

Estimação de Value at Risk para horizontes superiores a um dia por meio dos processos estocásticos GARCH e APARCH combinados com simulação de Monte Carlo

Guilherme Fernandes Sanches

<http://www.bndes.gov.br/bibliotecadigital>

Estimação de Value at Risk para horizontes superiores a um dia por meio dos processos estocásticos GARCH e APARCH combinados com simulação de Monte Carlo

Guilherme Fernandes Sanches*

Resumo

Neste artigo, mostra-se como é possível estimar o Value at Risk (VaR) de uma carteira de ativos para horizontes superiores a um dia por meio de simulação, em vez da adoção da regra da raiz quadrada do tempo. Utilizam-se os processos estocásticos *generalized autoregressive conditional heteroskedasticity* (GARCH), proposto por Bollerslev (1986), e *asymmetric power autoregressive conditional heteroskedasticity* (APARCH), proposto por Ding, Granger e Engle (1993), para a previsão da volatilidade um passo a frente. Observa-se a *performance* da medida de VaR diante da utilização de diferentes distribuições de probabilidade para a distribuição condicional dos retornos por meio dos testes de Cobertura Incondicional, de Kupiec (1995), e Aderência da Distribuição Analítica aos Dados, proposto por Berkowitz (2001). Apontam-se os problemas de confiar-se cegamente em testes baseados exclusivamente na ocorrência de falhas, como é o caso do teste de Kupiec. Modelam-se os processos estocásticos

*Economista do BNDES. Este artigo é de exclusiva responsabilidade do autor, não refletindo, necessariamente, a opinião do BNDES.

mencionados com três distribuições condicionais diferentes: normal-padrão, *t-student* e *t-student* assimétrica. A série de tempo utilizada compreende retornos do Ibovespa no período de 2 de janeiro de 2006 a 25 de novembro de 2013. A separação dos dados entre dentro e fora da amostra foi feita dinamicamente, de forma que se utilizam todas as informações disponíveis até o tempo “T” para produzir previsões para o tempo “T+H” no intervalo de observações fora da amostra.

Abstract

In this paper we show how to estimate the Value-at-Risk of a portfolio for time horizons over one day through Monte Carlo Simulation instead of using the square root of time rule. We implemented the stochastic processes generalized autoregressive conditional heteroskedasticity (GARCH) proposed by Bollerslev (1986) and asymmetric power autoregressive conditional heteroskedasticity (APARCH) proposed by Ding, Granger and Engle (1993) to obtain one step ahead volatility forecasting. We observed how the VaR measure performance behaved under different probability distributions for returns conditional distribution through both the Inconditional Coverage Test proposed by Kupiec (1995) and the Adherence of Theoretical Probability Distribution to Data Test proposed by Berkowitz (2001). We showed the problems involved when comparing alternative models through a test that solely accounts for the number of failures, like Kupiec Test. We modeled the previously mentioned stochastic processes under three different conditional distributions: Standard Gaussian, *t-student* and skewed *t-student*. The time series data comprehends returns of Ibovespa from 2 Jan 2006 to 25 Nov 2013. The separation between in-sample and out-of-sample data was performed dynamically so that we use all available information at time “T” to produce “T+H” forecasting for the out-of-sample data.

Introdução

Histórico e literatura

Na literatura de finanças quantitativas e gestão de risco financeiro, o VaR (ou Valor em Risco) é citado como uma medida de previsão de perda máxima esperada para uma determinada carteira de ativos financeiros. De forma mais detalhada, para um conjunto de variáveis composto pela carteira de ativos, nível de significância e horizonte de tempo, o VaR é definido como o valor cuja probabilidade de obtenção de um retorno inferior a ele no horizonte de tempo especificado é igual ao nível de significância utilizado. Isto é, ao se construir uma medida de VaR, espera-se que o retorno observado seja menor do que o VaR estimado apenas em fração equivalente ao nível de significância proposto.

A modelagem de VaR é muito importante para as instituições financeiras, entre outros motivos, porque ela é responsável pela construção de medidas de perda máxima esperada em um horizonte de tempo e nível de significância determinados para a carteira de participações societárias. Neste trabalho, utiliza-se uma série histórica de retornos do Índice da Bolsa de Valores de São Paulo (Ibovespa) como *proxy* para uma série de retornos obtidos por uma carteira de participações societárias.

A estimação de VaR teve seu início em 1989, na criação da metodologia RiskMetrics pelo banco JP Morgan [Longerstaey e Spencer (1996)]. Tal metodologia tinha como fundamento a estimação da volatilidade¹ pelo modelo *exponentially weighted moving average* (EWMA), em que a volatilidade segue um processo recursivo de-

¹ No caso univariado, estima-se apenas a volatilidade. No caso multivariado, devem-se estimar volatilidades e correlações.

pendente da volatilidade anterior multiplicada por um fator λ e pelo quadrado do retorno anterior multiplicado por $(1 - \lambda)$. A estimativa de VaR é obtida por meio da multiplicação da volatilidade estimada por um quantil da distribuição normal-padrão.²

A partir de então, várias metodologias de VaR foram propostas com base em diferentes modelos para a volatilidade e diferentes distribuições de probabilidade condicional para o retorno. Fernandes (2008) faz um breve relato do desenvolvimento dos modelos da família *autoregressive conditional heteroskedasticity* (ARCH), fundamentais na modelagem de volatilidade de séries temporais financeiras.

O primeiro modelo da família nasceu de uma pesquisa de Robert Engle, em 1982, no artigo “Autoregressive Conditional Heteroscedasticity with Estimates of the Variance of United Kingdom Inflation”. Tal modelo ficou conhecido como ARCH e seu diferencial era a possibilidade de capturar aglomerados de volatilidade que outros processos estocásticos não eram capazes de capturar. Embora a série utilizada tenha sido de dados de inflação, não levou muito tempo até os pesquisadores observarem sua aplicação em séries de retorno financeiro, que também apresentavam heterocedasticidade condicional.

Bollerslev (1986) aprimorou o modelo proposto por Engle (1982) por meio da criação do modelo GARCH, com previsão da volatilidade um pouco mais acurada por considerar a volatilidade como função não apenas da volatilidade passada, mas também do quadrado do retorno passado subtraído de sua média incondicional. Vários outros modelos foram propostos, como em Barndorff-Nielsen e Shephard (2001) e Taylor (1986). Berkowitz e O’Brien (2002), Danielsson (2002) e Lee e Saltoglu (2001) abordam a utilização de

² Para simplificar, assumiu-se média incondicional dos retornos igual a zero. Se fosse estimada sua média, bastaria somá-la à multiplicação proposta para a obtenção do VaR.

modelos do tipo ARCH para a estimação de VaR. Giot e Laurent (2004) fazem uma comparação entre modelos de volatilidade latente – que utilizam dados diários – e realizada – que utilizam dados intradiários – na estimação do VaR e concluem que não há ganhos significativos na estimação do VaR diário com base em modelos de volatilidade realizada.

Boudoukh *et al.* (2004) propõem uma estimativa chamada de MaxVaR para cálculo de VaR para longos horizontes de tempo. MaxVaR é definida como a perda máxima esperada de uma carteira **até** um horizonte de tempo especificado, e não apenas **naquele** horizonte de tempo, como é definido na metodologia tradicional de VaR. Por meio da modelagem do retorno como um processo de difusão contínuo lognormal, os autores mostram que o valor do MaxVaR será sempre igual ao valor do VaR sob um nível de significância dividido por dois. Para estimar o MaxVaR de uma carteira sob um nível de significância de 1%, por exemplo, deve-se estimar o VaR da mesma carteira sob um nível de significância de 0,5%.

Kaplanski e Levy (2010) propõem uma medida corretiva à regra da raiz quadrada do tempo proposta na modelagem Riskmetrics, de Longerstae y e Spencer (1996). O erro derivado da regra da raiz quadrada do tempo é positivo para curtos horizontes de tempo – provocando superestimação do VaR – e negativo para longos horizontes de tempo – provocando subestimação do VaR. Embora o erro seja relativamente pequeno para carteiras conservadoras e curtos horizontes de tempo, ele é bastante significativo para carteiras agressivas e longos horizontes de tempo.

Observa-se que os trabalhos relativos à estimação do VaR para horizontes superiores a um dia são bastante escassos na literatura de finanças empíricas. Grande parte dos artigos trata da previsão de volatilidade e de VaR para um dia, desprezando o detalhamento de previsões de perda para maiores horizontes de tempo, fundamentais

na prática das instituições financeiras sob um arcabouço regulatório que demanda tais estimativas.

Ainda que a metodologia de VaR seja coerente com as melhores práticas de modelagem de séries temporais financeiras, o próprio conceito de VaR é questionado por muitos autores a respeito de sua eficácia como instrumento de gestão de risco financeiro. O próprio Comitê do Sistema Financeiro Global (CGFS) do Banco de Compensações Internacionais (BIS) critica a adoção da metodologia de VaR para gestão de risco de mercado por sua incapacidade de funcionar bem sob flutuações extremas de preços [CGFS (1999)]. Isto é, o modelo falha quando mais se precisa dele, em situações de crise. Por isso, Artzner *et al.* (1997; 1999) propõem uma metodologia alternativa conhecida como Expected Shortfall (ES). O ES é definido como a perda esperada condicionada em um retorno abaixo do VaR previsto. No entanto, toda a regulação derivada de Basileia ainda é baseada em metodologias tradicionais de VaR, e esse é o motivo de focar os esforços deste trabalho no estudo da métrica mais acurada possível para estimação do VaR para horizontes superiores a um dia.

