

Problemas de otimização estocástica em horizonte periódico

Colaboração ONS – UFRJ

**Bernardo Freitas Paulo da Costa,
Ricardo Turano Figueiredo,
Anderson de Oliveira Calixto**

UFRJ – Universidade Federal do Rio de Janeiro
Departamento de Matemática Aplicada – Instituto de Matemática
Av. Athos da Silveira Ramos, 149
Centro de Tecnologia, Bloco C - Cidade Universitária
Caixa Postal 68530 – 21941-909 – Rio de Janeiro – RJ – Brasil

**Débora Dias Jardim Penna,
Alessandra Mattos Ramos de Oliveira,
Lucas de Souza Khenayfis**

ONS – Operador Nacional do Sistema Elétrico
Rua Júlio do Carmo, 251 - Cidade Nova
20211-160 – Rio de Janeiro – RJ – Brasil

Dezembro 2022

Conteúdo

Sumário	4
1 Introdução	6
I Fundamentação metodológica	8
2 Formulação periódica	8
2.1 Modelo simplificado de planejamento da operação de longo prazo	8
2.1.1 Função objetivo	8
2.1.2 Equação de balanço de energia hidráulica	9
2.1.3 Equação de balanço de carga	9
2.1.4 Limites de variáveis	10
2.1.5 Notação do modelo de planejamento	10
2.2 Modelo em horizonte infinito	11
2.3 Caso periódico	13
2.4 Caso pré-periódico	16
2.5 O fator de desconto	16
2.6 Generalizações do modelo	17
3 Algoritmo de solução	18
3.1 Definições preliminares	19
3.2 Recursão de programação dinâmica	19
3.2.1 Caso geral: período p	22
3.3 Modificações na PDDE	23
3.3.1 Convergência e etapa <i>Forward</i>	23
3.3.2 Construção de cortes e etapa <i>Backward</i>	25
3.4 Compartilhamento de cortes	26
4 Regime estacionário	26
4.1 Definições	26
4.2 Existência	28
4.3 Convergência	29
II Implementação e Resultados	30
5 Implementação	30
5.1 Desenvolvimento do protótipo	30
5.1.1 Adaptação para horizonte periódico	30
5.1.2 Migração do código para Julia 1.6	30
5.1.3 Outras modificações e melhorias	31
5.2 Ambiente de programação	31
5.3 Contribuições à comunidade Julia	33

6 Prova de conceito	34
6.1 Casos de estudo	34
6.2 Regime anual	34
6.2.1 Análise do estado estacionário	42
6.2.2 Comparando o estabelecimento do regime estacionário	45
6.3 Comparação dos regimes semestrais	53
6.4 Casos de estágios mensais	62
7 Resultados	71
7.1 Casos de estudo	71
7.2 Comparação periódico e não-periódico	71
7.3 Comparação com aversão ao risco	76
7.4 Comparação entre subsistemas	81
7.5 Construção de cortes e etapa <i>forward</i>	82
8 Conclusão	89
8.1 Propriedades do modelo periódico	89
8.2 Comparação com modelos de horizonte finito	90
8.3 Estabilidade numérica	91
8.4 Trabalhos futuros	91
A Envoltórias da função de custo futuro	92
A.1 Envoltória inferior	92
A.2 Envoltória superior	93
B Cadeias de Markov	93
B.1 Cadeias de Markov com espaço de estados discretos	94
B.2 Distribuição de probabilidade invariante para uma Cadeia de Markov	96
B.2.1 Cadeias de Markov irreduzíveis	99
B.2.2 Cadeias de Markov recorrentes e transientes	100
B.2.3 Cadeias de Markov periódicas e aperiódicas	103
B.2.4 Recorrência nula e positiva	104
B.3 Cadeias de Markov com espaço de estados contínuo	105
B.3.1 Cadeias de Markov irreduzíveis	107
B.3.2 Átomos e conjunto pequenos	109
B.3.3 Ciclos e períodos	110
B.3.4 Transiência a recorrência	110
B.3.5 Recorrência de Harris	112
B.3.6 Distribuição invariante	112
B.3.7 Convergência para a distribuição invariante	113
C Testes Estatísticos	114
C.1 Testes da Bondade do Ajuste	114
C.2 Processos Empíricos	118
C.3 Testes de hipótese na teoria da bondade do ajuste.	120
C.3.1 Poder de um teste estatístico	122
C.3.2 Critério de decisão via p-Valor	124
C.4 Testes de hipótese quando desejamos comparar duas distribuições.	125

C.5	Testes para uma amostra.	126
C.5.1	Teste de Kolmogorov-Smirnov.	126
C.5.2	Teste de Anderson-Darling para uma amostra.	127
C.6	Testes para duas amostras.	129
C.6.1	Teste de Kolmogorov-Smirnov para duas amostras.	129
C.6.2	Teste de Anderson-Darling para 2-amostras.	130

Sumário

Dando continuidade aos projetos de “*Metodologias de Programação Estocástica Inteira*”, estabelecidos entre o ONS e a COPPETEC/Instituto de Matemática da UFRJ, contratos GMC-CT-093/17 e PEM-CS-200/18, este projeto considerou na modelagem do problema de planejamento da operação a representação de um problema com horizonte periódico. Há duas motivações principais para tal representação: em primeiro lugar, o horizonte periódico evita os *efeitos de bordo* presentes no final do horizonte de planejamento. Em segundo lugar, é esperado que um modelo com horizonte periódico permita resolver problemas com maior acurácia, para um mesmo custo computacional, essencialmente por reaproveitar cortes obtidos em vários *estágios periódicos*.

A primeira parte deste relatório é metodológica, na qual introduzimos os modelos de otimização estocástica em horizonte periódico, e apresentamos uma discussão sobre as hipóteses subjacentes ao modelo de planejamento da operação que justificam sua aplicação. Também contemplaremos modelos pré-periódicos, que flexibilizam algumas destas hipóteses, permitindo uma maior gama de aplicações.

Em seguida, observaremos como as técnicas de programação dinâmica podem ser adaptadas para resolver problemas de horizonte periódico. Focaremos nas diferenças quanto à natureza do problema, mas também nas similaridades que permitem implementar um algoritmo similar à PDDE (Programação Dual Dinâmica Estocástica) para calcular soluções de problemas periódicos.

Finalmente, faremos uma breve discussão quanto ao estabelecimento de um regime estacionário das soluções: em que sentido este regime deve ser interpretado, e algumas de suas propriedades, especialmente relacionadas às noções de *Cadeias de Markov*. Aprofundando a abordagem probabilística, apresentaremos também a teoria de testes de hipóteses não-paramétricos, que são necessários para dar apoio estatístico à convergência - ou não - das soluções ao estado estacionário.

A segunda parte se dedica à avaliação do uso prático desta modelagem, analisando exemplos de problemas de planejamento da operação cada vez mais complexos. O protótipo computacional, desenvolvido através dos projetos supracitados, continuou sendo aprimorado. Notamos, principalmente, as diferentes adaptações necessárias para modelagem e simulação em horizontes periódicos; a sua atualização para as novas versões da linguagem *Julia* e dos pacotes *JuMP* e *SDDP.jl*; uma contribuição ao código do pacote *SDDP.jl* para aumentar a estabilidade numérica ao resolver problemas de grande porte; e a modificação para utilizar o *solver GLPK*, livre e gratuito, por default.

