_\ell^g} \right)^2 - \left(\frac{\#C_{j\lambda}^g}{k} \right)^2 \right\}.$$

It is important to stress that [LiaAl212] also presented a new k -prototypes algorithm¹³ [”the modified k -prototypes algorithm”] to create an optimal cluster partition $\hat{\mathcal{C}}^g = \{\hat{\mathcal{C}}_1^g, \dots, \hat{\mathcal{C}}_g^g\}$ for each value of g . Similar to the k -means algorithm [see [Har175]], the algorithm is a recursive strategy in which objects are allocated sequentially, resulting in new clustering partitions until the clustering partition remains invariant.¹⁴ The novelty in k -prototype algorithms is that they permits consider mixed data.

N.4 Quantization Error Modeling [QEM]: Let $W_g \equiv \frac{1}{k}W(g)$ be the error caused by grouping [within-class variance or Mean Square Error, MSE]. [KolAl216] proposed a new clustering goodness measure, entitled Parameterized Multiplicative Cost Function [PCF]:

$$PCF(g) \equiv W_g \times g^{-\hat{\beta}}, \quad (28)$$

where $\hat{\beta}$ is an OLS estimate of β , the slope coefficient in a linear regression of $\ln(W_g)$ on $\ln(g)$.¹⁵ A rate-distortion curve is the function that associates, for each possible number g of clusters, a clustering error measure. Although it is possible to choose another measures, they adopt W_g as clustering error measure and propose to use the PCF as the rate-distortion curve [cost function]. Their approach is based on the two main ideas that: (i) cost functions usually

¹³ See [Hua198].

¹⁴ For a given optimal cluster partition $\hat{\mathcal{C}}^g$, they propose identify the worst cluster $\hat{\mathcal{C}}_w^g \in \hat{\mathcal{C}}^g$ and remove it, allocating each object in the nearest cluster not removed. The new $g - 1$ clusters are taken to be the starting $g - 1$ clusters in the the modified k -prototypes algorithm. It will produce a new optimal clustering partition $\hat{\mathcal{C}}^{g-1} = \{\hat{\mathcal{C}}_1^{g-1}, \dots, \hat{\mathcal{C}}_{g-1}^{g-1}\}$ with $g - 1$ groups. The worst cluster $\hat{\mathcal{C}}_w^g \in \hat{\mathcal{C}}^g$ is defined as one that maximizes the sum of between cluster information entropies when exactly one cluster is removed. More details in [LiaAl212].

¹⁵ Based on $g = 1, \dots, g_{\max}$, where $g_{\max} \leq k$ is a constant to be choosed.

are decreasingly [in g] and (ii) their shapes should depend on the distribution of the data. Finally, the number of cluster is estimated by

$$\hat{g} = \arg \min_g PCF(g). \quad (29)$$

References

- [BenAl206] Ben-David, S.; Von Luxburg, U. & Pal, D. (2006): "A sober look at stability of clustering"; In Proc. 19th Ann. Conf. Learn. Theory (COLT 2006), Ed. G. Lugosi and H. Simon, pp. 5–19. Berlin: Springer.
- [CalHar174] Calinski, T. and Harabasz, J. (1974): "A dendrite method for cluster analysis"; Commun. Stat.—Theory Methods, 3 (1), 1–27.
- [EveAl211] Everitt, Brian S; Landau, Sabine; Leese, Morven; Stahl, Daniel (2011): "Cluster Analysis"; Wiley, Fifth Edition.
- [BreAl184] Breiman, Leo; Friedman, Jerome; Olshen, Richard; Stone, Charles (1984): "Classification and regression trees"; Chapman & Hall/CRC.
- [FanWan212] Fang, Yixin and Wang, Junhui (2012): "Selection of the number of clusters via the bootstrap method"; Computational Statistics and Data Analysis, 56, 468–477.
- [FujAl214] Fujita, André; Takahashi, Daniel Y.; Patriota, Alexandre G. (2014): "A non-parametric method to estimate the number of clusters"; Computational Statistics and Data Analysis, 73, 27–39.
- [Gor199] Gordon, A. (1999): "Classification"; 2nd edn. London: Chapman and Hall-CRC.
- [GluCor185] Gluck, M.A. and Corter, J.E. (1985): "Information, uncertainty, and the utility of categories"; in: Proceeding of the 7th Annual Conference of the Cognitive Science Society, Lawrence Erlbaum Associates, Irvine, CA, 1985, pp.283–287.
- [Har175] Hartigan, J.A. (1975): "Clustering algorithms"; John Wiley & Sons Inc., New York, 1975.
- [Hua198] Huang, Zhexue (1998): "Extensions to the k-means algorithm for clustering large data sets with categorical values"; Data Mining and Knowledge Discovery 2 (3), 283–304.
- [JamAl215] James, Gareth; Witten, Daniela; Hastie, Trevor; Tibshirani, Robert (2015): "An Introduction to Statistical Learning with Applications in R"; Springer.