O modelo de Luger

Em seminário eletrônico promovido pelo Professional Risk Management International Association (PRMIA), Luger (2013) apresenta uma forma de estimação de VaR para horizontes superiores a um dia sem a utilização da regra da raiz quadrada do tempo. A regra da raiz quadrada do tempo reside na afirmação de que o VaR e/ou a volatilidade de uma determinada carteira de ativos “H” passos a frente é igual ao VaR e/ou à volatilidade da carteira um passo a frente multiplicado pela raiz quadrada de “H”. Assumindo volatilidade constante no período entre “T+1” e “T+H” e definindo “ σ ” como a volatilidade da série de retornos, “R” como o próprio

retorno, “P” como o valor da carteira e “ Φ_{α}^{-1} ” como o quantil da distribuição normal-padrão no nível de significância “ α ”, obtém-se:

$$\begin{aligned}\sigma_{T+1:T+H} &= \sigma_{T+1} \sqrt{H} & R_{t+1} &= \log(P_{t+1}) - \log(P_t) = \sigma \varepsilon_{t+1}, \varepsilon_{t+1} \sim i.i.d. N(0,1) \\ VaR_{T+1:T+H}^{\alpha} &= VaR_{T+1}^{\alpha} \sqrt{H} & R_{T+1:T+H} &= \sum_{h=1}^H R_{T+h} = \sum_{h=1}^H \sigma \varepsilon_{T+h}\end{aligned}$$

$$VaR_{T+1:T+H}^{\alpha} = \sqrt{H} \times VaR_{T+1}^{\alpha} = -\sqrt{H} \times \sigma \times \Phi_{\alpha}^{-1}$$

A regra da raiz quadrada de “H” é válida, para a volatilidade, independentemente de qual seja a distribuição dos retornos, mas é necessário que eles sejam independentes e identicamente distribuídos. No entanto, para que seja válida para o VaR, é necessário que a distribuição incondicional dos retornos seja normal, o que nunca acontece na prática, uma vez que um dos fatos estilizados mais conhecidos da literatura de econometria financeira é o excesso de curtose da série de retornos financeiros.

Uma possibilidade para fugir da utilização da regra da raiz quadrada de “H” seria a modelagem dos retornos de longo prazo diretamente, da seguinte forma:

$$R_{t+1:t+H} = \sum_{h=1}^H R_{t+h} = \sigma_{t+1} \varepsilon_{t+1} + \sigma_{t+2} \varepsilon_{t+2} + \dots + \sigma_{t+H} \varepsilon_{t+H}$$

O lado direito da equação representa um processo estocástico do tipo *moving average* com número de períodos de defasagem igual a “H – MA(H)”. O processo “MA” apresenta correlação serial nas observações sobrepostas (*overlapping observations*), o que invalida a hipótese de independência entre os retornos. Por outro lado, o uso de informações não sobrepostas implica redução do tamanho da amostra. Além disso, o processo de agregação dos retornos no tempo altera suas propriedades dinâmicas. Retornos diários apresentam aglomerados de volatilidade bem mais expressivos do que retornos mensais.

Luger (2013) propõe, então, que a estimativa de VaR “H” passos a frente seja obtida por meio de técnicas de simulação e/ou reamostragem, fazendo inferência apenas sobre a distribuição dos retornos um passo a frente. A distribuição dos retornos “H” passos a frente é obtida pelo somatório das variáveis aleatórias simuladas um passo a frente. Este trabalho pretende não apenas implementar o modelo de Luger (2013) para séries temporais financeiras brasileiras, mas também estabelecer uma metodologia quantitativa de *backtesting* baseada na aderência do processo estocástico escolhido à série utilizada – como o teste de Berkowitz – e não apenas na observação do número de perdas maiores do que o VaR – como o teste de Kupiec. Além disso, disponibilizam-se, no Apêndice II, códigos em “R” para definição das funções utilizadas no artigo. Assim, é possível chamar tais funções para estimar o VaR e/ou realizar testes de aderência para qualquer série temporal univariada de forma extremamente simples.

Neste artigo, fixou-se “H” em 63 dias (úteis) com o objetivo de alinhamento à Circular 3.648 do Banco Central do Brasil, de 4 de março de 2013, que trata das abordagens *internal ratings-based* (IRB). Tais abordagens constituem uma metodologia de apuração da parcela de capital necessária para cobertura do risco de crédito com base em sistemas internos de classificação de risco. Como a carteira de ações classificada na carteira bancária deve compor parcela de capital de risco de crédito, é necessário que a modelagem de VaR esteja alinhada à mencionada circular. Entre outras determinações, ela estabelece que a metodologia VaR deve ser aplicada aos retornos trimestrais do valor das ações. Ainda que a instituição financeira não estime seu capital regulatório com base em modelos internos, as instituições sistemicamente importantes (SIFI) devem apurar capital econômico com base em modelos internos gerenciais para verificar sua adequação ao respectivo capital regulatório, no âmbito do processo interno de avaliação da adequação de capital (Icaap).

O nível de significância utilizado neste trabalho para a previsão do VaR será de 1%, por ser considerado o padrão da indústria e disciplinado pelo Banco Central do Brasil na Circular 3.646/2013 para a construção de modelos internos de risco de mercado.

Análise dos dados e apresentação das seções

Os dados utilizados para este estudo foram obtidos no *website* da BM&F Bovespa.³ São retornos diários do Ibovespa no período entre 2 de janeiro de 2006 e 25 de novembro de 2013, representando 1.959 observações. Separaram-se os dados entre dentro e fora da amostra de forma dinâmica. O período dentro da amostra inicialmente é composto de dados entre 2 de janeiro de 2006 e 19 de novembro de 2009. Conforme se percorreram as datas fora da amostra, incluíram-se todas as observações passadas àquela data na série dentro da amostra. Assim, dispõe-se sempre do máximo de informações disponíveis para construir-se a previsão de VaR “H” passos a frente. Ao todo, foram realizadas 938 estimações fora da amostra. No Apêndice II, apresentou-se o código na linguagem de computação “R” responsável pelo processo de estimação.

Observou-se, nos histogramas da primeira seção do Apêndice I, que a distribuição dos retornos diários não muda muito quando se observa apenas o período dentro ou fora da amostra e todos os dados possíveis. No entanto, a distribuição dos retornos trimestrais muda drasticamente quando se altera a amostra utilizada. Tal fato corrobora a modelagem dos retornos trimestrais por meio de simulação, uma vez que é bastante difícil encontrar uma distribuição analítica que se ajuste às várias formas que eles assumem ao longo do tempo, ainda que se observe sua distribuição condicional por meio da padronização pela volatilidade prevista.

³ <<http://www.bmfbovespa.com.br/>>.

Vale destacar que o padrão de dependência temporal identificado na série de retornos diários estudada está em linha com a literatura. Embora não haja dependência linear significativa, existe um padrão de dependência não linear caracterizado pela significativa autocorrelação da série dos quadrados dos retornos – como pode ser visto no Apêndice I –, o que caracteriza os aglomerados de volatilidade presentes nesse tipo de série.

A segunda seção será destinada à exposição dos processos estocásticos – GARCH e APARCH – e das distribuições de probabilidade – normal--padrão, *t-student* e *t-student* assimétrica. Na seguinte, serão descritas duas metodologias para estimação do VaR: uma que utiliza a regra da raiz quadrada do tempo e outra obtida por simulação de Monte Carlo. Os testes de aderência de Kupiec e Berkowitz são apresentados na quarta seção e os respectivos resultados na quinta. Dentre as conclusões mencionadas na sexta seção, destacam-se a melhor aderência da distribuição condicional dos retornos diários às distribuições analíticas quando comparados aos retornos trimestrais, a consequente superioridade da metodologia de VaR para longos horizontes de tempo obtida por simulação de Monte Carlo e os riscos de se confiar cegamente em testes baseados unicamente na ocorrência de falhas, como o teste de Kupiec.

Escolha dos processos estocásticos e distribuições de probabilidade condicionais

Processos estocásticos

GARCH

O processo estocástico GARCH (p,q), proposto por Bollerslev (1986), é considerado o padrão da indústria para a modelagem de volatilidade de séries financeiras. Sua especificação é a que segue:

$$GARCH(p, q)$$

$$R_t = \mu + \varepsilon_t$$

$$\varepsilon_t = h_t u_t, \quad u_t \sim D(.)$$

$$h_t^2 = \omega + \sum_{i=1}^p \alpha_i \varepsilon_{t-i}^2 + \sum_{j=1}^q \beta_j \sigma_{t-j}^2$$

Ela tem as seguintes restrições nos parâmetros:

$$\omega > 0, \quad \alpha_i \geq 0 \quad \forall i = 1, \dots, p$$

$$\beta_j \geq 0 \quad \forall j = 1, \dots, q$$

No modelo proposto por Bollerslev (1986), “D(.)” é uma distribuição de probabilidade normal-padrão. No entanto, neste artigo, serão utilizadas outras distribuições analíticas além da gaussiana para a modelagem de D(.).

APARCH

Existem vários outros modelos derivados do ARCH proposto por Engle (1982). Ding, Granger e Engle (1993) propuseram uma variação bastante interessante, chamada de *asymmetric power* ARCH ou APARCH.

$$APARCH(p, q)$$

$$R_t = \mu + \varepsilon_t$$

$$\varepsilon_t = h_t u_t, \quad u_t \sim D(.)$$

$$h_t^\delta = \omega + \sum_{i=1}^p \alpha_i \left(|\varepsilon_{t-i}| - \gamma_i \varepsilon_{t-i} \right)^\delta + \sum_{j=1}^q \beta_j h_{t-j}^\delta$$

Ela tem as seguintes restrições nos parâmetros:

$$\omega > 0, \quad \delta \geq 0, \quad \alpha_i \geq 0 \quad \forall i = 1, \dots, p$$

$$-1 < \gamma_i < 1 \quad \forall i = 1, \dots, p$$

$$\beta_j \geq 0 \quad \forall j = 1, \dots, q$$

No modelo proposto por Ding, Granger e Engle (1993), $D(\cdot)$ também é uma distribuição de probabilidade normal-padrão. Contudo, neste artigo, como já informado, serão utilizadas outras distribuições analíticas além da gaussiana para a modelagem de $D(\cdot)$.