Nos casos de estudo analisados, pudemos observar que o modelo proposto dá resultados muito bons, e consistentes. Primeiramente, notamos que ocorreu a convergência para o regime estacionário em todos eles. A variável principal para esta análise é, sempre, a energia armazenada, que, na maior parte das vezes, é a que mais demora a atingir este patamar. Em alguns casos estudados, este regime se estabelece, de forma clara, somente após mais de 8 anos. Também observamos, em casos simples, que esta convergência não depende da condição inicial: para três casos, o estado estacionário obtido após 8 anos é estatisticamente idêntico, o que sugere que o sistema de fato possui um único estado estacionário, atrator.

O regime estacionário depende do problema em si: o estudo de dois modelos com periodicidade semestral, correspondentes a agregações distintas dos 12 meses do ano, apontou para regimes marcadamente diferentes, o que era esperado.

Mais ainda, ao comparar o modelo usual de 120 estágios mensais com um modelo periódico de 12 estágios, notamos que o modelo periódico produz resultados similares, e sem efeitos de fim de horizonte, com o uso de apenas 12 funções de custo futuro, em vez de 120. Por não ter o efeito de fim de horizonte, o modelo periódico tem um sinal diferente de armazenamento, o que leva a uma operação levemente mais conservadora, e que, de fato, estabiliza após um certo tempo, enquanto

que o modelo de 120 estágios não tem um momento claro de estabilização da energia armazenada: após descontar as variações sazonais, ela sobe nos cinco primeiros anos de operação, e decai nos quatro seguintes.

Ademais, ao utilizar uma mesma função de custo futuro em meses correspondentes, o modelo periódico tem uma operação mais consistente, e com menos flutuação da geração térmica, do que o modelo de 120 estágios. Isso leva, entre outros, a uma melhor distribuição da geração térmica ao longo dos meses, o que gera menos picos de despacho térmico. Para tanto, a energia armazenada é maior e, naturalmente, mantém uma tendência de subida mais acentuada do que o modelo de horizonte finito. Com tudo isso, o modelo periódico tem um custo total de operação menor, por evitar despachar térmicas com alto custo marginal, e também reduz a ocorrência de deficit.

Por fim, também efetuamos comparações com um modelo incorporando aversão ao risco, que é perfeitamente compatível com modelos de horizonte infinito, e produz resultados significativamente diferentes, como esperado, inclusive do ponto de vista do regime estacionário atingido.

1 Introdução

Nos dois projetos anteriores, motivados por questões de modelagem matemática, exploramos e apresentamos diversas técnicas para tratar de problemas de otimização estocástica multi-estágio com variáveis inteiras [FPdCLdFCdC18, dCdFK⁺19]. Em ambos, sempre foi bastante claro que, com o aumento da capacidade de representação do modelo, graças à maior flexibilidade permitida pelas variáveis inteiras, também havia um maior custo computacional para a solução do problema.

Este custo se manifestava de duas formas, que ficaram bastante claras durante o primeiro projeto. A primeira é resultado da maior dificuldade intrínseca para resolver problemas de programação linear inteira-mista, se comparados a problemas de programação linear onde todas as variáveis são contínuas, como tradicionalmente usado para os problemas de planejamento de médio prazo da operação. A segunda é ligada ao método de solução: para aproximar a política ótima de operação, a decomposição temporal subjacente aos algoritmos de programação dinâmica estocástica necessita aproximar as funções de custo futuro. Ora, havendo variáveis de decisão inteiras, é possível que as funções de custo futuro deixem de ser convexas, o que aumenta tanto a dificuldade de aproximar tais funções, como de representá-las computacionalmente para serem usadas no estágio anterior.

Esta observação foi um dos pontos de partida do segundo projeto, que empregou um método, ainda de decomposição temporal, que utiliza apenas aproximações *convexas* das funções de custo futuro. Por não introduzir outras não-convexidades nos problemas de cada estágio, também reduzimos a complexidade dos problemas de programação linear inteira-mista resolvidos a cada iteração do algoritmo. Esta técnica trouxe bons resultados, tanto do ponto de vista do tempo computacional, como do ponto de vista da redução do *gap* de dualidade, fornecendo assim melhores estimativas para os custos de operação.

Ao longo deste segundo projeto, surgiu a ideia de buscar reduzir o custo computacional de algoritmos de programação dinâmica ao reduzir diretamente o número de funções de custo futuro a serem estimadas. Esta ideia também foi motivada pela estrutura periódica do problema de planejamento da operação, cuja sazonalidade mensal é repetida ao longo de um horizonte de 10 anos.

Mais ainda, do ponto de vista matemático, o uso de um horizonte periódico tem duas vantagens. Primeiramente, o horizonte atual de 10 anos, com 5 anos de configuração dinâmica e em seguida 5 anos de configuração estática contém uma escolha arbitrária para o final do horizonte de planejamento. Os 5 anos de configuração estática tem como propósito evitar efeitos de bordo, resultantes da inicialização com uma função de custo futuro nula. Ora, se a configuração é estática, e o processo de incertezas é sazonal, veremos que é possível modelar um problema periódico que está livre dos efeitos de bordo.

Em segundo lugar, o modelo periódico permite introduzir claramente a noção de um estado estacionário do problema de otimização estocástica. A partir da análise do estado estacionário, esperamos poder compreender de forma mais rigorosa os diversos comportamentos observados ao resolver o problema de planejamento, inclusive em horizonte finito, ao transferir resultados, mesmo que de forma qualitativa, entre ambos.

Neste projeto, estudamos como utilizar um modelo *periódico* para resolver o problema de planejamento da operação, assim como suas consequências do ponto de vista do tempo computacional e dos resultados obtidos. O uso de problemas periódicos em otimização é, de certa forma, clássico: no caso em que os estados e decisões são *discretos*, esta abordagem faz parte da área de processos Markovianos de decisão [Ber12]. No âmbito de problemas de otimização estocástica, onde em geral o estado e as decisões são *contínuos*, uma primeira abordagem foi dada por Birge em [BZ07]. Mais recentemente, estas ideias foram incorporadas ao algoritmo de Programação Dinâmica Dual

Estocástica (PDDE), em diversos trabalhos [Dow20, SD20].

Este relatório começa com uma revisão do modelo de planejamento da operação, para introduzir a notação, e daí apresentar a modelagem em horizonte infinito, e depois o caso de horizonte infinito e periódico, que será o foco deste projeto. Após termos estabelecido os modelos utilizados, iremos apresentar os algoritmos utilizados para resolver os problemas de otimização estocástica resultantes, enfatizando as similaridades e diferenças com o algoritmo de programação dual dinâmica estocástica (PDDE) atualmente empregado pelos modelos computacionais que o ONS utiliza para o planejamento da operação. Nesta parte, a noção de *Operador de Bellman*, ligando as diferentes funções de custo futuro, será fundamental tanto para a descrição matemática do modelo como para a construção do algoritmo de solução.

Em seguida, abordaremos o lado probabilístico do modelo, introduzindo a noção de *estado estacionário* para cadeias de Markov, e analisando as questões de convergência e unicidade. Vale ressaltar, desde já, que o cálculo deste estado estacionário não é uma tarefa simples, e por isto apresentaremos alguns *testes estatísticos* que nos permitirão analisar, de forma quantitativa, se houve de fato convergência para um eventual estado estacionário.