- [KauRou190] Kaufman, Leonard and Rousseeuw, Peter J. (1990): "Finding Groups in Data: An Introduction to Cluster Analysis"; Wiley, John Wiley & Sons.
- [KrzLai188] Krzanowski, W. J. and Lai, Y. T. (1988): "A criterion for determining the number of groups in a data set using sum-of-squares clustering"; *Biometrics*, Vol. 44, No. 1 (Mar.), pp.23–34.
- [LanAl204] Lange, T.; Roth, V.; Braun, M.; Buhmann, J. (2004): "Stability-based validation of clustering solutions"; *Neural Computatio*, 16, 1299–323.
- [LiaAl212] Liang, Jiye; Zhao, Xingwang; Li, Deyu; Cao, Fuyuan; Dang, Chuangyin (2012): "Determining the number of clusters using information entropy for mixed data"; *Pattern Recognition*, 45, 2251–2265.
- [Rou187] Rousseeuw, P. (1987): "Silhouettes: a graphical aid to the interpretation and validation of cluster analysis". *Journal of Computational and Applied Mathematics*, 20. 53-65.
- [TibAl201] Tibshirani, Robert; Walther, Guenther; Hastie, Trevor (2001): "Estimating the number of clusters in a data set via the gap statistic"; *J. R. Stat. Soc.: Ser. B (Stat. Methodol.)*, 63 (2), pp.411-423.
- [SugJam203] Sugar, Catherine A. and James, Gareth M. (2003): "Finding the number of clusters in a dataset: An information-theoretic approach"; *Journal of the American Statistical Association*, 98 (463), Sep., pp. 750–763.
- [KolAl216] Kolesnikov, Alexander; Trichina, Elena; Kauranne, Tuomo (2016): "Online frame-based clustering with unknown number of clusters"; *Pattern Recognition*, 48, pp. 941–952.
- [Wan210] Wang, Junhui (2010): "Consistent selection of the number of clusters via crossvalidation"; *Biometrika*, 97, 4, pp. 893–904.
- [ZhaAl217] Zhang, Yaqian; Mandziuk, Jacek; Quek, Chai Hiok; Goh, Boon Wooi (2017): "Curvature-based method for determining the number of clusters". *Information Sciences*, 415–416. 414–428.

1 Ranking Data: context, Plackett-Luce and the Bayesian Approach

Ranking data typically arise from a context where a collection of k objects $\mathcal{O}^k \equiv \{\mathcal{O}_1, \dots, \mathcal{O}_k\}$ are compared in some(s) aspect(s) in such a way that is possible to define an ordered list with the best one in the first position, the second best in the second position, *etc.* Let $\mathcal{O} = \mathcal{O}^k$ [we will suppress the length of the object list]. We will use alternatively the label i for identify the i -th object $\mathcal{O}_i$. Thus, it is possible to rewrite the list as $\mathcal{O} \equiv \{1, \dots, k\} = \mathbb{N}_k$.¹ The same objects can be compared/judged/ordered in different ways. Suppose we observe n [eventually distinct] such judgments without ties [there is a single best object, a single second best, *etc.*]. For each $j = 1, \dots, n$, the associated ranking will be denoted by the k -dimensional column **rank vector**

$$R_j = [R_j(1) \ \cdots \ R_j(k)]^T, \quad (1)$$

where $R_j(i)$ denotes the ranking of object i in the j -th judgment.² Observe that $R_j(i) = l$ if the i -th object $[\mathcal{O}_i]$ is considered the l -th best in the j -th judgment. If there are no ties, $R_j(i) = l$ if, and only if, $R_j^{-1}(l) = i$, where $R_j^{-1} = [R_j^{-1}(1) \ \cdots \ R_j^{-1}(k)]^T$ is the order vector which indicates in the l -th coordinate the object classified as the l -th best. Our sample of ranking data will be denoted by R^n , an $k \times n$ matrix $R^n = [R_1 \ \cdots \ R_n]$. [AlvYu214] points that ranking data arises naturally from “situations where it is desired to rank a set of individuals or objects” according with an criterion.³ Suppose that the specific criterion is a univariate quantitative measure. Denote by $X_j = [X_j(1) \ \cdots \ X_j(k)]^T$ the **score vector** evaluated by the j -th judge, where $X_j(i)$ is the score of i -th object in the j -th judgment. Then, R_j satisfies

$$R_j(i) < R_j(i') \Leftrightarrow X_j(i) > X_j(i'). \quad (2)$$

This relationship defines uniquely the rank vector R_j as a permutation of $(1, \dots, k)$ for a given score vector X_j . Although we are not necessarily assuming that the observed rank vectors are generated from some score vector, this is a central idea in the Ranking Model we will present further up.

¹ $\mathbb{N}_k$ denotes the set of the first k natural numbers. $\mathbb{N}$ denotes the set of all natural numbers.

² R_j is the rank vector of the j -th judge. Judges could be: i) people [comparing objects like films, books, ...], ii) animals [comparing, by use or consumption frequencies, objects like foods, colors,...] or iii) inanimate entities as, for example, sport competitions, or clinical experiments producing ranking of medicines or medical procedures.

³ They also recognizes that ranking data may arise as a powerful data transformation strategy, as common in nonparametric inference.

Here, we will follow an approach identified in [Mar195] as “Model the Ranking Process” instead of the alternative approach based on “Model the Population of Rankers”. In the approach adopted here, we are more interested in describe mathematically the process that a representative judge evaluate the objects more than evaluate the distribution of eventually distinct judgment patterns. We presume that variations in observed ranking data are not associated with the judges. Rather, they stem from random factors affecting each object’s performances. We assume that ties are not observed [in another way, if the observed ranking is based on scores, suppose that scores are drawn from continuous distributions]. However, we allow distinct objects to be indistinguishable in a sense that equivalent objects tends to be ranked as similar in different judgments [in particular, alternations between the rankings of equivalent objects should be more often than those between non-equivalent objects].