Por meio da realização de estudos de Monte Carlo, Ding, Granger e Engle (1993) concluíram que, embora os modelos da família ARCH propostos até então fossem capazes de capturar o padrão de dependência não linear existente em séries de retorno financeiro, havia espaço para melhorar a caracterização de tal dependência. Tanto o modelo GARCH, proposto por Bollerslev (1986), como o modelo GARCH de valor absoluto, proposto por Taylor (1986), são capazes de reproduzir alguns padrões de dependência dessas séries. Logo, não existe razão para não utilizar a hipótese de que a variância condicional seja função linear dos retornos defasados ao quadrado [Bollerslev (1986)] ou de que a volatilidade seja função linear dos retornos defasados em módulo [Taylor (1986)]. Assim, Ding, Granger e Engle (1993) propõem um modelo mais geral capaz de reproduzir, com algumas restrições, os modelos de Bollerslev (1986), Taylor (1986) e outros cinco processos estocásticos usualmente citados na literatura de séries temporais financeiras:

(i) ARCH (p), proposto por Engle (1982), com as seguintes restrições:

$$\delta=2, \gamma_i = 0 \quad i=1,\dots,p, \beta_j = 0 \quad \forall j=1,\dots,q$$

(ii) GARCH (p,q), proposto por Bollerslev (1986), com as seguintes restrições:

$$\delta=2, \gamma_i = 0 \quad \forall i=1,\dots,p$$

(iii) GARCH, proposto por Taylor (1986), com as seguintes restrições:

$$\delta=1, \gamma_i = 0 \quad \forall i=1,\dots,p$$

(iv) Glosten-Jagannathan-Runkle (GJR) GARCH, proposto por Glosten, Jagannathan e Runkle (1989), com as seguintes restrições:

$$\delta=2, 0 \leq \gamma_i < 1 \quad \forall i=1,\dots,q$$

(v) *Threshold autoregressive conditional heteroskedasticity* (TARCH), proposto por Zakoian (1994), com as seguintes restrições:

$$\delta = 1, \beta_j = 0 \quad \forall j=1, \dots, q$$

(vi) *Nonlinear autoregressive conditional heteroskedasticity* (NARCH), proposto por Higgins e Bera (1992), com as seguintes restrições:

$$\gamma_i = 0 \quad \forall i=1, \dots, p, \quad \beta_j = 0 \quad \forall j=1, \dots, q$$

(vii) *Logarithmic autoregressive conditional heteroskedasticity* (Log-ARCH), proposto por Geweke e Porter-Hudak (1983), com a seguinte restrição:

$$\delta \rightarrow 0 \text{ (caso-limite)}$$

Entre outras conclusões, o trabalho de Ding, Granger e Engle (1993) evidencia a propriedade de memória longa nas séries de retorno financeiro ao encontrar autocorrelação positiva e significativamente diferente de zero para “ $|r_t|^d$ ” ($d > 0$). Além disso, para um *lag* fixo “ τ ”, a função “ $\rho_t(d) = \text{corr}(|r_t|^d, |r_t + \tau|^d)$ ” tem um único ponto de máximo quando “ d ” está próximo de um, o que vai contra a assunção de outros modelos da família ARCH que trabalham com retornos quadráticos. No artigo desses autores, o valor estimado para o parâmetro “ δ ” foi de 1,43 para dados do índice S&P 500, significativamente diferente do modelo de Taylor (1986), em que “ δ ” = 1, e do modelo de Bollerslev (1986), em que “ δ ” = 2. Isto é, o modelo ideal estaria entre esses dois processos estocásticos, invalidando a opção de se utilizar qualquer um deles. Neste trabalho, compara-se a *performance* dos modelos GARCH (1,1) – por ser o padrão da indústria – e do APARCH (1,1) – por ser reconhecidamente mais flexível e capaz de capturar padrões de dependência que o modelo GARCH não é capaz de capturar.

Distribuições de probabilidade condicionais

Neste estudo, optou-se por utilizar, além da normal-padrão, distribuições *t-student* e *t-student* assimétrica, proposta por Fernández e

Steel (1998), para a distribuição condicional dos retornos. A distribuição *t-student* assimétrica tem a capacidade de capturar, além do típico excesso de curtose de séries de retorno financeiro, alguma assimetria que possa existir para tal série.

De acordo com Lambert e Laurent (2001), o processo estocástico “ z_t ” tem distribuição *t-student* assimétrica, isto é, “ $z_t \sim \text{SKST}(0, 1, \xi, v)$ ”, se, e somente se:

$$f(z_t | \xi, v) = \begin{cases} \frac{2}{\xi + \frac{1}{\xi}} \text{sg}[\xi(sz_t + m)|v], & \text{se } z_t < -\frac{m}{s} \\ \frac{2}{\xi - \frac{1}{\xi}} \text{sg}[(sz_t + m)/\xi|v], & \text{se } z_t \geq -\frac{m}{s} \end{cases}$$

O termo “ $g(. | v)$ ” representa a função de densidade de probabilidade de uma *t-student* (simétrica); e “ ξ ” é o coeficiente de assimetria. Os parâmetros “ m ” e “ s^2 ” são, respectivamente, a média e a variância da distribuição *t-student* assimétrica. O parâmetro “ ξ ” é responsável por modelar a assimetria, enquanto o parâmetro “ v ” modela o excesso de curtose da distribuição.

$$m = \frac{\Gamma\left(\frac{v-1}{2}\right) \sqrt{v-2}}{\sqrt{\pi} \Gamma\left(\frac{v}{2}\right)} \left(\xi - \frac{1}{\xi}\right)$$

$$s^2 = \left(\xi^2 + \frac{1}{\xi^2} - 1\right) - m^2$$

Utilizar distribuições não gaussianas para a modelagem da distribuição condicional dos retornos é muito comum na literatura de séries temporais financeiras. Para citar alguns artigos, há os trabalhos de Giot e Laurent (2004), Lambert e Laurent (2001) e Barndorff-Nielsen e Shephard (2001). Ainda que se utilizem mo-

delos um pouco mais complexos como o APARCH, é muito difícil obter resíduos normalmente distribuídos.⁴ Ainda existe excesso de curtose e assimetria. Por conta disso, efetua-se o teste de Berkowitz para saber qual distribuição ajusta-se melhor ao erro do modelo.

Técnicas de estimação do VaR para horizontes superiores a um dia

Forma analítica

Nessa metodologia, calcula-se a volatilidade “H” passos a frente como função da volatilidade um passo a frente. A volatilidade estimada pela regra da raiz quadrada do tempo é uma metodologia bem mais simples do que a proposta por Luger (2013) e utiliza premissas não verificadas na realidade, conforme relatado na primeira seção.

$$\sigma_{T+H} = \sqrt{H} \cdot \sigma_{T+1}$$

Depois de obter a volatilidade, basta multiplicá-la pelo quantil da distribuição condicional dos retornos e somar a média para se chegar à medida de VaR. No caso da *t-student* e da *t-student* assimétrica, utilizam-se os parâmetros estimados de graus de liberdade e coeficiente de assimetria iguais aos estimados para a série de retornos diários na maximização da verossimilhança como *proxy*. Vai se notar que a distribuição dos retornos trimestrais padronizados

⁴ Em modelos de volatilidade realizada, em que a volatilidade é estimada com base em dados intradiários, é possível obter uma distribuição condicional mais próxima da normal-padrão quando a padronização ocorre pela volatilidade integrada, que é uma *proxy* da verdadeira volatilidade do dia. No entanto, quando a padronização é feita pela volatilidade realizada estimada um passo a frente, ainda é difícil obter uma série sem excesso de curtose e sem assimetria. Para mais informações, ver Fernandes (2008).

em nada se assemelha a qualquer dessas distribuições. Ainda que a volatilidade fosse estimada de outra forma, seria muito difícil, para não dizer impossível, obter uma distribuição dos retornos trimestrais padronizados parecida com as distribuições analíticas mencionadas.

$$\frac{R_{T+H} - \mu_H}{\sigma_{T+H}} \sim D(.)$$

$$VaR_{T+H}^\alpha = \mu_H + \sigma_{T+H} \cdot D^{-1}(\alpha)$$

Em que:

R_{T+H} = retornos “H” passos a frente

μ_H = média incondicional dos retornos “H” passos a frente

σ_{T+H} = volatilidade “H” passos a frente

$D(.)$ = distribuição dos retornos padronizados “H” passos a frente

$D^{-1}(\alpha)$ = quantil da distribuição $D(.)$ avaliado no nível de significância α

VaR_{T+H}^α = Value at Risk estimado “H” passos a frente com nível de significância α

Simulação de Monte Carlo

A técnica de estimação de VaR “H” passos a frente proposta por Luger (2013) é descrita a seguir para o caso do modelo GARCH (1,1) e distribuição condicional dos retornos normal-padrão. O processo estocástico GARCH (1,1) é especificado da seguinte forma:

$$R_{t+1} = \mu + \sigma_{t+1} \varepsilon_{t+1}, \text{ sendo } \varepsilon_{t+1} \sim i.i.d. N(0,1)$$

$$\sigma_{t+1}^2 = \omega + \alpha_1 (R_t - \mu)^2 + \beta_1 \sigma_t^2$$

1) Gere (“N” x “H”) variáveis aleatórias normais-padrão “ $\tilde{\varepsilon}_{i,t+h}$ ”, em que $i = 1, \dots, N$ e $h = 1, \dots, H$. “N” é o número de iterações na simulação de Monte Carlo e “H” é o número de passos a frente utilizado na previsão do VaR.

2) Simule a distribuição do retorno um passo a frente da seguinte forma:

$$\tilde{R}_{i,T+1} = \mu + \sigma_{T+1} \tilde{\varepsilon}_{i,T+1} \text{ para } i = 1, \dots, N$$

3) De posse dos retornos simulados, atualize as variâncias de acordo com a equação recursiva do modelo:

$$\tilde{\sigma}_{i,T+2}^2 = \omega + \alpha_1 (\tilde{R}_{i,T+1} - \mu) + \beta_1 \sigma_{T+1}^2$$

4) Por meio dos retornos simulados a seguir, obtenha os H-ésimos retornos simulados:

$$\tilde{R}_{i,T+1:T+H} = \sum_{h=1}^H \tilde{R}_{i,T+h}, \text{ para } i = 1, \dots, N$$

5) Estime o VaR para o H-ésimo dia da seguinte forma:

$$VaR_{T+1:T+H}^{\alpha} = \text{Quantile}^{\alpha} \left(\left\{ \tilde{R}_{i,T+1:T+H} \right\}_{i=1}^N \right)$$

Destaca-se que a inferência a respeito da distribuição dos retornos é feita apenas para a previsão um passo a frente. A distribuição “H” passos a frente é obtida por meio do somatório de retornos simulados um passo a frente. Esse é um detalhe que faz toda a diferença. Retornos diários padronizados apresentam frequência de observação muito mais semelhante a distribuições analíticas do que retornos trimestrais padronizados. Para estimar o VaR de posse dos retornos trimestrais obtidos por processo de simulação, basta obter o quantil desses dados simulados.