Daí, o relatório prossegue para uma parte mais computacional. Em primeiro lugar, faremos um resumo dos aprimoramentos trazidos ao protótipo ao longo deste projeto, assim como uma descrição do ambiente computacional empregado. Com isto, passaremos aos casos de estudo analisados, começando por um modelo simples, com apenas uma variável de estado, que é equivalente ao SIN totalmente agregado em um mercado e um reservatório equivalente de energia (REE), portanto também sem restrições de transmissão. Este caso é suficientemente simples para que possamos analisar em profundidade os resultados, pois haverá menos variáveis de decisão. Além de uma comparação entre o modelo periódico e um modelo de horizonte finito, também faremos comparações com modelos semestrais, para destacar o impacto da sazonalidade no regime estacionário.

Em seguida, abordaremos um caso mais próximo da realidade operativa, onde consideramos os 4 REEs tradicionais (SE, S, NE, N) e os submercados correspondentes. Aí, veremos o impacto, por um lado, das restrições de transmissão, mas também dos limites máximo e mínimo de armazenamento em cada reservatório, na política obtida. Além de uma comparação com o caso de horizonte finito, veremos a importância da implementação *adaptada ao horizonte infinito* das etapas *forward* e *backwards* da PDDE para construir as funções de custo futuro corretas para o problema.

Enfim, também efetuamos uma análise com aversão ao risco, que é compatível com esta metodologia de problemas periódicos, para observar os impactos tanto na política como no estado estacionário.

Enfim, encerramos o relatório elencando trabalhos futuros que podem ser realizados, a partir das conclusões tiradas destes desenvolvimentos.

Parte I

Fundamentação metodológica

2 Formulação periódica

Para apresentarmos a formulação periódica adaptada ao modelo de planejamento da operação, iremos fazer um breve resumo da estrutura fundamental do problema, considerando um modelo simplificado do mesmo. A seguir, destacaremos sob quais condições este problema pode ser reformulado a partir de uma perspectiva de horizonte infinito e periódico. Então, apresentaremos uma formulação mista, e mais geral, composta por um conjunto de estágios periódicos, precedidos de estágios não periódicos, que também são conhecidos como *pré-periódicos*. Este modelo misto, combinando ambos tipos de estágios, permite maior generalidade de aplicações, e é o mais próximo do problema de planejamento da operação tal qual utilizado pelo ONS atualmente.

2.1 Modelo simplificado de planejamento da operação de longo prazo

A formulação em programação dinâmica do modelo de planejamento de longo prazo da operação, considerando a representação agregada das usinas hidrelétricas em reservatórios equivalentes, é dada por:

$$Q_t(v_t, a_t) = \min_{(v_{t+1}, q_t, s_t, g_t, df_t, f_t)} c^\top g_t + c_{df}^\top df_t + \beta \bar{Q}_{t+1}(v_{t+1}) \quad (2.1)$$

s.a

$$v_{t+1} = v_t + a_t - q_t - s_t,$$

$$q_t + M_I g_t + df_t + M_D f_t = d_t,$$

$$0 \leq v_{t+1} \leq \bar{v}, \quad 0 \leq q_t \leq \bar{q}, \quad \underline{g} \leq g_t \leq \bar{g},$$

$$0 \leq f_t \leq \bar{f}, \quad 0 \leq s_t, \quad 0 \leq df_t,$$

$$\bar{Q}_{t+1}(v_{t+1}) = \begin{cases} \mathbb{E}[Q_{t+1}(v_{t+1}, a_{t+1})] & , t \in \{1, \dots, T-1\} \\ 0 & , t = T \end{cases} \quad (2.2)$$

para todo t desde 1 a T . Para cada estágio t , o vetor de variáveis de decisão é $x_t = (v_{t+1}, q_t, s_t, g_t, df_t, f_t)$. Neste modelo simplificado, a única incerteza considerada é a afluência a_t , que é assumida “*stagewise independent*”. Ao longo desta seção, detalharemos as variáveis, restrições e função objetivo deste problema de planejamento.

2.1.1 Função objetivo

A *função objetivo*

$$c^\top g_t + c_{df}^\top df_t + \beta \bar{Q}_{t+1}(v_{t+1})$$

representa a soma dos custos de geração térmica, déficit (corte de carga) e o custo futuro $\bar{Q}_{t+1}(v_{t+1})$ onde:

- c vetor de custo unitário térmico (\$/MWmês);
- g_t vetor de geração térmica no estágio t (MWmês);
- c_{df} vetor de custo unitário do déficit (\$/MWmês);
- df_t vetor de déficit de energia no estágio t (MWmês);
- β fator de desconto;
- $\overline{Q}_{t+1}(\cdot)$ função de custo futuro do estágio $t + 1$ correspondente ao valor esperado do custo total do estágio $t + 1$ em diante (\$);
- v_{t+1} vetor de energia armazenada ao final do estágio t (MWmês);

Em cada estágio a função objetivo é composta por duas parcelas. A primeira é uma função de custo imediato, $c^\top g_t + c_{df}^\top df_t$, que dá o preço a ser pago pela quantidade de energia gerada através de térmicas mais o custo atribuído ao déficit de energia que ocorre naquele estágio. Já a segunda é uma função de custo futuro, $\beta \overline{Q}_{t+1}(\cdot)$, que, para um dado volume armazenado nos reservatórios no *final* do estágio, dá a média dos custos imediatos de todos os próximos estágios. Assim, a decisão no estágio t leva em consideração não apenas os custos atuais, mas também uma estimativa dos custos decorrentes da decisão operativa naquele instante de tempo, descontada por um fator β .

2.1.2 Equação de balanço de energia hidráulica

A equação de *balanço de energia hidráulica* no estágio t é

$$v_{t+1} = v_t + a_t - q_t - s_t,$$

onde

- v_t vetor de energia armazenada no início do estágio t (MWmês);
- a_t vetor de energia afluente durante o estágio t (MWmês);
- q_t vetor de energia turbinada durante o estágio t (MWmês);
- s_t vetor de energia vertida durante o estágio t (MWmês);

O significado dessa restrição é físico: conservação de energia. A quantidade final de energia (ou água) disponível em um reservatório, denotada por v_{t+1} , é igual à quantidade que já estava lá no começo do estágio (v_t) mais quanta energia chegou (a_t) menos o que foi defluído, seja por turbinamento (q_t) ou vertimento (s_t).

2.1.3 Equação de balanço de carga

A equação de *balanço de carga* para todos os subsistemas e estágio t é:

$$q_t + M_I g_t + df_t + M_D f_t = d_t$$

onde

- d_t vetor de carga durante o estágio t (MWmês);
- q_t vetor de energia turbinada durante o estágio t (MWmês);
- M_I matriz indicadora das térmicas associadas a um mesmo subsistema;
- g_t vetor de geração térmica durante o estágio t (MWmês);
- $M_I g_t$ vetor de geração térmica total em cada subsistema (MWmês);
- df_t vetor de déficit de energia durante o estágio t (MWmês);
- M_D matriz de incidência representando as direções de fluxo de energia entre os subsistemas;
- f_t vetor de intercâmbio de energia durante o estágio t , em cada conexão (MWmês);
- $M_D f_t$ vetor de intercâmbio líquido de energia entre os subsistemas (MWmês);

Esta equação corresponde à restrição de caráter físico que nos diz que a soma da geração térmica com a geração hidrelétrica num dado estágio deve ser igual à demanda daquele estágio.