For the sake of simplicity, we will adopt similar notations here for both data and the random elements of which they are a realization. In particular, both observed score vector X_j and rank vector R_j are assumed to be realizations from the random vectors X_j and R_j , respectively. We will assume that $X, X_1, \dots, X_n$ are iid [identically independently distributed] as well as $R, R_1, \dots, R_n$. Define Π_k as the set of all permutations of $(1, \dots, k)$, what is,

$$\Pi_k \equiv \{\pi = (\pi(1), \dots, \pi(k)); \pi \text{ is a permutation of } (1, \dots, k)\}. \quad (3)$$

Note that, as pointed in [FliVer193], if ties are not allowed, the random vector R must be an element of Π_k , whose cardinality is $\#(\Pi_k) = k!$. Moreover, if $\mathcal{P}_{k!}$ denotes the standard simplex set, described by

$$\mathcal{P}_{k!} = \left\{ p = (p_1, \dots, p_{k!}) \in \mathbb{R}^{k!}; \sum_{i=1}^{k!} p_i = 1 \right\}, \quad (4)$$

then any probability distribution P for the random vector R is a discrete probability distribution associated with some specific vector $p \in \mathcal{P}_{k!}$ [where $\Pr(R = \pi_j) = p_j$ for every $j = 1, \dots, k!$]. [Mar195] defines a probability model $\mathcal{P}$ for R as a subset of $\mathcal{P}_{k!}$ that, typically, will be parametrized by θ in a space Θ : $\mathcal{P} = \{P_\theta; \theta \in \Theta\}$.

1.1 Ranking Data and the Plackett-Luce Model

According to [AlvYu214], there are different classes of models for ranking data [e.g., Order Statistics Models, Paired Comparison Models, Distance-Based

Models, ...]. They describe a generic Order Statistics Model [OSM, for short] $\mathcal{P} = \{P_\theta; \theta \in \Theta\}$ as a probability model satisfying, for each $\pi \in \Pi_k$,

$$P_\theta(R = \pi) = \Pr(X_{\pi^{-1}(1)} > X_{\pi^{-1}(2)} > \cdots > X_{\pi^{-1}(k)}), \forall \theta \in \Theta. \quad (5)$$

Thereby, an OSM can be defined implicitly by models for the score vector X , whose distribution will depends on θ . For example, if is assumed that $X \sim N(\mu, \Omega)$, then, $\mathcal{P}$ would be the classical Thurstone model - [Thu127]. In the **Plackett-Luce model** [[Luc159] and [Pla175]], it is assumed that the scores are independent random variables, each having Gumbel distribution with possibly different parameters. The Plackett-Luce model [**PL** model henceforth] is formally described in Definition 1.

Definition 1 [PL 1] Let $\Theta = [0, +\infty)^k$. The PL model $\mathcal{P} = \{P_\theta; \theta \in \Theta\}$ is defined as the class of distributions P_θ induced by the following assumptions:
i) $\forall \theta \in \Theta$, $P_\theta(R = \pi) = \Pr(X_{\pi^{-1}(1)} > X_{\pi^{-1}(2)} > \cdots > X_{\pi^{-1}(k)}), \pi \in \Pi_k$;
ii) $X_1, \dots, X_k$ are independent;
iii) For a fixed $\theta = (\theta_1, \dots, \theta_k)$, $X_i - \ln(\theta_i) \sim \text{Gumbel}$, or
 $P(X_i - \ln(\theta_i) \leq \xi) = P(X_i \leq \ln(\theta_i) + \xi) = \exp(-\exp(-\xi))$.

In PL model, for every $\theta \in \Theta$, each distribution P_θ can be easily described by its associated probability mass function [pmf] - it has a closed form. Wherefore, it is also possible to define PL model alternatively by displaying the associated pmf, cf. Definition 2 - both Definitions 1 and 2 are equivalent [[Mar195]].

Definition 2 [PL 2] Let $\Theta = [0, +\infty)^k$. As before, let PL model be denoted by $\mathcal{P} = \{P_\theta; \theta \in \Theta\}$. For each θ , P_θ is described by the following pmf

$$P_\theta(R = \pi) = \prod_{i=1}^{k-1} \frac{\theta_{\pi^{-1}(i)}}{[\theta_{\pi^{-1}(i)} + \cdots + \theta_{\pi^{-1}(k)}]} = \prod_{i=1}^k \frac{\theta_{\pi^{-1}(i)}}{[\sum_{l=i}^k \theta_{\pi^{-1}(l)}]}, \forall \pi \in \Pi_k. \quad (6)$$

As mentioned in [AlvYu214], for every $2 \leq k' \leq k$ and $1 \leq i'_1 < i'_2 < \cdots < i'_{k'} \leq k$, if $i \in \{i'_1, \dots, i'_{k'}\}$, the PL model implies in the following equality

$$\Pr(X_i = \max(X_{i'_1}, \dots, X_{i'_{k'}})) = \frac{\theta_i}{\theta_{i'_1} + \cdots + \theta_{i'_{k'}}}. \quad (7)$$

In particular, the ratio $\frac{\theta_i}{\theta_1 + \cdots + \theta_k}$ returns the probability of i -th object O_i be ranked as first among all. The larger the parameter θ_i , the higher is its chance to be better ranked. In PL model, the relative order of O_i with respect to the $O_{i'_1}, \dots, O_{i'_{k'}}$ does not depend on other objects that have not been contemplated in the comparison.

In [Pla175], the adoption of such a model was motivated by an analogy with a k -horse race. Illustratively, for each horse O_i [labeled as i] is associated a positive number θ_i . If there are 4 horses, the horse with label 3 will be the first with probability $\theta_3/(\theta_1 + \dots + \theta_4)$. Conditionally to the triumph of the horse labeled as 3, the probability of the horse labeled as 1 to finish in second place does not depend on θ_3 and is given by $\theta_1/(\theta_1 + \theta_2 + \theta_4)$. Continuing this process we obtain the probability function associated with the second definition of PL model. [Mar195] stresses that in PL model “parameters are identifiable only up to a multiplicative constant”. He further emphasizes that this model does not bellows to a regular exponential family model and presents the following form to the log-likelihood:

$$\ell(\theta) \equiv \ell(\theta; R^n) \equiv \sum_{j=1}^n \left\{ \sum_{i=1}^{k-1} \left[\ln \left(\frac{\theta_{R_j^{-1}(i)}}{\sum_{\zeta=i}^k \theta_{R_j^{-1}(\zeta)}} \right) \right] \right\}. \quad (8)$$