Para mudar a distribuição condicional dos retornos, basta alterar o passo (1). Em vez de gerar variáveis aleatórias normais-padrão para “ $\varepsilon_{i,T+h}$ ”, basta gerar ocorrências da distribuição desejada. Para alterar o processo estocástico gerador dos retornos simulados, basta alterar o passo (3). No caso do APARCH (1,1), a equação da volatilidade é especificada da seguinte forma:

$$\tilde{\sigma}_{i,T+2} = \left\{ \omega + \alpha_1 \left(\left| \tilde{R}_{i,T+1} - \mu \right| - \gamma_1 (\tilde{R}_{i,T+1} - \mu) \right)^\delta + \beta_1 \sigma_{i+1}^\delta \right\}^{1/\delta}$$

Testes de aderência

Com o objetivo de avaliar a *performance* das diferentes metodologias de VaR, utilizaram-se dois testes presentes na literatura: teste de Cobertura Incondicional, proposto por Kupiec (1995) e de Aderência da Distribuição Analítica aos Dados, proposto por Berkowitz (2001).

Teste de Cobertura Incondicional de Kupiec (1995)

O teste de Cobertura Incondicional é o mais difundido no mercado e é baseado na frequência de ocorrências de eventos de perda que excedem o VaR estimado. Seja “x” o número de falhas, isto é, o número de ocorrências de eventos de perda que excedem o VaR em uma amostra de tamanho “n”. Se o modelo de VaR estiver correto e se as ocorrências de falhas forem independentes, então “x” segue uma distribuição binomial de parâmetros “n” e “p”. Sob a hipótese nula, o modelo de previsão é correto e a frequência observada de falhas é consistente com o nível de significância utilizado no modelo. Trata-se de um teste de razão de verossimilhança com a seguinte estatística de teste:

$$LR = -2 \log[(1 - p^*)^{n-x} (p^*)^x] + 2 \log \left[\left(1 - \left[\frac{x}{n} \right] \right)^{n-x} \left(\frac{x}{n} \right)^x \right]$$

O termo “p*” denota a probabilidade de falhas sob a hipótese nula, “n” é o tamanho da amostra e “x” é o número de falhas na amostra. Sob a hipótese nula, a probabilidade de ocorrência de falha (“p”) é igual ao nível de significância (“p*”) do modelo e a estatística de teste tem distribuição qui-quadrado com um grau de liberdade:

$$LR_{uc} \sim X_1$$

A principal crítica ao teste de Kupiec (1995) reside no fato de ele se basear única e exclusivamente na frequência de falhas, sem fazer qualquer tipo de análise em relação à aderência da distribuição aos dados ou à dependência temporal da série de ocorrências de falhas, assumindo até que os eventos de falha sejam independentes.

Teste de Aderência da Distribuição Analítica aos Dados de Berkowitz (2001)

O teste de Berkowitz (2001) é, sem dúvida, o teste mais interessante presente neste artigo. Ele não observa a ocorrência de falhas; sequer necessita do valor estimado do VaR. Ele vai além. O objetivo do teste de Berkowitz é checar se a distribuição condicional dos retornos utilizada na modelagem está de fato aderente aos dados observados ou não.

Para isso, é preciso inicialmente criar uma série de retornos padronizados da seguinte forma:

$$y_t = \frac{R_t - \mu}{\sigma_t}$$

Em que:

y_t = retorno padronizado no tempo “t”

μ = média incondicional do processo estocástico

σ_t = volatilidade prevista para o tempo “t” de acordo com o modelo utilizado

Para a realização do teste proposto por Berkowitz (2001), é necessário transformar a série de retornos padronizados em uma série “ z_t ”, tal que:

$$z_t = \Phi^{-1}(\hat{F}(y_t))$$

Em que:

ϕ^{-1} = quantil da distribuição normal-padrão

$\hat{F}(\cdot)$ = função de distribuição acumulada da previsão

y_t = retorno padronizado realizado no tempo “t”

Berkowitz (2001) usa resultados provados por outros autores para afirmar que “ z_t ” tem distribuição normal-padrão independentemente de qual seja a função “ $\hat{F}(\cdot)$ ”, desde que o modelo esteja corretamente especificado. Isto é, o teste de Berkowitz nada mais é do que a verificação se “ z_t ” segue uma distribuição próxima da normal-padrão. Quanto mais próxima a série “ z_t ” estiver de uma normal-padrão, mais próxima de zero estará a estatística de teste e, consequentemente, mais favorável será o resultado do teste à hipótese nula. A hipótese nula reside na afirmação de que “ $\hat{F}(\cdot)$ ” caracteriza fielmente os dados observados.

Como o objetivo deste estudo é fazer o teste de aderência para uma metodologia de VaR, que foca nas realizações de perdas, vai se realizar o teste truncado em algumas regiões da cauda esquerda da variável aleatória transformada “ z_t ”, conforme definição da variável “Q” a seguir:

$$z_t^* = \begin{cases} Q, & \text{se } z_t \geq Q \\ z_t, & \text{se } z_t < Q \end{cases}$$

$$Q = \phi^{-1}(\alpha)$$

A função de verossimilhança é definida como:

$$L(\mu, \sigma | z^*) = \sum_{z_t^* < Q} \left(-\frac{1}{2} \ln(2\pi\sigma^2) - \frac{1}{2\sigma^2} (z_t^* - \mu)^2 \right) + \sum_{z_t^* = Q} \left(\ln \left(1 - \Phi \left(\frac{Q - \mu}{\sigma} \right) \right) \right)$$

A estatística de teste é baseada na diferença entre as funções de verossimilhança restrita e irrestrita:

$$LR_{tail} = -2(L(0,1) - L(\hat{\mu}, \hat{\sigma}))$$

$$LR_{tail} \sim X^2(2) \text{ sob } H_0$$

A implementação do teste de Berkowitz se dá de forma diferente para a modelagem por meio de simulação de Monte Carlo e para aquela que utiliza a regra da raiz quadrada do tempo. Para a modelagem com simulação, testa-se se os retornos diários padronizados têm frequência de observação semelhante à distribuição analítica utilizada, uma vez que a distribuição trimestral é construída pela soma de retornos diários simulados. Então, obtém-se:

$$\hat{y}_t = \frac{R_t - \hat{\mu}}{\hat{\sigma}_t}$$

Em que:

y_t = retorno diário padronizado no tempo “t”

μ = média incondicional do processo estocástico

σ_t = volatilidade prevista para o tempo “t” de acordo com o modelo utilizado para a série de retornos diários

Já para a modelagem que utiliza a regra da raiz quadrada do tempo, a assunção que se faz é a de que os retornos trimestrais têm algum tipo de distribuição analítica, e a estimação de VaR passa pela obtenção de um quantil dessa distribuição. Por isso, o teste de Berkowitz para essa segunda metodologia de modelagem passa pela construção da variável transformada “ z_t ” com série de retornos trimestrais padronizados. Dessa forma, obtém-se:

$$\hat{y}_t = \frac{R_t - \hat{\mu}}{\hat{\sigma}_t}$$

Em que:

y_t = retorno trimestral padronizado no tempo “t”

μ = média incondicional do processo estocástico multiplicada por “H”

σ_t = volatilidade prevista para o tempo “t” de acordo com a regra da raiz quadrada de “H”

Em ambos os casos, são utilizados apenas dados fora da amostra para a construção de “ y_t ”.

Resultados obtidos para a série de retornos financeiros

Utilizam-se três tipos de modelagem para a distribuição condicional dos retornos: normal, *t-student* e *t-student* assimétrica. Combinam-se tais distribuições com os processos estocásticos GARCH (1,1) e APARCH (1,1), conforme discutido na segunda e na terceira seções.

Apresenta-se, na sequência, tabela comparativa dos modelos utilizados. No teste de Berkowitz, a coluna “Corte” diz respeito aos níveis de significância utilizados no truncamento, uma vez que se está focando na aderência da cauda, e não de toda a distribuição, aos dados. Para cada valor de corte e para cada modelo, obteve-se um p-valor diferente, conforme evidenciado na tabela. No teste de Kupiec, o percentual de falhas indica a fração de dias em que o retorno trimestral foi inferior ao VaR estimado e “p-valor” indica o p-valor encontrado para determinado modelo.

O baixo percentual de falhas obtido pela modelagem por simulação de Monte Carlo se deve primordialmente ao fato de os retornos aberrantes de 2008 estarem presentes na faixa de dados dentro da amostra. Isto é, estimaram-se parâmetros com dados estressados e estes foram testados em um período de relativa calma, uma vez que não se observou frequência semelhante de perdas a partir de 2009 quando comparada a 2008.

No teste de Berkowitz para os modelos que utilizam simulação de Monte Carlo na construção da previsão “H” passos a frente,

isto é, com a transformação “ z_t ” realizada sobre a série de retornos diários padronizados, verifica-se que o processo estocástico APARCH (1,1) com distribuição condicional *t-student* teve p-valor superior a 23% para todos os níveis de corte. Nenhum outro modelo foi tão aderente para vários níveis diferentes de corte. O processo estocástico GARCH (1,1) com distribuição condicional *t-student* assimétrica teve p-valores incríveis para os cortes de 1%, 5% e 10%: todos acima de 90%.

Entre os modelos que utilizam a regra da raiz quadrada de “H” para a previsão da volatilidade, o único que apresentou p-valor significativo foi o APARCH com distribuição condicional normal, mas apenas para o nível de corte de 1%. No nível de corte subsequente, de 5%, o p-valor ficou cravado em zero, o que demonstra pouca aderência da frequência dos retornos trimestrais padronizados pelo GARCH com distribuição normal a essa distribuição – o que pode ser visto facilmente na subseção “Retornos trimestrais – regra da raiz quadrada de ‘H’” do Apêndice I.

Vale destacar que não se encontrou p-valor significativo no teste de Kupiec na estimação combinada com simulação de Monte Carlo, mas sim na estimação com a regra da raiz quadrada do tempo. Tal fato evidencia o problema de se utilizar o teste de Cobertura Incondicional de Kupiec de forma indiscriminada como único fator de escolha. Os modelos APARCH normal, APARCH *t-student* e GARCH *t-student* assimétrica combinados com a regra da raiz quadrada de “H” apresentaram p-valores altíssimos em Kupiec e baixíssimos em Berkowitz. Isto é, mesmo observando que os retornos trimestrais padronizados por essas volatilidades não se parecem nem um pouco com as distribuições analíticas utilizadas, não se pode rejeitar a hipótese nula de Kupiec, ainda que se utilize nível de significância na casa de 25% nos testes.