Aqui, isso se traduz na equação de balanço de carga. Ou seja, a carga total d_t para atender a demanda em cada subsistema será igual à soma da energia turbinada naquele estágio q_t , da geração térmica $M_I g_t$ de cada subsistema, da quantidade de déficit no estágio df_t e do intercâmbio líquido de energia entre os subsistemas $M_D f_t$.

2.1.4 Limites de variáveis

Os limites das variáveis de decisão são:

- $0 \leq v_{t+1} \leq \bar{v}$ vetores de limites superior e inferior de energia armazenada (MWmês);
- $0 \leq q_t \leq \bar{q}$ vetores de limites superior e inferior de energia turbinada (MWmês);
- $\underline{g} \leq g_t \leq \bar{g}$ vetores de limites superior e inferior de geração térmica (MWmês);
- $0 \leq f_t \leq \bar{f}$ vetores de limites superior e inferior de intercâmbio de energia (MWmês);
- $0 \leq df_t$ vetor de limites inferior do déficit (MWmês); e
- $0 \leq s_t$ vetor de limites inferior da energia vertida (MWmês).

Para manter a coerência com a realidade, as variáveis relacionadas a volume armazenado ou geração de energia devem ser sempre não-negativas. Além disso, as variáveis podem possuir algum limite superior ou inferior nos valores que podem assumir. Estas restrições permitem modelar os limites reais na produção de energia pelas usinas ou no fluxo possível nas linhas de intercâmbio.

2.1.5 Notação do modelo de planejamento

De forma a enfatizar as diferenças entre variantes deste modelo básico, representaremos as restrições de limite (também conhecidas como restrições “de caixa”) pelo conjunto $\mathcal{X}_t$, e as variáveis pelo vetor x_t :

$$x_t := (v_{t+1}, q_t, s_t, g_t, df_t, f_t) \in \mathcal{X}_t.$$

Sempre que for conveniente, reescrevemos o modelo de planejamento (2.1) como

$$\begin{aligned}
 Q_t(v_t, a_t) = \min_{x_t} \quad & c^\top g_t + c_{df}^\top df_t + \beta \bar{Q}_{t+1}(v_{t+1}) \\
 \text{s.a} \quad & v_{t+1} = v_t + a_t - q_t - s_t, \\
 & q_t + M_I g_t + df_t + M_D f_t = d_t, \\
 & x_t = (v_{t+1}, q_t, s_t, g_t, df_t, f_t) \in \mathcal{X}_t.
 \end{aligned} \tag{2.3}$$

Observe que também omitimos a definição da função de custo futuro $\bar{Q}_{t+1}(v_{t+1})$ como média dos custos $Q_{t+1}(v_{t+1}, a_{t+1})$.

Este modelo pode ser representado graficamente como indicado na figura 2.1. Nele, cada círculo representa um estágio, e as flechas indicam o fluxo do tempo. Optamos por não representar a árvore de cenários completa, porque neste momento vamos focalizar nos estágios e nas funções de custo futuro. E como no caso de incertezas *stagewise independent* há uma função de custo futuro *por estágio*, este gráfico também indica a dependência entre as funções de custo futuro: a função $\bar{Q}_2$ depende de $\bar{Q}_3$, e assim por diante, até que a última ($\bar{Q}_{120}$) não depende de mais nada, o que é indicado pelo símbolo terminal.

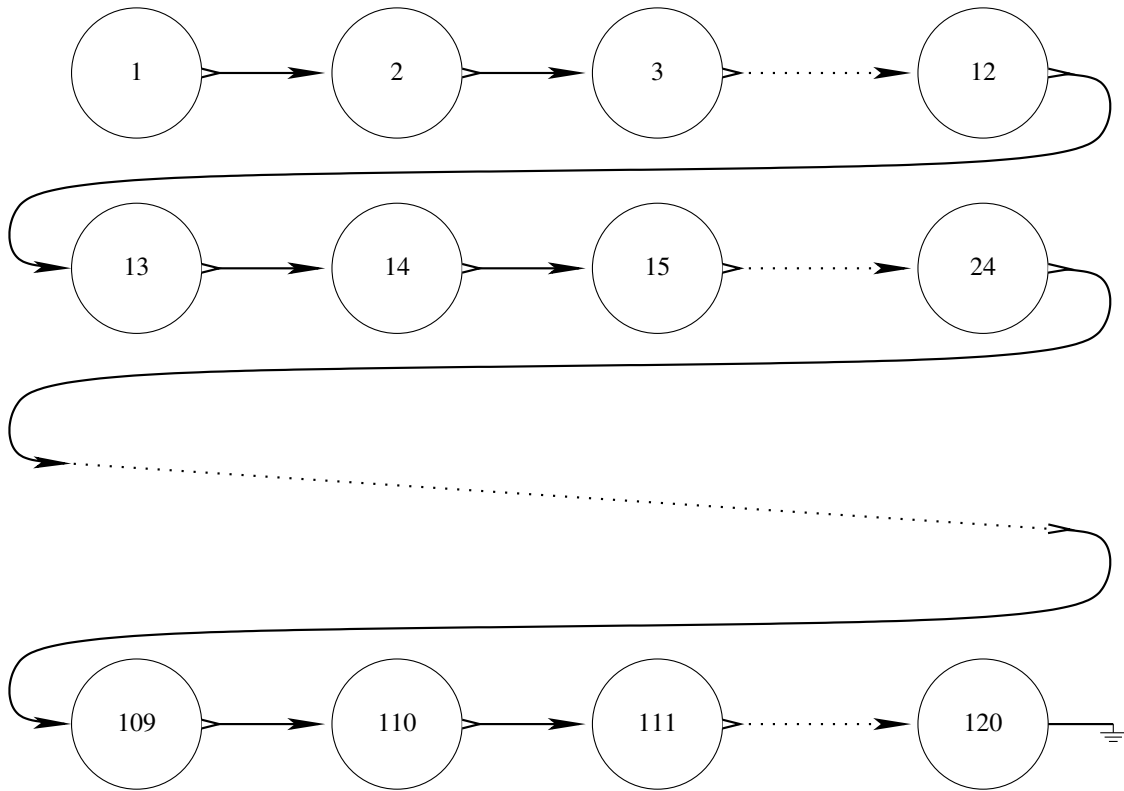


Figura 2.1: Problema de otimização com 120 estágios.

2.2 Modelo em horizonte infinito

Vamos generalizar o modelo acima, de programação dinâmica, para problemas de horizonte infinito: para isso, suprimimos a condição terminal

$$\bar{Q}_{T+1} = 0,$$

e estendemos o problema indefinidamente à frente. Entretanto, não é claro que esta generalização é válida, pois ao incluir infinitos estágios a própria definição das funções de custo futuro depende, sempre, de infinitos estágios.