It is not a simple task to maximizes the log-likelihood appearing in equation 8. Numerical methods to deal with the direct optimization of it can be highly time consuming. [Hun204] proposed an Minorization–Maximization [MM] Algorithm to deal with similar problems. This procedure is iterative and generates a sequence of estimates that converge to the maximum likelihood solution. Basically, it is necessary to find a [minorizing] function $\mathcal{Q}^m$ such that: i) $\mathcal{Q}^m$ can be easily maximized [more than likelihood function ℓ] and ii) the maximization of $\mathcal{Q}^m$ leads to maximization of ℓ . For the specific likelihood function associated with the PL model, [Hun204] identified the following specific minorizing function $\mathcal{Q}^m$:

$$\mathcal{Q}^m(\theta) \equiv \mathcal{Q}(\theta; R^n; \theta^{(m)}) \equiv \sum_{j=1}^n \left\{ \sum_{i=1}^{k-1} \left[\ln \left(\theta_{R_j^{-1}(i)} \right) - \frac{\sum_{\zeta=i}^k \theta_{R_j^{-1}(\zeta)}}{\sum_{\zeta=i}^k \theta_{R_j^{-1}(\zeta)}^{(m)}} \right] \right\}, \quad (9)$$

where $\theta^{(m)} = (\theta_1^{(m)}, \dots, \theta_k^{(m)})$ is a known vector [optimal in m -th iteration]. Given $\theta^{(m)}$, we can find another point $\theta^{(m+1)}$ such that

$$\theta^{(m+1)} \equiv \arg \max_{\theta} \mathcal{Q}^m(\theta). \quad (10)$$

Then, starting from an arbitrary point $\theta^{(0)}$, for each $m = 1, 2, \dots$, we should to define a minorizing function $\mathcal{Q}^m = \mathcal{Q}(\theta; R^n; \theta^{(m)})$ and maximizes it to defines $\theta^{(m+1)}$. The problem (10) has an closed form solution. For $t = 1, \dots, k$, let us define $w_t \equiv \sum_{j=1}^n \mathbb{I}(R_j(t) < k)$. For a given $\theta^{(m)} = (\theta_1^{(m)}, \dots, \theta_k^{(m)})$, if $w_t \neq 0, \forall t \in \{1, \dots, k\}$, then,

$$\theta^{(m+1)} = (\theta_1^{(m+1)}, \dots, \theta_k^{(m+1)}) , \text{ where } \theta_i^{(m+1)} = \frac{w_i}{\sum_{j=1}^n \sum_{\zeta=1}^{k-1} \delta_{j\zeta i} \left[\sum_{\xi=\zeta}^k \theta_{R_j^{-1}(\xi)}^{(m)} \right]^{-1}}$$

and $\delta_{j\zeta i} = \mathbb{I}(i \in \{R_j^{-1}(\zeta), \dots, R_j^{-1}(k)\}) = \mathbb{I}(\zeta \leq R_j(i) \leq k)$.

[Hun204] shows that MM algorithm leads to the ML estimates.⁴

1.2 A Hierarchical Bayesian approach to the PL model

In the classical PL model, it is common - see [AlvYu214] - to adopt an alternative parameterization by letting $\alpha_i = \ln(\theta_i)$, $\alpha = (\alpha_1, \dots, \alpha_k)$. In this case, the log-likelihood (8) takes the form:

$$\ell(\alpha) \equiv \sum_{j=1}^n \left\{ \sum_{i=1}^{k-1} \left[\alpha_{R_j^{-1}(i)} - \ln \left(\sum_{\zeta=i}^k \exp(\alpha_{R_j^{-1}(\zeta)}) \right) \right] \right\}. \quad (11)$$

The true parameter vector $\alpha = (\alpha_1, \dots, \alpha_k)$ [or $\theta = (\theta_1, \dots, \theta_k)$] contains valuable information about the cluster structure of the objects. For example, objects labeled i and i' should be classified in the same cluster [there is a true non observed tie] if and only if $\alpha_i = \alpha_{i'}$ [equivalently, $\theta_i = \theta_{i'}$]. In this case, both the probability of observing $R_i > R_{i'}$ and the probability of observing $R_{i'} > R_i$ are equal to 50% - see equation (7). Otherwise, if $\alpha_i \neq \alpha_{i'}$ [$\theta_i \neq \theta_{i'}$], then objects labeled i and i' should be classified in distinct clusters [there is no true tie between objects i and i']. Suppose for example, without loss of generality, that $\alpha_i > \alpha_{i'}$ [$\theta_i > \theta_{i'}$]. In the latter case, the probability of observing $R_i < R_{i'}$ [i ranked better than i'] is bigger than 50%. Based on this discussion, we can define a true **classification vector** - Definition 3.

Definition 3 (Classification Vector) *The true vector α [or θ] induces a true **classification vector** $c = (c_1, \dots, c_k)$ where c_i is an positive integer indicating the cluster or group of object labeled i . c is characterized by:*

- i) *there is an integer g such that $g \equiv \max_{1 \leq i \leq k} c_i$; g is the number of clusters;*
 - ii) *for each $l \in \{1, \dots, g\}$, there exists at least one object i such that $c_i = l$;*
 - iii) *for any pair of objects i and i' , $c_i < c_{i'}$ if and only if $\alpha_i > \alpha_{i'}$.*
- Even more, $c_i = c_{i'}$ if and only if $\alpha_i = \alpha_{i'}$.*

⁴ In “PlackettLuce” R package, the MM algorithm is available for deal with PL model estimation task - [TurAl218]. More recently, a new method [weaver algorithm] was proposed in [DonYin218] for estimate the PL model parameters.