Tabela 1

Simulação de Monte Carlo (%)

Tabela 1a Berkowitz

Corte	APARCH normal	APARCH <i>t-student</i>	APARCH <i>t-student</i> assimétrica	GARCH normal	GARCH <i>t-student</i>	GARCH <i>t-student</i> assimétrica
1	3,53	87,24	83,55	0,05	83,44	91,11
5	0,14	23,20	0,79	0,50	48,45	93,37
10	1,52	68,20	8,16	0,33	45,57	97,96
15	1,81	62,09	8,64	2,50	3,60	11,45
20	5,72	49,20	3,92	21,80	14,36	17,75

Tabela 1b Kupiec

	APARCH normal	APARCH <i>t-student</i>	APARCH <i>t-student</i> assimétrica	GARCH normal	GARCH <i>t-student</i>	GARCH <i>t-student</i> assimétrica
Percentual de falhas	0,00	0,00	0,00	0,10	0,32	0,11
p-valor	0,00	0,00	0,00	0,04	1,46	0,04

Fonte: Elaboração própria, com base em valores estimados em R.

Nota: Tabela com os p-valores e frequência de falhas dos testes realizados fora da amostra.

Tabela 2

Regra da raiz quadrada de “H” (%)

Tabela 2a Berkowitz

Corte	APARCH normal	APARCH <i>t-student</i>	APARCH <i>t-student</i> assimétrica	GARCH normal	GARCH <i>t-student</i>	GARCH <i>t-student</i> assimétrica
1	50,68	1,47	0,76	0,00	0,00	0,32
5	0,00	0,00	0,00	0,00	0,00	0,00
10	0,00	0,00	0,00	0,00	0,00	0,00

(Continua)

(Continuação)

Corte	APARCH normal	APARCH <i>t-student</i>	APARCH <i>t-student</i> assimétrica	GARCH normal	GARCH <i>t-student</i>	GARCH <i>t-student</i> assimétrica
15	0,00	0,00	0,00	0,00	0,00	0,00
20	0,00	0,00	0,00	0,00	0,00	0,00

Tabela 2b Kupiec

	APARCH normal	APARCH <i>t-student</i>	APARCH <i>t-student</i> assimétrica	GARCH normal	GARCH <i>t-student</i>	GARCH <i>t-student</i> assimétrica
Percentual de falhas	0,75	1,39	0,32	5,01	4,37	1,28
p-valor	41,35	26,17	1,46	0,00	0,00	40,99

Fonte: Elaboração própria, com base em valores estimados em R.

Nota: Tabela com os p-valores e frequência de falhas dos testes realizados fora da amostra.

Conclusões

Neste artigo, mostrou-se como estimar o VaR de uma carteira para horizontes superiores a um dia de acordo com metodologia proposta por Luger (2013). Em tal metodologia, o VaR “H” passos a frente é estimado pela obtenção de um quantil de uma distribuição simulada por meio do processo estocástico gerador da série. Ela tem significativos ganhos em relação à adoção da regra da raiz quadrada do tempo, que tem premissas não verificadas na realidade e muitas vezes ignoradas na prática da gestão de risco financeiro.

Como se vê nos histogramas presentes nos itens “GARCH (1,1) combinado com simulação de Monte Carlo” e “APARCH (1,1)

combinado com regra da raiz quadrada de ‘H’ do Apêndice I, a série de retornos padronizados “H” passos a frente não tem um comportamento semelhante a qualquer das distribuições analíticas citadas neste artigo. É muito mais fácil aceitar a hipótese nula do teste de Berkowitz de correta especificação do processo gerador da série para retornos padronizados diários do que trimestrais. Assim, quando se gera a distribuição da previsão dos retornos trimestrais por simulação, parte-se de um pressuposto com elevadíssima aceitação por Berkowitz para se construir a série de retornos “H” passos a frente. Verifica-se, de acordo com os histogramas presentes no item “Exemplos de histogramas de retornos trimestrais previstos, gerados pela combinação do processo estocástico APARCH (1,1) *t-student* assimétrica com simulação de Monte Carlo” do Apêndice I, que o aspecto da distribuição simulada muda ao longo do tempo, de forma que é bastante difícil – para não dizer impossível – modelá-la diretamente por meio de função de distribuição analítica e obter resultados tão bons quanto os obtidos por simulação, sobretudo quando se utiliza como parâmetro o p-valor de Berkowitz.

Uma das principais conclusões deste trabalho é a de que a escolha de modelos com base exclusivamente em testes de ocorrência de falhas, como o teste de Kupiec, pode trazer graves problemas. Observaram-se distribuições de retornos padronizados que muito diferem de as respectivas distribuições analíticas receberem p-valores altíssimos em Kupiec e baixíssimos em Berkowitz. O teste de Berkowitz é muito mais abrangente, uma vez que ele checa se a distribuição analítica é, de fato, aderente aos dados.

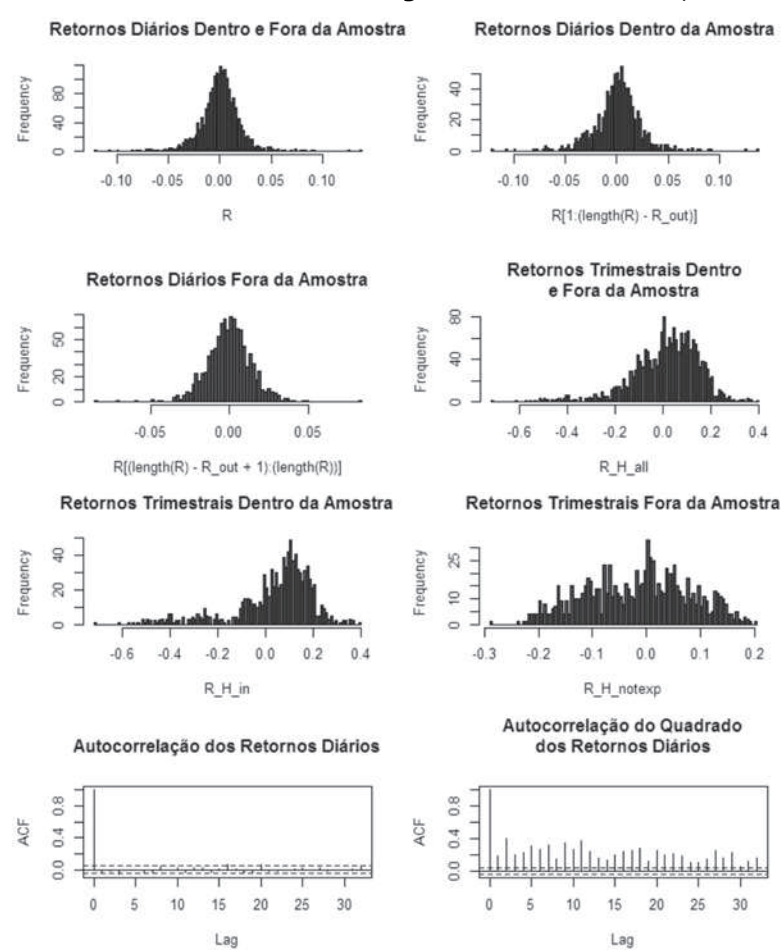
Uma interessante extensão deste trabalho seria a aplicação dos mesmos modelos para o caso multivariado, isto é, em que o VaR da carteira seja decomposto em seus fatores de risco. Tal abordagem é fundamental na análise de sensibilidade do VaR em relação a mu-

danças nas volatilidades e correlações, o que é bastante interessante para a realização de testes de estresse.

Apêndice I

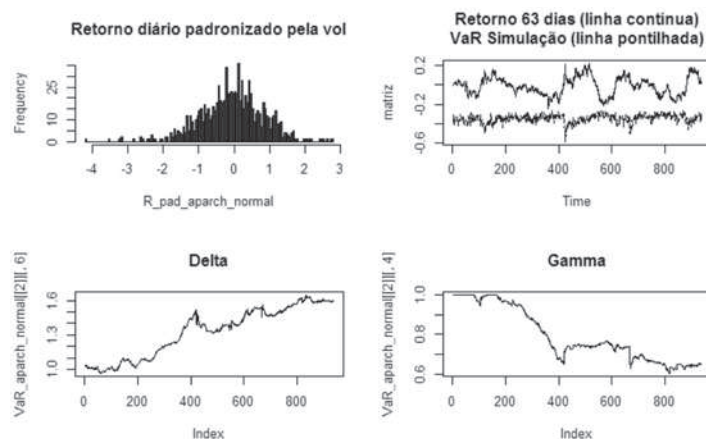
Gráficos

Dados sem tratamento – histogramas e autocorrelações

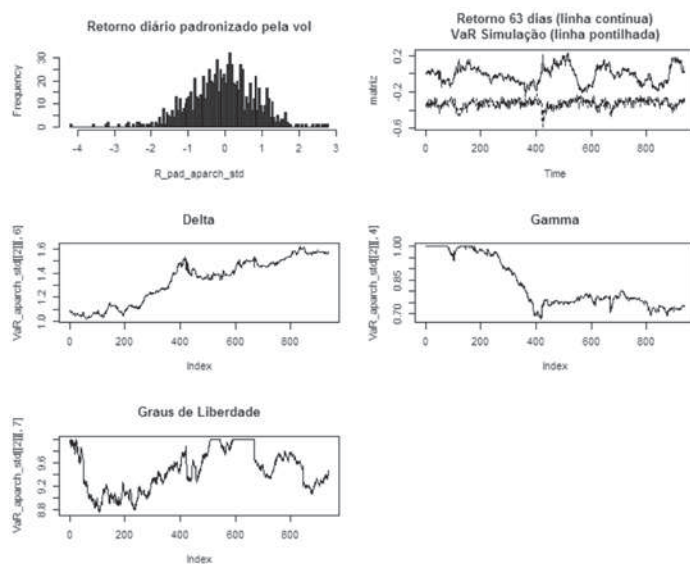


APARCH (1,1) combinado com simulação de Monte Carlo

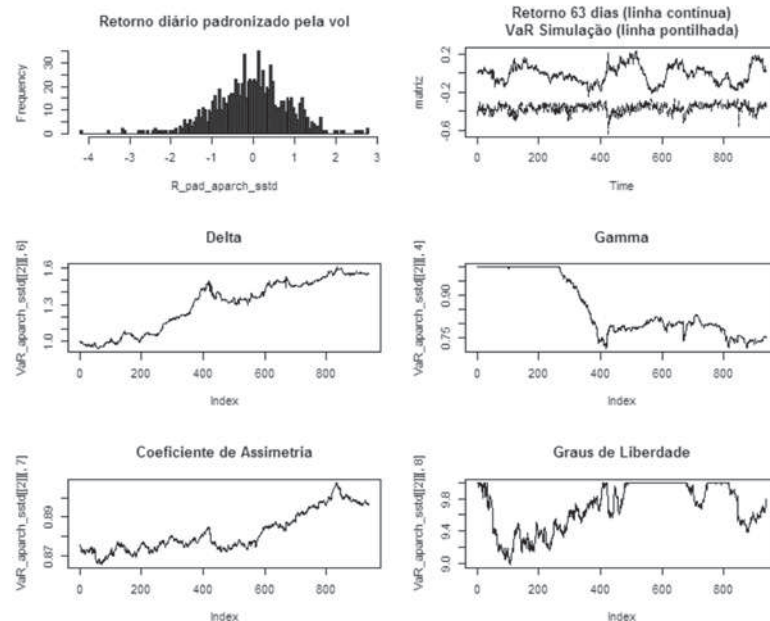
APARCH (1,1) com distribuição condicional normal-padrão