Para provar que este modelo está bem definido, vamos usar a *formulação extensiva* [SDR09] de problemas de otimização estocástica, incluindo todos os estágios e cenários em um único problema:

$$\begin{aligned} \min_{(v_{t+1}, q_t, s_t, g_t, df_t, f_t)} \quad & \mathbb{E}_\xi \left[\sum_{t=1}^{\infty} \beta^{t-1} \left(c^\top g_t(\xi) + c_{df}^\top df_t(\xi) \right) \right] \\ \text{s.a} \quad & v_{t+1} = v_t + a_t - q_t - s_t, \\ & q_t + M_I g_t + df_t + M_D f_t = d_t, \\ & 0 \leq v_{t+1} \leq \bar{v}, \quad 0 \leq q_t \leq \bar{q}, \quad \underline{g} \leq g_t \leq \bar{g}, \\ & 0 \leq f_t \leq \bar{f}, \quad 0 \leq s_t, \quad 0 \leq df_t, \end{aligned} \quad (2.4)$$

onde enfatizamos a dependência das variáveis de decisão g_t e df_t na realização da incerteza ξ dentro da função objetivo, que está sendo minimizada *em média*. Esta dependência também poderia ser explicitamente incluída nas restrições, por exemplo escrevendo a primeira como

$$v_{t+1}(\xi) = v_t(\xi) + a_t(\xi) - q_t(\xi) - s_t(\xi), \quad \forall t, \forall \xi$$

mas não o faremos, em geral, para simplificar a notação. Como no modelo original das equações (2.1) e (2.2), a única incerteza que levamos em conta são as afluências a_t , e portanto cada cenário ξ deve conter todas as realizações para $t \geq 1$.

Observamos, também, que os custos de um estágio t aparecem multiplicados por um fator β^{t-1} , que corresponde à correção dos preços no tempo t para preços no tempo 1. Ao escrever o modelo de programação dinâmica como nas equações (2.1) e (2.2), há apenas um fator β , pois este transporta os preços do tempo $t+1$ de volta ao tempo t , independente do valor de t . Com isto, notamos que a função de custo futuro $\bar{Q}_{t+1}(v_{t+1})$ representa o custo futuro, *no tempo* $t+1$, da operação também a partir do tempo $t+1$.

Com a formulação extensiva, é mais fácil estabelecer hipóteses para que o problema de horizonte infinito esteja bem posto. A mais comum é que os *custos imediatos* sejam limitados, ou seja, que, para qualquer estado inicial v_t , existam decisões viáveis no tempo t tais que o custo

$$c^\top g_t(\xi) + c_{df}^\top df_t(\xi)$$

seja limitado por alguma constante C , independente do estágio t .

Como consequência direta dessa hipótese, podemos demonstrar a existência de uma cota superior para o custo total da operação, ao longo do tempo. Isso também garante que, apesar de ter infinitos estágios, temos soluções cuja função de custo futuro será finita. De fato, o termo dentro do valor esperado em (2.4) será limitado por

$$\sum_{t=1}^{\infty} \beta^{t-1} (c^\top g_t(\xi) + c_{df}^\top df_t(\xi)) \leq \sum_{t=1}^{\infty} \beta^{t-1} C = C \frac{1}{1-\beta}, \quad (2.5)$$

onde usamos que $\beta < 1$ para poder realizar a soma da progressão geométrica. Assim, independente do cenário, o custo operativo será limitado, e portanto ao calcular a esperança o custo médio também será limitado da mesma forma.

Uma vez que os custos de operação são sempre finitos, podemos definir uma função de custo futuro para cada estágio t , analogamente ao que fazemos no caso de horizonte finito. A função de

custo futuro do estágio 1, levando em conta o estado final v_2 do primeiro estágio será:

$$\begin{aligned} \overline{\mathcal{Q}}_2(v_2) = & \min_{(v_{t+1}, q_t, s_t, g_t, df_t, f_t)} \mathbb{E}_\xi \left[\sum_{t=2}^{\infty} \beta^{t-2} \left(c^\top g_t(\xi) + c_{df}^\top df_t(\xi) \right) \right] \\ \text{s.a. } & v_{t+1} = v_t + a_t - q_t - s_t, \\ & q_t + M_I g_t + df_t + M_D f_t = d_t, \\ & 0 \leq v_{t+1} \leq \bar{v}, \quad 0 \leq q_t \leq \bar{q}, \quad \underline{g} \leq g_t \leq \bar{g}, \\ & 0 \leq f_t \leq \bar{f}, \quad 0 \leq s_t, \quad 0 \leq df_t, \end{aligned} \quad (2.6)$$

onde as equações serão válidas apenas para $t = 2, 3, \dots$, sendo excluído o primeiro estágio correspondente a $t = 1$. Assim, a soma da função objetivo começa no estágio $t = 2$, e o fator de desconto também será ajustado, para corresponder ao preço corrente em $t = 2$.

Analogamente, definiremos as funções de custo futuro $\overline{\mathcal{Q}}_n(v_n)$ do n -ésimo estágio no modelo de horizonte infinito, como o mínimo do custo esperado de

$$\sum_{t=n}^{\infty} \beta^{t-n} (c^\top g_t(\xi) + c_{df}^\top df_t(\xi)),$$

sujeitos às mesmas restrições a partir de $t = n$, e dado o volume inicial v_n .

Feito isto, podemos recuperar a formulação de programação dinâmica da seção anterior, e obter a decomposição temporal do problema de otimização, separando em um custo presente e custo futuro a cada estágio, sujeito às restrições e incertezas *daquele estágio*.

Comentário 2.1. Seria possível adotar uma outra convenção para as funções de custo futuro, onde o valor das mesmas sempre se referisse ao tempo inicial $t = 1$. Neste caso, a definição das mesmas não necessitaria alterar o fator de desconto, mas, com isto, teríamos outra interpretação para as mesmas. A convenção que usamos será mais conveniente para o tratamento dos casos de horizonte periódico.

A grande dificuldade deste modelo de otimização estocástica em horizonte infinito é a presença de um número infinito de funções de custo futuro, uma para cada estágio. Assim, mesmo dispondo de uma formulação em termos de programação dinâmica, ainda seria necessário determinar uma infinidade destas funções de custo futuro, o que levaria a um custo computacional muito elevado. Por isto, destacamos nas próximas seções dois casos particulares, mas computacionalmente mais abordáveis, onde sempre teremos um número *finito* de funções de custo futuro a serem calculadas para obtermos a solução.

2.3 Caso periódico

No modelo apresentado na equação (2.1), cada estágio t é, a princípio, independente dos demais: a demanda d_1 pode ser diferente de d_2 , assim como as capacidades de geração $\bar{q}$, armazenamento $\bar{v}$, ... também podem evoluir com o tempo. Mais ainda, as afluências, dadas por variáveis aleatórias A_t , também podem ser diferentes. Aqui, queremos dizer que a variável aleatória A_1 pode ter *distribuição* diferente da distribuição de A_2 , e não que uma realização particular (em um cenário, por exemplo) destas variáveis foi diferente.

Entretanto, no caso do problema de planejamento da operação, alguns destes parâmetros aparecem de forma *periódica*. Os mais notáveis são: as afluências, que seguem um processo periódico, refletindo a sazonalidade das estações chuvosas e secas, e a demanda, que reflete aumentos e reduções de consumo de energia segundo as estações. Além disto, vários parâmetros podem ser

constantes: por exemplo, as capacidades das usinas e seus custos mudam apenas quando muda a configuração do sistema.