The classification vector c has only natural numbers as coordinates. We emphasize that the number of clusters g is defined as the maximum value of c coordinates. The true vector α uniquely defines both the true classification vector c and the true number of clusters g . Unfortunately, the parameter vector α is unknown. As usual, we could use the ranking data to estimate its value or, more interesting, directly estimate g as a parameter of the ranking data distribution.

Likelihood: The first task to estimate the number of groups g is to relate the likelihood function to the parameter g . As usual in Bayesian inference [see, for example, [MigGam199] or [BerSmi194]], let us to denote the **likelihood** by $p(R^n|\alpha)$. Using equation (11) it is easy to see that ⁵:

$$p(R^n|\alpha) \equiv \prod_{j=1}^n \left\{ \prod_{i=1}^{k-1} \left[\frac{\exp\left(\alpha_{R_j^{-1}(i)}\right)}{\sum_{\zeta=i}^k \exp\left(\alpha_{R_j^{-1}(\zeta)}\right)} \right] \right\}. \quad (12)$$

As the likelihood function does not depend explicitly on g we propose an hierarchical approach [see [Con210]] relating vector α to scalar g . We use five intermediate stages that will be explained below.

Stage 1: In the first stage we relate α to the classification vector c . As we have discussed before, α determines c . The reverse is not true: a given classification vector c does not determines α . However, we propose the following degenerate distribution for the first stage:

$$p(\alpha|c, s) = \mathbb{I}(\alpha = \boldsymbol{\alpha}(c, s)), \quad (13)$$

where $c = (c_1, c_2, \dots, c_k)$ [where $c_i \in \mathbb{N}_g$, $\forall i = 1, \dots, k$], $s > 0$ is a positive scalar and $\boldsymbol{\alpha}: \mathbb{N}^k \times \mathbb{R}_+^* \mapsto \mathbb{R}_+^k$ is a deterministic function satisfying $\alpha_i > \alpha_{i'}$ if and only if $c_i < c_{i'}$, when $\alpha = \boldsymbol{\alpha}(c, s)$. The linkage function $\boldsymbol{\alpha}$ must be chosen. In particular, we have chosen

$$\boldsymbol{\alpha}(c, s) = (\tilde{\alpha}_{c_1}, \dots, \tilde{\alpha}_{c_k}), \quad (14)$$

where

$$\tilde{\alpha} = \left((g-1) \ln \left(1 + \frac{1}{s} \right), (g-2) \ln \left(1 + \frac{1}{s} \right), \dots, \ln \left(1 + \frac{1}{s} \right), 0 \right)$$

⁵ Observe that $R_j^{-1}(i)$ indicates the label of the object ranked as the i -th best in the j -th judgment, as before.

and $\tilde{\alpha}_{c_i}$ indicates the c_i -th coordinate of the vector $\tilde{\alpha}$.

The scale parameter s controls the degree of separation between distinct groups. We will propose a prior distribution for s in the last [5-th] stage.

Stage 2: For a given classification vector $c = (c_1, \dots, c_k)$, we already know that $g = \max_{1 \leq i \leq k} c_i$. We also can define a distribution vector $d = (d_1, \dots, d_g)$, where d_l indicates the number of coordinates in c that assume value l . Obviously, $d_l \in \{1, \dots, k\}, \forall l = 1, \dots, g$ and $\sum_{l=1}^g d_l = k$. Repair that the length of d also indicates the value of g . d does not determines c . In stage 3 we propose a prior distribution for the classification vector c depending on the distribution vector d . For a fixed $d = (d_1, \dots, d_g), [d_l \in \mathbb{N}_k, l = 1, \dots, g \text{ and } \sum_{l=1}^g d_l = k]$, let C^d denote the set of compatible classification vectors, where:

$$C^d = \{(c_1, \dots, c_k) \in \mathbb{N}_g^k; \#\{i; c_i = l\} = d_l, \forall l = 1, \dots, g\}. \quad (17)$$

Then we adopt an uniform prior distribution on C^d , as follows.

$$p(c|d) = \frac{\mathbb{I}(c \in C^d)}{\#C^d}, \text{ where } \#C^d = \binom{k}{d_1, \dots, d_g} = \frac{k!}{d_1! \dots d_g!}. \quad (18)$$

Stage 3: For a given distribution vector $d = (d_1, \dots, d_g)$, it is true that $g = \text{length}(d)$. Although g does not determine d , there exists a finite set D^g of distribution vectors compatible with each possible g . For a given natural g , D^g is described by

$$D^g = \left\{ (d_1, \dots, d_g) \in \mathbb{N}_k^g; \sum_{i=1}^g d_i = k \right\}. \quad (19)$$

We propose an uniform prior distribution on D^g for d :

$$p(d|g) = \frac{\mathbb{I}(c \in D^g)}{\#D^g}, \text{ where } \#D^g = \binom{k-1}{k-g}. \quad (20)$$

Stage 4: The three previous stages allow us to relate the likelihood function to the parameter g . Obviously, g take values in $\{1, \dots, k\}$. Then a possible choice for the prior distribution for g is

$$p(g) = \frac{\mathbb{I}(g \in \mathbb{N}_k)}{k}. \quad (21)$$

Stage 5: Finally, we choose to adopt an independent prior distribution for the scale parameter s . As s is strictly positive, we propose to use a gamma distribution with hyperparameters (α, β) , described by the probability density function

$$p(s) = \frac{s^{\alpha-1} e^{-\frac{s}{\beta}}}{\Gamma(\alpha) \beta^\alpha} \mathbb{I}(s \in \mathbb{R}_+^*). \quad (22)$$

Posterior, PLHB Method: Following [WesHar197], the **joint posterior distribution** is given by

$$p(\alpha; c; d; s; g | R^n) \propto p(R^n | a) \times p(a | c, s) \times p(c | d) \times p(d | g) \times p(g) \times p(g). \quad (23)$$