APARCH (1,1) com distribuição condicional *t-student* com número de graus de liberdade estimado a cada passo pela função “GARCHFit”

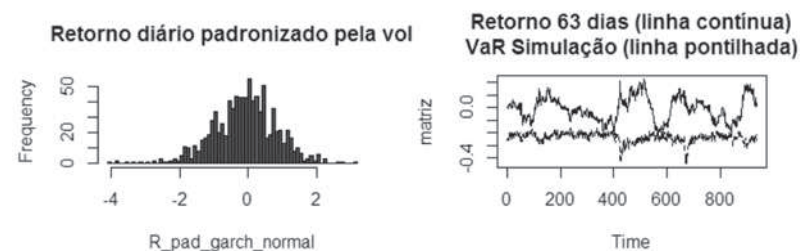


APARCH (1,1) com distribuição condicional *t-student* assimétrica com número de graus de liberdade e coeficiente de assimetria estimados a cada passo pela função “GARCHFit”



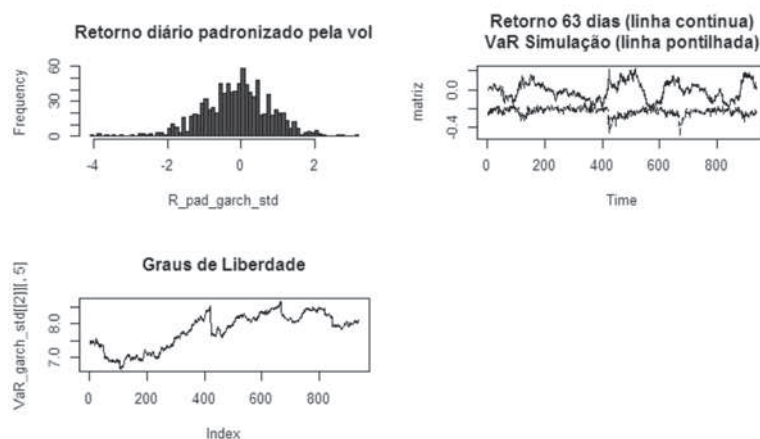
GARCH (1,1) combinado com simulação de Monte Carlo

GARCH (1,1) com distribuição condicional normal-padrão

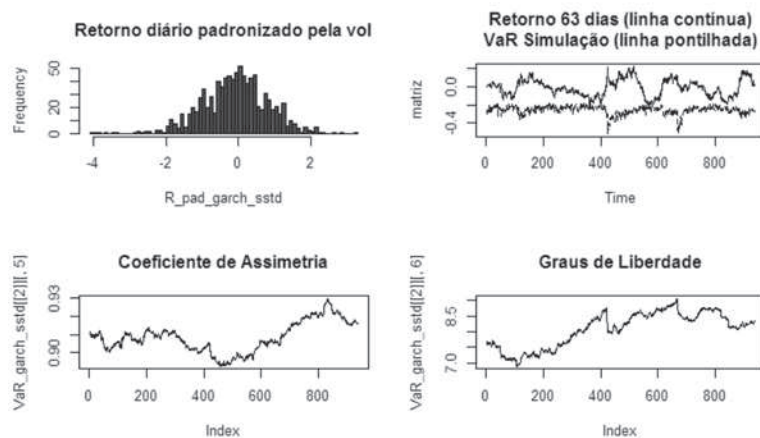


Estimação de Value at Risk para horizontes superiores a um dia por meio dos processos estocásticos GARCH e APARCH combinados com simulação de Monte Carlo

GARCH (1,1) com distribuição condicional *t-student* com número de graus de liberdade estimado a cada passo pela função “GARCHFit”

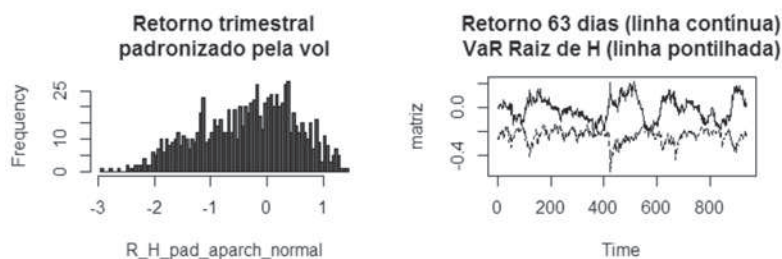


GARCH (1,1) com distribuição condicional *t-student* assimétrica com número de graus de liberdade e coeficiente de assimetria estimados a cada passo pela função “GARCHFit”

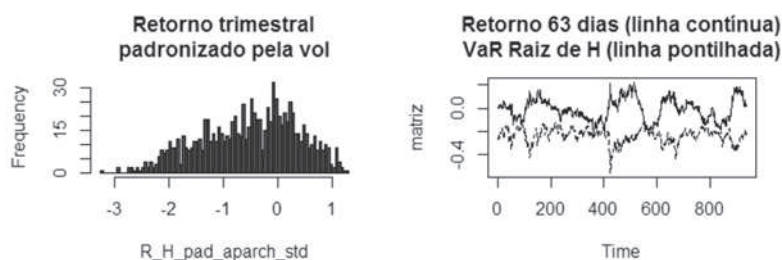


APARCH (1,1) combinado com regra da raiz quadrada de “H”

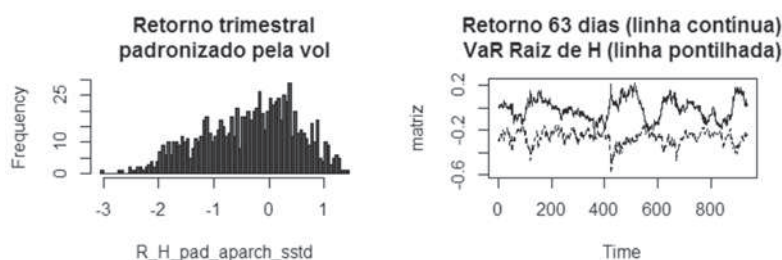
APARCH (1,1) com distribuição condicional normal-padrão



APARCH (1,1) com distribuição condicional *t-student* com número de graus de liberdade estimado a cada passo pela função “GARCHFit”

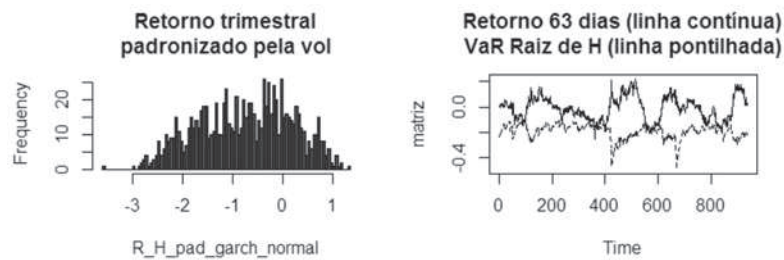


APARCH (1,1) com distribuição condicional *t-student* assimétrica com número de graus de liberdade e coeficiente de assimetria estimados a cada passo pela função “GARCHFit”

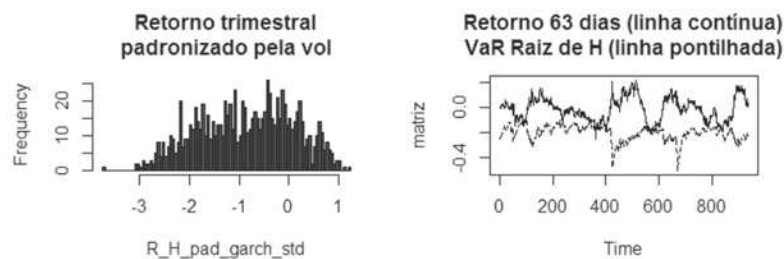


GARCH (1,1) combinado com regra da raiz quadrada de “H”

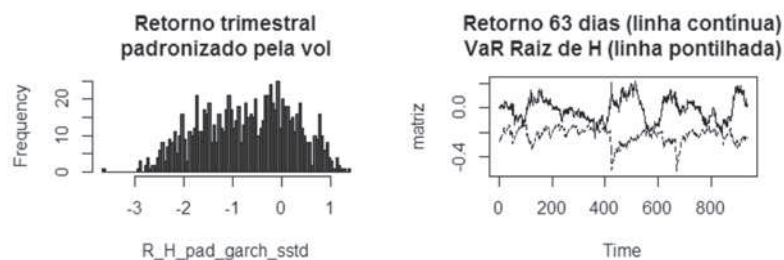
GARCH (1,1) com distribuição condicional normal-padrão



GARCH (1,1) com distribuição condicional *t-student* com número de graus de liberdade estimado a cada passo pela função “GARCHFit”