Comentário 2.2. Aqui, estamos analisando a periodicidade dos *parâmetros* do problema, ou seja, dos dados que determinam o problema de otimização específico que buscamos resolver. É a periodicidade destes parâmetros que determina se um problema será periódico ou não.

É importante frisar que, neste momento, não estamos considerando o valor que as *soluções* terão, nem as *variáveis de decisão*, que dependem não apenas do estágio, mas também de todo o cenário até elas. De fato, vale lembrar que as decisões $(v_{t+1}, q_t, s_t, g_t, df_t, f_t)$ dependem, além da realização particular da incerteza no estágio t , também do estado inicial v_t que, por sua vez, depende das decisões e incertezas passadas.

Estas características sugerem investigar a formulação do problema no caso em que *todos* os parâmetros do problema sejam periódicos. O problema específico de planejamento da operação, tal qual utilizado atualmente pelo ONS, não é estritamente periódico, mas pré-periódico. Veremos como este caso se comporta na próxima seção.

Neste caso periódico, podemos simplificar bastante a análise do modelo de otimização estocástica resultante. De fato, suponha que todos os *parâmetros* do problema (2.4) tenham período p .

Em primeiro lugar, a verificação de que o problema de planejamento está bem posto é mais simples: basta analisar cada um dos estágios que compõem um período quanto à

- presença de recurso relativamente completo (usualmente dado por variáveis de folga); e
- existência de uma cota superior C , independente do estado inicial e realização da incerteza ξ , para os custos.

Como visto na equação (2.5), para qualquer valor do fator de desconto β , menor do que 1, o problema será limitado.

Recapitulando a definição das funções de custo futuro, podemos notar que as funções $\bar{Q}_{2+p}$ e $\bar{Q}_2$ são idênticas. De fato, o problema de otimização correspondente a cada uma delas é o mesmo, já que os *parâmetros* que os definem são iguais. Aqui, em particular, vale lembrar a observação sobre o fator de desconto: a soma em $\bar{Q}_2$ começa em $t = 2$, mas é descontada com β^{t-2} para ser o custo em $t = 2$. Por sua vez, a soma de $\bar{Q}_{2+p}$ começa em $t = 2 + p$, mas o fator de desconto é alterado para $\beta^{t-(2+p)}$, já que $\bar{Q}_{2+p}$ representa o custo em valores no tempo $t = 2 + p$. Assim,

$$\begin{aligned}\bar{Q}_{2+p}(v) &= \min \sum_{t=2+p}^{\infty} \beta^{t-(2+p)} (c^\top g_t + c_{df}^\top df_t) \\ &= \min \sum_{\tau=2}^{\infty} \beta^{(\tau+p)-(2+p)} (c^\top g_{\tau+p} + c_{df}^\top df_{\tau+p}) \\ &= \min \sum_{\tau=2}^{\infty} \beta^{\tau-2} (c^\top g_{\tau+p} + c_{df}^\top df_{\tau+p}) = \bar{Q}_2(v).\end{aligned}$$

Isso garante que, dado um valor qualquer v , qualquer solução de $\bar{Q}_2(v)$ pode ser usada, transladada no tempo de p estágios, para resolver $\bar{Q}_{2+p}(v)$, e reciprocamente.

Naturalmente, o mesmo vale para qualquer estágio subsequente:

$$\bar{Q}_{t+p}(v) = \bar{Q}_t(v), \text{ para } t \geq 2.$$

Em particular, temos $p = 12$ no caso do planejamento da operação de médio prazo, o que mostra que, para resolver o caso periódico correspondente, basta calcular as funções de custo futuro correspondentes a

$$\overline{Q}_2, \overline{Q}_3, \dots, \overline{Q}_{12}, \text{ e } \overline{Q}_{13}.$$

Notemos que, como o problema não tem $\overline{Q}_1$, pois o primeiro estágio é determinístico, a formulação do problema de horizonte infinito e periódico exige a construção de 13 estágios: o primeiro, determinístico, e 12 estocásticos em seguida, para fornecer as 12 funções de custo futuro que descrevem um período completo. Este modelo está ilustrado na figura 2.2, onde indicamos com as setas retas a passagem do tempo. A última seta, ligando o estágio 13 ao segundo, deve ser interpretada não como uma volta ao passado, mas apenas representando o fato que o futuro do estágio 13 é análogo ao que aconteceria se começássemos do estágio 2, *com o estado de saída do estágio 13*.

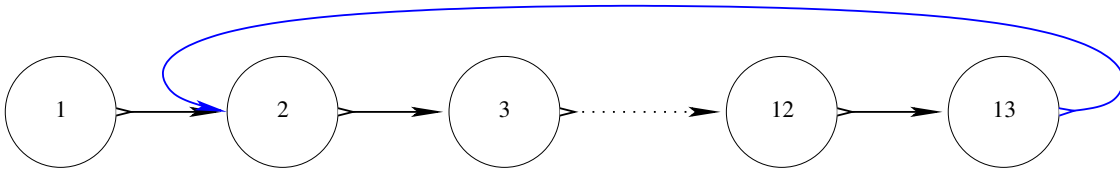


Figura 2.2: Problema de otimização de período 12.

As equações de programação dinâmica são, essencialmente, as mesmas do caso de horizonte finito e também do horizonte infinito: elas representam o custo total como soma de um custo presente e o custo futuro. No caso periódico, entretanto, podemos simplificar a recorrência, sem a necessidade de termos infinitas funções como no caso de horizonte infinito. De fato, como $Q_{2+p} = Q_2$, a equação que define $\overline{Q}_{1+p}$ é

$$\begin{aligned} Q_{1+p}(v_{1+p}, a_{1+p}) &= \min_{x_{1+p}} c^\top g_{1+p} + c_{df}^\top df_{1+p} + \beta \overline{Q}_2(v_{2+p}) \\ \text{s.a. } v_{2+p} &= v_{1+p} + a_{1+p} - q_{1+p} - s_{1+p}, \\ q_{1+p} + M_I g_{1+p} + df_{1+p} + M_D f_{1+p} &= d_{1+p}, \\ x_{1+p} &= (v_{1+p+1}, q_{1+p}, s_{1+p}, g_{1+p}, df_{1+p}, f_{1+p}) \in \mathcal{X}_{1+p}. \end{aligned}$$

Assim, a recorrência de programação dinâmica, que em horizonte finito começa em $t = T$, e é calculada até $t = 1$, passa a ser uma recorrência entre p funções de custo futuro, onde $\overline{Q}_3$ é usada para calcular $\overline{Q}_2$, $\overline{Q}_4$ para calcular $\overline{Q}_3$, e assim sucessivamente, até que a primeira função $\overline{Q}_2$ é, finalmente, utilizada para calcular $\overline{Q}_{p+1}$. No caso do problema de planejamento da operação, onde os últimos 5 anos são periódicos, isso se resulta numa redução de 60 funções de custo futuro para apenas 12.

Retornando à figura 2.2, também podemos interpretar as setas como indicando a relação entre as funções de custo futuro: Q_1 depende de Q_2 para ser calculada, Q_2 de Q_3 e assim por diante, até Q_{13} de volta a Q_2 . Neste caso, a similaridade é ainda maior: calculamos Q_{13} a partir de Q_2 exatamente da mesma forma que calculamos Q_2 a partir de Q_3 .