If $\theta = (\alpha; c; d; s)$ we can rewrite the joint posterior distribution as $p(\theta; g | R^n) = p(\alpha; c; d; s; g | R^n)$. The **posterior distribution for g** is

$$p(g | R^n) = \int p(\theta, g | R^n) \partial \theta. \quad (24)$$

Our proposal is to **define the estimated number of groups as the mode of prior distribution $p(g | R^n)$** :

$$\hat{g} = \arg \max_{1 \leq g \leq k} p(g | R^n). \quad (25)$$

Practical issues: Let $\phi = (\theta; g)$. It is not a simple task to obtain the posterior distribution $[p(\phi | R^n) \text{ or } p(g | R^n)]$ directly. However, as discussed in [Gam206], using Markov Chain Monte Carlo algorithms [MCMC] it is possible to simulate a Markov chain $\{\phi^{(m)}\}_m$ with the posterior distribution $p(\phi | R^n)$ as its stationary distribution. We have chosen here to use more specifically the Metropolis Hastings [MH] algorithm, described first in [Has170]. Basically, starting from a initial condition $\phi^{(1)}$ and a proposal distribution q the MH algorithm allows us to generate $\theta^{(2)}, \theta^{(3)}, \dots, \theta^{(M+BI)}$ by employ the recursive rule

$$\phi^{(m)} \equiv \begin{cases} \phi^{*m}, & \text{if } u_m^* < \tau^{(m)} \\ \phi^{(m-1)}, & \text{otherwise} \end{cases}, \text{ where} \quad (26)$$

$$\tau^{(m)} = \min \left\{ 1, \frac{p(\phi^{*m} | R^n)}{p(\phi^{(m-1)} | R^n)} \times \frac{q(\phi^{(m-1)} | \phi^{*m})}{q(\phi^{*m} | \phi^{(m-1)})} \right\}, \quad (27)$$

u_m^* is drawn from a uniform(0, 1) distribution and ϕ^{*m} is simulated from the

proposal distribution⁶ $q(\cdot|\phi^{(m-1)})$. Here, we opt by using

$$q(\phi|\cdot) = p(\phi) \text{ [prior for } \phi]. \quad (28)$$

After discarding a certain phase of burn-in [of length BI], the sampled parameters $\theta^{(BI+1)}, \dots, \theta^{(M+BI)}$ can be viewed as a random sample from $p(\phi|R^n)$. As discussed in [BroRob198] and [CowCar196], it is necessary use an sufficiently large value of $M + BI$ to attain convergence, and several convergence diagnostic tools have been proposed in the literature that allows us to choose the appropriated M value. In our simulation study, we adopt $M = 9 \times 10^3$ and $B = 10^3$.⁷

References

- [AlvYu214] Alvo, Mayer and Yu, Philip L.H. (2014): “Statistical Methods for Ranking Data”; Springer.
- [BerSmi194] Bernardo, J. M. and Smith, A. F. M. (1994): “Bayesian Theory”; Chichester: Wiley.
- [BroRob198] Brooks, STEPHEN and Roberts, GARETH O. (1998): “Convergence assessment techniques for Markov chain Monte Carlo”; Statistics and Computing, December , Volume 8, Issue 4, pp 319–335.
- [Con210] Congdon, P.D. (2010): “Applied Bayesian Hierarchical Methods”; CRC, Boca Raton, FL.
- [CowCar196] Cowles, Mary Kathryn and Carlin, Bradley P. (1996): “Markov Chain Monte Carlo Convergence Diagnostics A Comparative Review”; Journal of the American Statistical Association, Volume 91, Issue 434.
- [DonYin218] Dong, Fanghu and Yin, Guosheng (2018): “Maximum likelihood estimation for incomplete multinomial data via the weaver algorithm”; Stat Comput, 28, pp.1095–1117.
- [FliVer193] Fligner, Michael A. and Verducci, Joseph S. (1993): “Probability Models and Statistical Analyses for Ranking Data”; Lecture Notes in Statistics, Springer-Verlag.

⁶ The proposal $q(\cdot|\text{cond})$ is a density for each value specified in “cond”- “cond” could be interpreted as a parameter of this distribution.

⁷ Also we discard 9 out the remaining 10 observations to reduce a possible dependency structure in the chain, thereby obtaining an effective chain $\theta^{(BI+1)}, \theta^{(BI+11)}, \theta^{(BI+21)}, \dots$ with size 900.

- [Gam206] Gamerman, D. (2006): “Markov Chain Monte Carlo: Stochastic Simulation for Bayesian Inference”; Chapman and Hall/CRC, London, U.K, 2006.
- [Has170] Hastings, W. K. (1970): “Monte Carlo Sampling Methods using Markov Chains and their Applications”; *Biometrika*, 57, 97–109.
- [Hun204] Hunter, David R. (2004): “MM Algorithms for Generalized Bradley-Terry Models”; *The Annals of Statistics*, Vol. 32, No. 1, 384–406.
- [Luc159] Luce, R. D. (1959): “Individual choice behavior: A Theoretical Analysis”; New York: Wiley.
- [Mar195] Marden, John I. (1995): “Analyzing and Modeling Rank Data”; Chapman & Hall/CRC Monographs on Statistics & Applied Probability, London, 1st Edition.
- [MigGam199] Migon, H. S. and Gamerman, D. (1999): “Statistical Inference: an Integrated Approach”; Arnold, London.
- [Pla175] Plackett, Robert L. (1975): “The Analysis of Permutations”; *Appl. Statist* 24 (2):193–202.
- [Thu127] Thurstone, L. L. (1927): “A law of comparative judgement”; *Psychological Reviews*, Vol.34, 4, pp.273-286.
- [TurAl218] Turner, Heather; Kosmidis, Ioannis; Firth, David; Etten, Jacob van (2018): “PlackettLuce: Plackett-Luce Models for Rankings”; R package, version 0.2-3.
- [WesHar197] West, M. and Harrison, J.F. (1997): “Bayesian forecasting and dynamic models”; Springer, 2nd ed., Springer Verlag Series in Statistics.