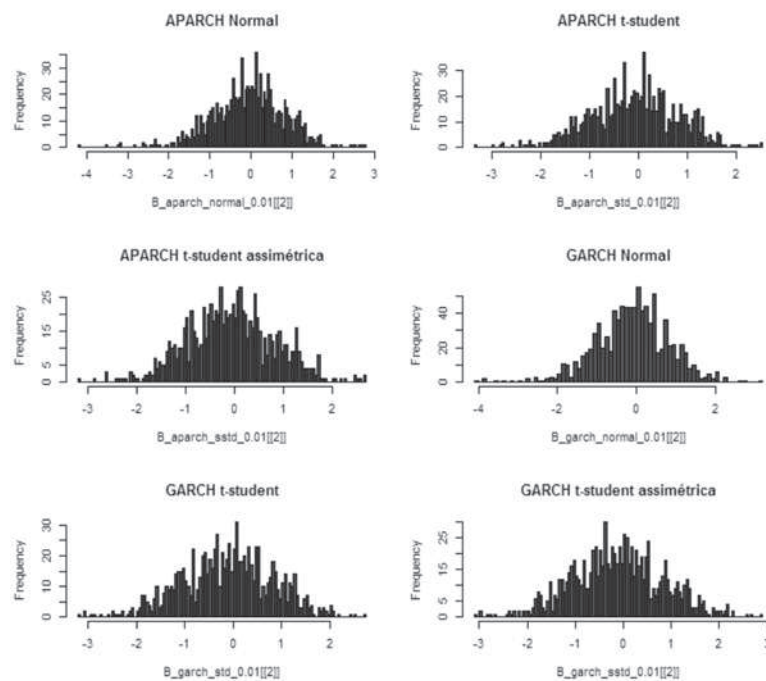
GARCH (1,1) com distribuição condicional *t-student* assimétrica com número de graus de liberdade e coeficiente de assimetria estimados a cada passo pela função “GARCHFit”



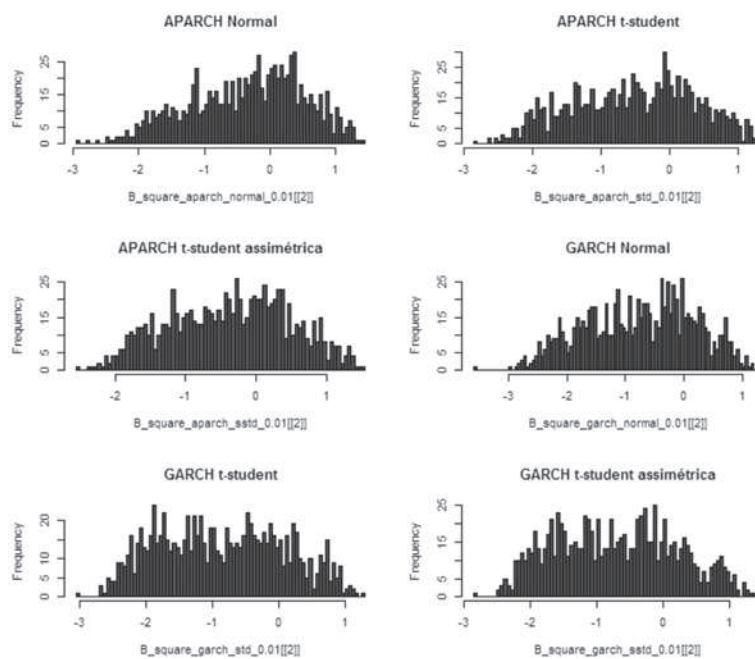
Variável transformada “ z_t ” utilizada no teste de Berkowitz

Obs.: Quanto mais próxima a “ z_t ” estiver de uma distribuição normal-padrão, mais a estatística de teste estará aderente à hipótese nula de correta especificação do modelo.

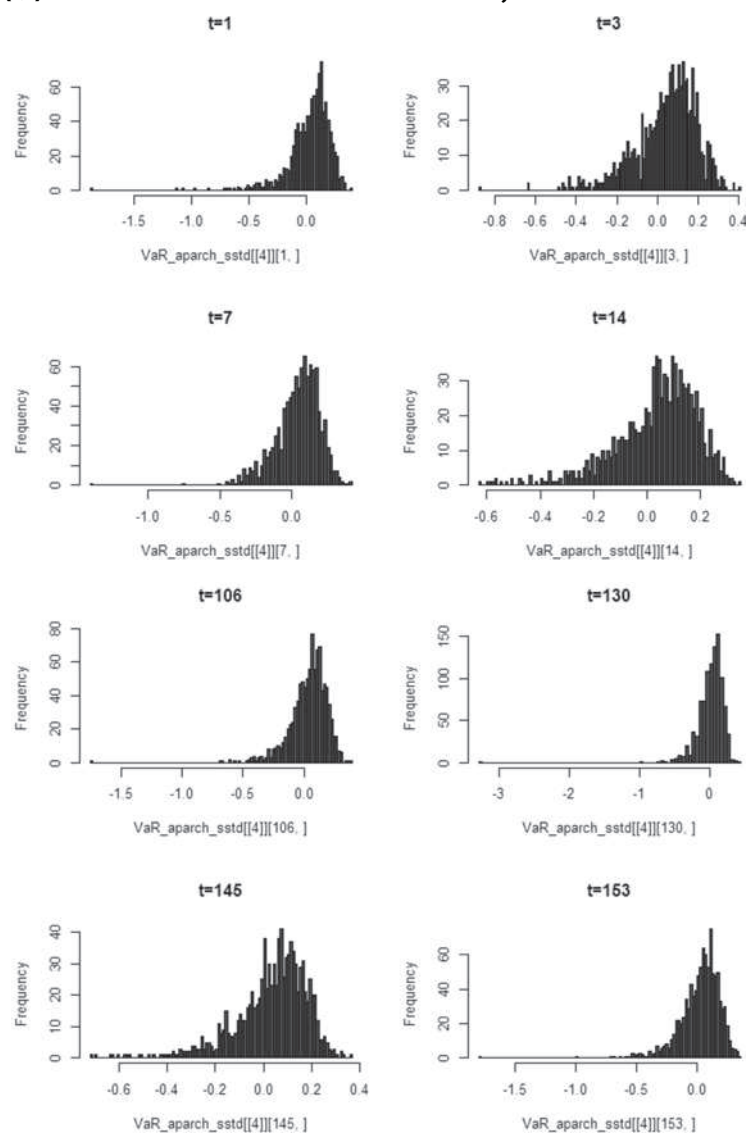
Retornos diários – simulação de Monte Carlo



Retornos trimestrais – regra da raiz quadrada de “H”



Exemplos de histogramas de retornos trimestrais previstos, gerados pela combinação do processo estocástico APARCH (1,1) *t-student* assimétrica com simulação de Monte Carlo



Estimação de Value at Risk para horizontes superiores a um dia por meio dos processos estocásticos GARCH e APARCH combinados com simulação de Monte Carlo

Apêndice II

Códigos em “R” para a implementação dos modelos

Foram utilizados os pacotes estatísticos “fGARCH” e “ruGARCH” na implementação das funções.

1. Função criada para estimar o VaR “H” passos a frente por meio da técnica de simulação de Monte Carlo, proposta por Luger (2013), em região fora da amostra, com o objetivo de se realizarem testes de aderência (*backtesting*)

```
VaR_H = function(N, H, R, R_out, modelo, dist, alpha, exp)
{
  # N é o número de iterações na simulação de Monte Carlo
  # H é o número de passos a frente na estimação do VaR
  # R é a série histórica completa
  # R_out é o tamanho da série fora da amostra
  # modelo é o processo estocástico utilizado: ~GARCH (1,1) ou
  ~APARCH (1,1)
  # dist é a distribuição condicional dos retornos: “norm” (normal-
  -padrão) , “std” (t-student) ou “sstd” (t-student assimétrica)
  # alpha é o nível de significância adotado na previsão do VaR
  # se exp==TRUE, então a função retorna a transformação exponencial
  do VaR. Caso contrário, não há tal transformação
  # R_in é o tamanho da série dentro da amostra no momento inicial,
  uma vez que a cada passo novas informações são incorporadas à
  amostra
  R_in = length(R) - R_out
  # Z é o número de estimações realizadas
  Z = R_out-(H-1)
  # R_sim é a matriz de retornos simulados um passo a frente
```

```

R_sim = array(data=0 , dim=c(N,H))
# sigma_sim é a matriz de volatilidades atualizadas
sigma_sim = array(data=0 , dim=c(N,H))
# vol é o vetor de volatilidades estimadas
vol = array(data=0 , dim=c(Z))
# VaR é o vetor de medidas de VaR estimadas
VaR = array(data=0 , dim=c(Z))
# o tamanho da matriz de coeficientes depende do modelo e da
distribuição. Quatro é o número mínimo de coeficientes (modelo
GARCH com distribuição condicional normal)
length_coefs = 4
if (modelo=="aparch(1,1)")
    length_coefs = length_coefs + 2
if (dist=="std")
    length_coefs = length_coefs + 1
else if (dist=="sstd")
    length_coefs = length_coefs + 2
# coefs é a matriz de coeficientes das estimações
coefs = array(data=0.0 , dim=c(Z,length_coefs))
if (dist=="norm")
    # Epsilon_sim é um vetor aleatório Normal-Padrão de ta-
    manho N x H
    Epsilon_sim = rnorm(N*H)
# R_H_sim a distribuição de retornos simulados "H" passos a frente
para cada observação "z", isto é, para cada observação fora da
amostra
R_H_sim = array(data=0, dim=c(Z,N))

```

```

# z é o contador do número de estimações
for (z in 1:Z)
{
# “garchFit” retorna a estimação dos parâmetros realizada por
modelo da família “ARCH”
fit = garchFit(formula = modelo , data=R[1:(R_in+z-1)] , cond.
dist=dist , trace=FALSE)
# Construção dos valores estimados dos parâmetros para cada “z”
mu = coef(fit)[[1]]
omega = coef(fit)[[2]]
alpha1 = coef(fit)[[3]]
# Caso APARCH (1,1)
if (modelo==~aparch(1,1)) {
gamma1 = coef(fit)[[4]]
beta1 = coef(fit)[[5]]
delta = coef(fit)[[6]]
if (dist=="std")
shape = coef(fit)[[7]]
else if (dist=="sstd") {
skew = coef(fit)[[7]]
shape = coef(fit)[[8]] } }
# Caso GARCH (1,1)
else if (modelo==~garch(1,1)) {
beta1 = coef(fit)[[4]]
gamma1 = 0.0
delta = 2.0
if (dist=="std")
shape = coef(fit)[[5]]