Esta redução do número de funções de custo futuro é, certamente, uma vantagem do modelo de horizonte periódico. Entretanto, vale ressaltar que, do ponto de vista matemático, o problema de horizonte finito é um problema *direto*: a partir da função de custo futuro terminal, $\overline{Q}_{T+1} \equiv 0$, calculamos sucessivamente as funções $\overline{Q}_T, \overline{Q}_{T-1}, \dots, \overline{Q}_3$ e $\overline{Q}_2$. Agora, devemos resolver um sistema de equações, onde a função de custo futuro $\overline{Q}_2$ entra no cálculo de $\overline{Q}_{p+1}$. Como veremos na seção 3, esta diferença leva a uma modificação do algoritmo de programação dual dinâmica

estocástica (PDDE). E, no âmbito *computacional*, esperamos observar, ao longo dos experimentos realizados nas próximas metas, que a redução de complexidade resultante de um menor número de funções de custo futuro compense o acréscimo resultante de $\overline{Q}_2$ não ser mais a “última” função a ser calculada.

2.4 Caso pré-periódico

Atualmente, o ONS utiliza um horizonte de planejamento de 10 anos. Nos 5 primeiros anos, a configuração é dinâmica, o que quer dizer que, a princípio, todos os parâmetros podem mudar. Na prática, isso corresponde a instalação de novas usinas e redes de transmissão, assim como o crescimento (projetado) da demanda de energia. Daí em diante, a configuração é fixa, a demanda é periódica ao longo do ano, e igual à demanda do quinto ano, e as incertezas continuam sendo *stagewise independent*, e periódicas de período 12.

Assim, o problema de planejamento da operação possui 60 estágios iniciais, que correspondem à configuração dinâmica, e ao crescimento da demanda, enquanto que os 60 últimos estágios são usados como condição de fronteira. Mas estes 60 últimos também são periódicos, de período 12, e podemos aplicar as mesmas ideias desenvolvidas anteriormente. Assim, temos ao todo 59 funções de custo futuro, correspondentes a $\overline{Q}_2$ até $\overline{Q}_{60}$, e mais 12 funções do segmento periódico, dadas por $\overline{Q}_{61}$ até $\overline{Q}_{72}$. Graficamente, esse problema pode ser representado como na figura 2.3, onde indicamos com τ o último estágio antes de começar o período (igual a 60 no caso do ONS), além de denotar o comprimento do período por p como de hábito.

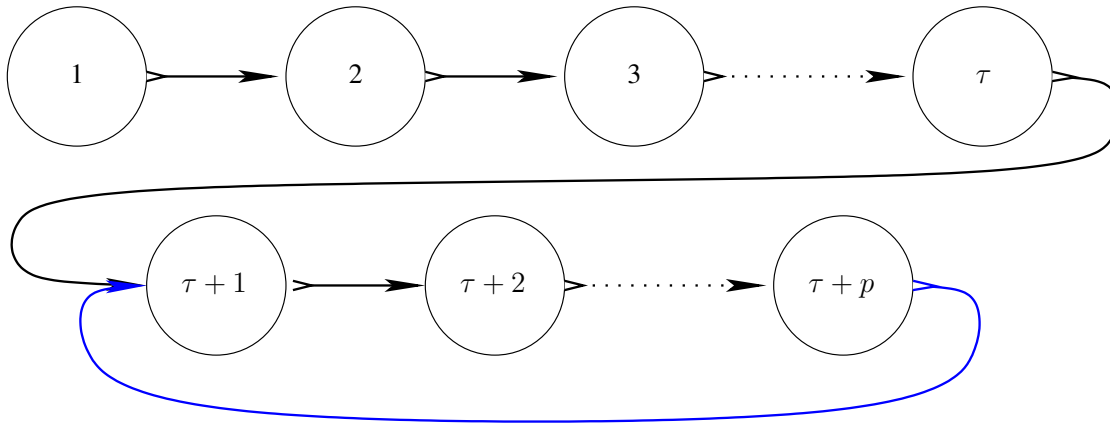


Figura 2.3: Problema de otimização pré-periódico: uma parte com τ estágios e uma parte de período p .

Este caso geral permite utilizar as técnicas de horizonte periódico em uma grande quantidade de problemas, pois podemos incorporar uma parte inicial com mais detalhes, ao mesmo tempo que consideramos a operação futura através do regime periódico final.

2.5 O fator de desconto

Vimos ao longo desta seção que o fator de desconto β possui duas interpretações: a primeira, econômica, ao transportar para o tempo t os custos do tempo $t + 1$. A segunda, mais matemática, ao garantir que problemas de otimização com infinitos estágios, e portanto os problemas periódicos, são coerentes, e em particular que as funções de custo futuro destes problemas estão bem definidas.

Veremos, na próxima seção, que o fator de desconto também será importante para a construção do algoritmo, garantindo que este converge para uma solução do problema de horizonte periódico.

Neste momento, vale observar que o fator de desconto aparece, nas estimativas de custo total do problema, através da desigualdade

$$\overline{Q} \leq \frac{C}{1 - \beta},$$

correspondente à soma de uma série geométrica. Assim, quanto mais próximo de 1 estiver β , maior será o termo da direita. Ora, β próximo de 1 corresponde a uma taxa de juros próxima de zero, já que $\beta = \frac{1}{1+r}$. Substituindo na equação acima, obtemos

$$\overline{Q} \leq C \frac{1+r}{r}.$$

A maior parte destes cálculos pode ser invertido para considerar o custo mínimo: se, a cada estágio, o custo é de, pelo menos c , também podemos garantir que

$$\overline{Q} \geq \frac{c}{1 - \beta} = c \frac{1+r}{r},$$

o que mostra que o fator de desconto β (ou a taxa de juros r) determinam a ordem de grandeza dos custos.

2.6 Generalizações do modelo

De forma bastante geral, a metodologia de horizonte periódico pode ser adaptada para qualquer problema que seja dado por uma recorrência de programação dinâmica da forma

$$Q_t(x_t, \xi_t) = \min_{(x_{t+1}, y_t)} c^\top y_t + \beta \overline{Q}_{t+1}(x_{t+1}) \quad (2.7)$$

$$\text{s.a. } A_t x_{t+1} + B_t x_t + W_t y_t = d_t,$$

$$\overline{Q}_{t+1}(x_{t+1}) = \rho_t [Q_{t+1}(x_{t+1}, \xi_{t+1})] \quad (2.8)$$

onde x_t é o estado, y_t são as variáveis de decisão local, ξ_t representa as incertezas (podendo afetar A , B , W , c e d), e ρ_t é uma medida coerente de risco, por exemplo a média ou o CV@R [STdCS13].

Da mesma forma que anteriormente, se todos os parâmetros do problema forem p periódicos, (ou se forem periódicos a partir do estágio $\tau + 1$) poderemos também simplificar o problema, ao identificar as funções de custo futuro da parte periódica do problema. Isso inclui, naturalmente, as matrizes e vetores A , B , W , c e d , que devem depender periodicamente de t , assim como as incertezas, que devem ser processos de período p , e também as medidas de risco ρ_t . Vale notar que, em geral, as medidas de risco são sempre as mesmas em todos os estágios, já que não faz sentido, em geral, alterar os parâmetros de aversão a risco. Portanto, a generalização para o caso avesso a risco pode ser feita na maioria dos casos.