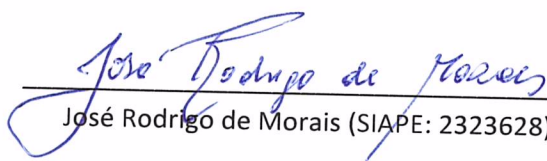
Da: Comissão de Ensino do GET

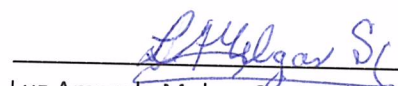
Para: Chefe do Departamento de Estatística

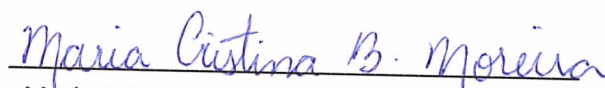
Assunto: Avaliação do Relatório Final associado ao Projeto de Ensino do Prof. Wilson Calmon Almeida dos Santos.

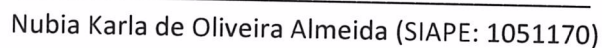
A Comissão de Ensino, no uso de suas atribuições, avaliou o Relatório Final do projeto de ensino intitulado "Estudo Dirigido de Análise Real", submetido pelo docente Wilson Calmon Almeida dos Santos (SIAPE: 2197625). Ao considerar a Instrução de Serviço GET 03-2016 e por entender que o projeto cumpriu seu propósito, a Comissão de Ensino considera validada a execução do mesmo.

Niterói, 19 de dezembro de 2018.


José Rodrigo de Moraes (SIAPE: 2323628)


Luz Amanda Melgar Santander (SIAPE: 1193497)


Mária Cristina Bessa Moreira (SIAPE: 106164)


Nubia Karla de Oliveira Almeida (SIAPE: 1051170)



Relatório Final do Projeto de Iniciação à Docência

Título do Projeto: Estudo Dirigido de Análise Real

Identificação do Projeto: Projeto de Iniciação à Docência

Nome do Docente Responsável (Coordenador do Projeto): Wilson Calmon Almeida dos Santos (2197625)

Período Global de Execução e Avaliação: 01/02/2018 a 20/12/2018 [avaliação contínua]

Resultados Obtidos: O projeto visava o estudo, de forma independente, porém supervisionada, do conteúdo de Análise Real, em nível básico. Os alunos envolvidos corresponderam amplamente às expectativas do projeto. Do ponto de vista prático, semana após semana, dedicaram-se, mostraram desenvoltura e alta capacidade de compreensão dos tópicos propostos pelo professor a partir da mera leitura do livro “Análise Real” [Lima, Elon]. Em cada encontro, revezaram-se na apresentação de conteúdos previamente indicados, expondo e dirimindo [muitas vezes, eles próprios] dúvidas que surgiam. A parte a ser exposta por cada aluno só era divulgada no momento da exposição. Ambos os alunos, porém, não se intimidavam quando solicitados em cada momento. Estavam sempre prontos para expor e colaborar com a exposição um do outro. Além da dedicação, eles foram proativos na busca pelo conhecimento. Recorreram a materiais adicionais como o livro “Curso de Análise” e a vídeo-aulas como as disponibilizadas pelo IMPA [<http://video.impa.br/>]. O professor, mero expectador das exposições feitas pelos alunos, conclui o projeto com a certeza de que o principal objetivo foi cumprido: os alunos se tornaram mais independentes na busca pelo conhecimento, dando passos mais firmes em termos do autodidatismo e tornando-se mais íntimos dos livros. Destaco ainda, no que tange ao aprendizado, a evolução dos alunos em termos de maturidade, criticidade e autoconhecimento.

Aluno(s) Vinculado(s):

PEDRO FILIPE TEIXEIRA DE ARAUJO (matrícula 116054011)

RODRIGO MOURAO CALAND DIAS (matrícula 217054058)

*** Relatos dos alunos vinculados anexados.**



Relato Final - Projeto de Iniciação à Docência

Estudo Dirigido de Análise Real

Aluno: PEDRO FILIPE TEIXEIRA DE ARAUJO (matrícula 116054011)

Orientador: Wilson Calmon

Relato: O projeto proporcionou bons resultados em relação ao objetivo de tornar o aluno mais independente e autodidata em relação aos estudos. A cada semana foram abordados tópicos diferentes do livro Análise Real, onde se notou um bom entendimento da grande maioria das seções apresentadas por parte dos alunos, enquanto o professor não deixou de ajudar em caso de dúvidas, orientando os dois alunos de modo que ambos conseguissem esclarecer as questões por conta própria, e até mesmo respondendo a indagações de temas relacionados, mas não abordados no livro em questão. A apresentação revezada ajudou a manter uma dinâmica equilibrada e bem distribuída entre os dois alunos, ao passo que criou maior desenvoltura e facilidade ao apresentar os resultados. Também foi notório o bom relacionamento e a troca de ideias entre todos os envolvidos, sempre dispostos a se ajudar. Conforme passaram os capítulos foram propostas Listas de Exercício, selecionados de modo que, além de atingir um bom conhecimento teórico através da leitura do livro, os alunos também fossem instigados à prática e aplicação dos resultados. Destaca-se a melhora no entendimento das partes teóricas não apenas do conteúdo apresentado, como também nas disciplinas cursadas que são relacionadas ao mesmo. Conclui-se que ao final do projeto não só os objetivos iniciais de independência e autodidatismo nos estudos foram alcançados, como também a melhora na desenvoltura e forma de apresentar conceitos e resultados, além de ter sido um ótimo ambiente para troca de ideias e crescimento no que tange a parte de estudos e relacionamentos.