```

```

else if (dist=="sstd") {
    skew = coef(fit)[[5]]
    shape = coef(fit)[[6]] } }
if (dist=="std")
    # Epsilon_sim é um vetor aleatório t-student de tamanho N
    x H, com número de graus de liberdade (nu) estimado
    Epsilon_sim = rstd(n=N*H , nu=shape)
else if (dist=="sstd")
    # Epsilon_sim é um vetor aleatório t-student assimétrico de
    tamanho N x H, com número de graus de liberdade (nu) e
    coeficiente de assimetria (xi) estimados
Epsilon_sim = rsstd(n=N*H , nu=shape , xi=skew)
# "predict" retorna a previsão um passo a frente da vol
vol[z] = predict(fit, n.ahead = 1)[[3]]
# i é o contador do número de iterações na simulação de Monte
Carlo
for (i in 1:N)
{
    # h é o contador do número de passos a frente na estimação do VaR
    for (h in 1:H)
    {
        if (h==1)
            sigma_sim[i,h] = vol[z]
        else
            sigma_sim[i,h] = (omega + alpha1*((abs(R_sim[i,h-1]-mu) -
            gamma1*(R_sim[i,h-1]-mu))^delta) + beta1*(sigma_
            sim[i,h-1]^delta))^(1/delta)
    }
}

```

```

R_sim[i,h] = mu + sigma_sim[i,h] * Epsilon_sim[i + (h-1)*N]
R_H_sim[z,i] = R_H_sim[z,i] + R_sim[i,h]
}
}
VaR[z] = quantile(R_H_sim[z,], probs=alpha)
coefs[z,] = coef(fit)
}
# se exp==TRUE, então a função retorna a transformação
exponencial do VaR
if (exp==TRUE)
    VaR = exp(VaR) - 1
# a função retorna quatro elementos: o vetor de previsões do VaR, a
matriz de coeficientes estimados, o vetor de volatilidades estimadas
e a matriz com realizações da previsão do retorno “H” passos a
frente para cada dia fora da amostra
lista = list(VaR = VaR, Coeficientes = coefs, Volatilidade = vol,
R_H_sim = R_H_sim)
return(lista)
}

```

Função criada para estimar o VaR “H” passos a frente por meio da regra da raiz quadrada de “H” em região fora da amostra, com o objetivo de se realizarem testes de aderência (*backtesting*)

```

VaR_square = function(H, mu, vol, dist, skew, shape, alpha, exp)
{

```

```

# q é um vetor de quantis utilizados na estimação do VaR
q = array(data=0.0, dim=c(length(vol)))
for (i in 1:length(q))
    if (dist=="norm")
        q[i] = qnorm(p=alpha)
    else if (dist=="std")
        q[i] = qstd(p=alpha, nu=shape[i])
    else if (dist=="sstd")
        q[i] = qsstd(p=alpha, xi=skew[i], nu=shape[i])
# estimação do VaR "H" passos a frente pela regra da raiz de "H"
VaR = mu*H + sqrt(H)*vol*q
if (exp==TRUE)
    VaR = exp(VaR)-1
return(VaR)
}

```

Função criada para realizar o teste de Kupiec

```

Kupiec = function(R_H, VaR, alpha)
{
# R_H representa o retorno realizado "H" passos a frente
# VaR representa o VaR calculado "H" passos a frente
# n é o número de observações fora da amostra
n = length(R_H)
# x é o contador de dias em que o retorno é inferior ao VaR
x=0.0
for (i in 1:n)
    if (R_H[i] < VaR[i])
        x = x + 1.0
}

```

```

# LR é a estatística de teste
LR = (-2)*log( ((1-alpha)^(n-x)) * (alpha^x) ) +
2*log( ((1-(x/n))^(n-x)) * ((x/n)^x) )
# sob H0, LR tem distribuição qui-quadrado com um grau de
liberdade. Logo, o p-valor é definido como a região à direita da
estatística de teste da função de distribuição acumulada da qui-
quadrado com um grau de liberdade. A função “pchisq” retorna
P(X>x) quando lower.tail=FALSE. “q” é o quantil da distribuição e
“df” é o número de graus de liberdade
p_valor = pchisq(q=LR, df=1, lower.tail=FALSE)
# a função retorna o p-valor do teste e o número de falhas “x”
return(list(p_valor=p_valor, x=x))
}

```

Função criada para realizar o teste de Berkowitz

```

Berkowitz = function(R_pad, dist, shape, skew, corte)
{
# R_pad é a série de retornos padronizados
# dist é a distribuição da previsão: “norm”, “std” ou “sstd”
# shape é o vetor de números de graus de liberdade estimados
# skew é o vetor de coeficientes de assimetria estimados
# z é a variável transformada utilizada como entrada para o teste
z = array(data=0.0 , dim=c(length(R_pad)))
# será aplicado o quantil da distribuição normal-padrão sobre a
distribuição acumulada do retorno padronizado
if (dist=="norm")
    z = qnorm(pnorm(R_pad))

```

```

else
    for (i in 1:length(R_pad))
        if (dist=="std")
            z[i] = qnorm(pstd(R_pad[i], nu=shape[i]))
        else
            z[i] = qnorm(psstd(R_pad[i], nu=shape[i],
                               xi=skew[i]))
# B retorna o resultado da função BerkowitzTest (pacote "ruGARCH")
# aplicada sobre a série transformada "z". A função BerkowitzTest
# utiliza a série "z" para calcular a estatística de teste e o p-valor do
# teste de Berkowitz. Quando tail.test=TRUE, o teste é aplicado apenas
# sobre a cauda esquerda da distribuição, com precisão determinada
# pelo parâmetro "alpha"
B = BerkowitzTest(data=z, tail.test=TRUE, alpha=corte, lags=0)
lista=list(B, z=z)
return(lista)
}

```

Referências

ARTZNER, P. *et al.* Thinking coherently. *Risk*, v. 10, n. 11, p. 68-71, 1997.

_____. Coherent measures of risk. *Mathematical Finance*, v. 9, n. 3, p. 203-228, 1999.

BANCO CENTRAL DO BRASIL. *Circular 3.646*, de 4 de março de 2013.

Estabelece os requisitos mínimos e os procedimentos para o cálculo, por meio de modelos internos de risco de mercado, do valor diário referente à parcela RWA_{MINT} dos ativos ponderados pelo risco (RWA), de que trata

a Resolução 4.193, de 1º de março de 2013, e dispõe sobre a autorização para uso dos referidos modelos.

_____. *Circular 3.648*, de 4 de março de 2013. Estabelece os requisitos mínimos para o cálculo da parcela relativa às exposições ao risco de crédito sujeitas ao cálculo do requerimento de capital mediante sistemas internos de classificação do risco de crédito (abordagens IRB) (RWA_{CIRB}), de que trata a Resolução 4.193, de 1º de março de 2013.

BARNDORFF-NIELSEN, O. E.; SHEPHARD, N. Non-Gaussian Ornstein-Uhlenbeck-Based Models and some of their uses in financial economics. *Journal of the Royal Statistical Society*, B, 63(2), p. 167-241, 2001.

BERKOWITZ, J. Testing density forecasts with applications to risk management. *Journal of Business & Economic Statistics*, 19(4), p. 465-474, 2001.

BERKOWITZ, J.; O'BRIEN, J. How accurate are Value at Risk models at commercial banks? *Journal of Finance*, 57, p. 1093-1112, 2002.

BOLLERSLEV, T. Generalized autoregressive conditional heteroskedasticity. *Journal of Econometrics*, 31, p. 307-327, 1986.

BOUDOUKH, J. *et al.* MaxVaR: Long horizon Value at Risk in a mark-to-market environment. *Journal of Investment Management*, v. 2, n. 3, p. 14-19, 2004.

CGFS – COMMITTEE ON THE GLOBAL FINANCIAL SYSTEM. *A review of financial market events in autumn 1998*. BIS, 1999.

DANIELSSON, J. The emperor has no clothes: limits to risk modelling. *Journal of Banking and Finance*, 26, p. 1273-1296, 2002.

DING, Z.; GRANGER, C. W. J.; ENGLE, R. F. A long memory property of stock market returns and a new model. *Journal of Empirical Finance*, 1, p. 83-106, 1993.

ENGLE, R. Autoregressive conditional heteroscedasticity with estimates of the variance of United Kingdom inflation. *Econometrica*, 50, p. 987-1007, 1982.

FERNANDES, G. *Volatilidade realizada: previsão e aplicações a medidas de Value at Risk*. Monografia (Graduação em Ciências Econômicas) – Pontifícia Universidade Católica do Rio de Janeiro, Rio de Janeiro, 2008.

FERNÁNDEZ, C.; STEEL, M. On Bayesian modelling of fat tails and skewness. *Journal of the American Statistical Association*, 93, p. 359-371, 1998.

GEWEKE, J.; PORTER-HUDAK, S. The Estimation and Application of Long Memory Time Series Models. *Journal of Time Series Analysis*, 4, p. 221-238, 1983.

GIOT, P.; LAURENT, S. Modelling daily Value at Risk using realized volatility and ARCH type models. *Journal of Empirical Finance*, 11, p. 379-398, 2004.

GLOSTEN, L.; JAGANATHAN, R.; RUNKLE, D. *Relationship between the expected value and the volatility of the nominal excess return on stocks*. Unpublished manuscript. J. L. Kellogg Graduate School, Northwestern University, 1989.

HIGGINS, M.; BERA, A. A class of nonlinear ARCH models. *International Economic Review*, 33, 1992.

KAPLANSKI, G.; LEVY, H. The two-parameter long-horizon Value at Risk. *Frontiers in Finance and Economics*, v. 7, n. 1, p. 1-20, 2010.

KUPIEC, P. Techniques for verifying the accuracy of risk measurement models. *Journal of Derivatives*, 2, p. 73-84, 1995.

LAMBERT, P.; LAURENT, S. *Modelling financial time series using GARCH-type models and a skewed student density*. Mimeo. Université de Liège, 2001.

LEE, T. H.; SALTOGLU, B. *Evaluating predictive performance of Value at Risk models in emerging markets: a reality check*. Working paper. Marmara University, 2001.

LONGERSTAEY, J.; SPENCER, M. *Riskmetrics – technical document*. 4. ed. New York: J.P. Morgan/Reuters, 1996.

LUGER, R. Managing Financial Risk over Long Horizons.

In: PROFESSIONAL RISK MANAGEMENT INTERNATIONAL ASSOCIATION WEBINAR.

19 jun. 2013. Disponível em: <<http://www.prmia.org/civicrm/event/info?id=4107&reset=1>>. Acesso em: 29 mai. 2014.

TAYLOR, S. *Modelling financial time series*. New York: Wiley, 1986.

ZAKOIAN, J. M. Threshold heteroskedasticity models. *Journal of Economic Dynamics and Control*, 15, p. 931-955, 1994.