Relato Final - Projeto de Iniciação à Docência

Estudo Dirigido de Análise Real

Aluno: RODRIGO MOURAO CALAND DIAS (matrícula 217054058)

Orientador: Wilson Calmon

Relato:

Para um aluno situado no segundo período da graduação, estudar Análise Real pode parecer uma tarefa bem complicada. No entanto, o que esse projeto foi capaz de mostrar é que, com bastante esforço e uma orientação personalizada, barreiras desse porte podem ser ultrapassadas seguramente.

Em relação ao método pedagógico empregado, gostaria de destacar o estímulo à postura ativa dos participantes através da leitura individual e de apresentações semanais no quadro. Tal aspecto não somente tornou o projeto mais estimulante e divertido, como também proporcionou uma certa evolução aos participantes no campo da oratória. Na maioria dos encontros, tinha a plena sensação de estar simplesmente conversando em meio à amigos interessados em Matemática. Nesse sentido, afirmo que o Estudo Dirigido passou longe de ser uma experiência tediosa ou desgastante, apesar da dificuldade dos tópicos estudados. Destaco ainda a qualidade do material selecionado, que se ajustou perfeitamente à necessidade de um processo de aprendizado sucinto, sem muitas divagações.

Segundamente, devo tecer elogios aos envolvidos. Meu colega de projeto, apesar de tão assustado quanto eu em relação ao livro, demonstrou simpatia e extrema boa vontade para trocar experiências. Quanto ao professor, distinguiu-se pelo domínio impecável do conteúdo, sempre capaz de tirar as dúvidas mais complexas com a paciência de um monge tibetano.

Diante do que foi exposto, fica evidente que a iniciativa foi um grande sucesso em todos os aspectos, proporcionando aos alunos envolvidos maior coragem e inteligência emocional para enfrentar as dificuldades que estão por vir no decorrer da graduação. Talvez seja possível criticar ligeiramente o projeto por não estabelecer, em alguns casos, conexão entre os assuntos abordados e a Estatística. Contudo, reconheço que isso já era esperado, haja vista o tempo curto que se dispunha semanalmente.

Quadro de Horários

2019 - J

Código	Disciplinas 2019-01	Turma	Curso	Prof em 2019-1	Seg	Ter	Qua	Qui	Sex
GET00040	Estatística V	A1	Química + Química Industrial + Geofísica	Luz Amanda	16-18		16-18		
GET00040	Estatística V	B1	Geofísica + Ciências ambientais	Ludmilla		11-13		11-13	
GET00041	Bioestatística	A1	Biomedicina	Víctor (sub.)		09-11		11-13	
GET00053	Estatística Básica Aplicada às Ciências Humanas	A1	Psicologia	Rebecca (sub.)		18-20		18-20	
GET00059	Bioestatística I	A1	Medicina Veterinária	Víctor (sub.)	16-18	16-18			
GET00116	Fundamentos de Estatística Aplicada	A1	Ciências Contábeis	Marco Aurélio	16-18		16-18		
GET00116	Fundamentos de Estatística Aplicada	B1	Administração	Yasmin (sub.)	20-22		20-22		
GET00117	Métodos Estatísticos Aplicados à Economia I	A1	Ciências Econômicas	Eduardo			11-13		11-13
GET00117	Métodos Estatísticos Aplicados à Economia I	B1	Ciências Econômicas	Adrian			20-22		20-22
GET00118	Métodos Estatísticos Aplicados à Economia II	A1	Ciências Econômicas	Marcio			11-13		11-13
GET00118	Métodos Estatísticos Aplicados à Economia II	B1	Ciências Econômicas	Marcio			20-22		20-22
GET00119	Estatística Básica para Engenharia II	A1	Desenho Industrial	Jessica	09-11		09-11		
GET00121	Introdução à Probabilidade e à Estatística	A1	Matemática	Adrian		11-13		11-13	
GET00121	Introdução à Probabilidade e à Estatística	B1	Matemática	Ana Maria	18-20		18-20		
GET00122	Probabilidade e Estatística	A1	Ciências da Computação	Luz Amanda	14-16		14-16		
GET00122	Probabilidade e Estatística	B1	Sistemas de Informação	Moisés	20-22		20-22		
GET00149	Delimitação de Experimentos	A1	Estatística	Patrícia	09-11		09-11		
GET00158	Estatística Básica Aplicada às Ciências da Vida	A1	Ciências Biológicas	Luis Guillermo		11-13		11-13	
GET00169	Estatística Básica para Ciências Humanas I	A1	Relações Internacionais	José Rodrigo	07-09		07-09		
GET00170	Estatística Geral*	B1	Turismo	Keila	14-16		14-16		
GET00170	Estatística Geral*	C1	Ciências Sociais	Rebecca (sub.)		20-22		20-22	
GET00170	Estatística Geral*	D1	Sociologia	Hugo	18-20		18-20		
GET00170*	Estatística Geral*	A1	Arquivologia + Biblioteconomia	Keila	16-18		16-18		
GET00176	Estatística Aplicada às Ciências da Vida	A1	Nutrição	Ana Beatriz		09-11		09-11	
GET00177	Estatística Básica para Engenharia	A1	Engenharias	Yasmin (sub.)	14-16		14-16		
GET00177	Estatística Básica para Engenharia	B1	Engenharias	Valentin	16-18		16-18		
GET00177	Estatística Básica para Engenharia	C1	Engenharias	Wilson		09-11		09-11	
GET00177	Estatística Básica para Engenharia	D1	Engenharias	Douglas		11-13		11-13	
GET00177	Estatística Básica para Engenharia	E1	Engenharias	Víctor (sub.)		14-16		14-16	
GET00177	Estatística Básica para Engenharia	F1	Engenharias	Wilson		16-18		16-18	