A mesma observação pode ser feita para o caso de problemas onde a incerteza segue um processo autorregressivo. Aqui, não é verdade que a variável aleatória natural para o problema é dada por um processo periódico, já que esta depende da realização no tempo anterior. Entretanto, é possível realizar uma mudança de variáveis, reescrevendo a incerteza através de uma equação linear. Por exemplo, para um processo de ordem 2, dito AR(2), temos

$$\xi_t = \Phi_{t,1} \xi_{t-1} + \Phi_{t,2} \xi_{t-2} + \eta_t.$$

Agora, as variáveis η_t são, além de *stagewise independent*, também parte de um processo periódico, pois sua distribuição depende do ajuste da série ξ_t , dada pelos coeficientes $\Phi_{t,1}$ e $\Phi_{t,2}$.

Como no caso da PDDE em horizonte finito, este tipo de dependência é incorporado através da inclusão de *novas variáveis de estado*, que representam o estado do processo ξ_t , e no qual passamos a considerar apenas a inovação η_t como parte da incerteza. Assim, o problema passa a ser:

$$V_t(x_t, \xi_{t-1}, \xi_{t-2}; \eta_t) = \min_{(x_{t+1}, y_t)} c^\top y_t + \beta \bar{V}_{t+1}(x_{t+1}, \xi_t, \xi_{t-1}) \quad (2.9)$$

$$\begin{aligned} \text{s.a} \quad & A_t x_{t+1} + B_t x_t + W_t y_t = d_t \\ & \xi_t = \Phi_{t,1} \xi_{t-1} + \Phi_{t,2} \xi_{t-2} + \eta_t, \end{aligned}$$

$$\bar{V}_{t+1}(x_{t+1}, \xi_t, \xi_{t-1}) = \rho_t [V_{t+1}(x_{t+1}, \xi_t, \xi_{t-1}; \eta_{t+1})] \quad (2.10)$$

Esta mudança de variáveis (e a consequente mudança das funções de custo futuro) é essencial para preservar a hipótese de que os ruídos são *stagewise independent*. E só nestas condições podemos garantir que as funções de custo futuro dependem apenas do estágio em questão, e do estágio atual. Neste sentido, o aumento das variáveis de estado também corresponde ao acréscimo de informação necessária para representar o estado do sistema: se as afliências seguem um processo AR2, então não apenas o estado do reservatório deve ser levado em conta para o planejamento futuro, mas as afliências dos dois estágios anteriores também.

3 Algoritmo de solução

O algoritmo de PDDE, utilizado para resolver problemas de otimização estocástica multiestágio, possui três componentes fundamentais:

1. A representação das funções de custo futuro;
2. A construção de aproximações destas funções usando dualidade para programação linear; e
3. A amostragem dos cenários, independentemente e a cada estágio, para a exploração da árvore de cenários.

O uso das funções de custo futuro é assegurado pela hipótese sobre as incertezas: a cada estágio, elas são dadas por uma variável aleatória ξ_t *independente* das anteriores. Esta hipótese é conhecida como *stagewise independence*, para significar que as variáveis ξ_t são independentes ao longo dos estágios.¹ Isto permite utilizar uma única função de custo futuro por estágio, levando desta forma às equações de programação dinâmica (2.1) e (2.2).

O mesmo se dá no caso de horizonte infinito: teremos uma função de custo futuro $\bar{Q}_t$ por estágio sempre que as incertezas do modelo forem *stagewise independent*.

Nesta seção, veremos como as componentes da PDDE devem ser adaptadas para resolver problemas *periódicos*. Começaremos introduzindo os *Operadores de Bellman*, que representam de forma exata uma etapa da recorrência de programação dinâmica, para em seguida analisar como tais operadores permitem construir uma solução do problema de otimização estocástica periódico. Enfim, veremos como implementar um método de aproximações sucessivas que permite construir as funções de custo futuro em problemas periódicos, e com isto calcular uma política ótima para os mesmos.

¹Esta distinção é relevante no caso em que as incertezas ξ_t são dadas por um *vetor aleatório*. Neste caso, as coordenadas de ξ_t *podem* ser dependentes entre si — elas não podem ser dependentes das coordenadas observadas em ξ_{t-1} .

3.1 Definições preliminares

A construção das funções de custo futuro, através da solução de problemas de otimização correspondentes a um passo a frente, pode ser resumida pelo conceito de *Operador de Bellman*, que representa a transmissão dos custos de um estágio para o anterior. Mais concretamente:

Definição 3.1 (Operador de Bellman). O Operador de Bellman correspondente ao estágio t do problema de otimização estocástica (2.1) e (2.2) é a transformação que leva uma função ϕ na função $\mathcal{B}_t(\phi)$, se usarmos ϕ como função de custo futuro do estágio t . Ou seja, dada $\phi(v_{t+1})$, $(\mathcal{B}_t(\phi))(x_t)$ é:

$$\mathbb{E}_{\xi_t} \left[\begin{array}{l} \min_{(v_{t+1}, q_t, s_t, g_t, df_t, f_t)} \quad c^\top g_t + c_{df}^\top df_t + \beta \phi(v_{t+1}) \\ \text{s.a.} \quad v_{t+1} = v_t + a_t - q_t - s_t, \\ \quad \quad q_t + M_I g_t + df_t + M_D f_t = d_t, \\ \quad \quad 0 \leq v_{t+1} \leq \bar{v}, \quad 0 \leq q_t \leq \bar{q}, \quad \underline{g} \leq g_t \leq \bar{g}, \\ \quad \quad 0 \leq f_t \leq \bar{f}, \quad 0 \leq s_t, \quad 0 \leq df_t \end{array} \right]$$

Comentário 3.2. A notação usual para os operadores de Bellman não inclui os parêntesis, que são comuns para a aplicação de uma função, para não sobrecarregar as equações. Assim, sempre que for conveniente usaremos $\mathcal{B}_t \phi$ em vez de escrever $\mathcal{B}_t(\phi)$.

Observe que a única diferença entre a definição acima e a formulação das equações (2.1)–(2.2) é que substituímos a *verdadeira* função de custo futuro $\bar{Q}_{t+1}$ por uma função qualquer ϕ . Esta mudança é fundamental para analisar a programação dinâmica, em especial no caso de horizonte infinito, pois permite falar sobre as aproximações das funções de custo futuro que serão construídas ao longo do algoritmo de programação dinâmica.

Podemos observar esta diferença através das figuras a seguir. Na figura 3.1, destacamos como a programação dinâmica relaciona as funções de custo futuro $\bar{Q}_{t+1}$ e $\bar{Q}_t$. O fluxo do tempo vai de t para $t+1$, e a função de custo futuro $\bar{Q}_{t+1}$, correspondente ao custo médio a partir do estágio $t+1$, é usada no estágio t para calcular $\bar{Q}_t$, que representará o custo médio a partir de t (e será usada no estágio $t-1$). Já na figura 3.2, mostramos como esta operação é generalizada para observar o que acontece apenas no estágio t . Dada uma função ϕ que representa o valor do estado final no tempo t , e que a princípio pode ser uma função qualquer, o problema de otimização do estágio t permite calcular uma função $\mathcal{B} \phi$, que corresponde ao valor do estado *inicial* do estágio t .

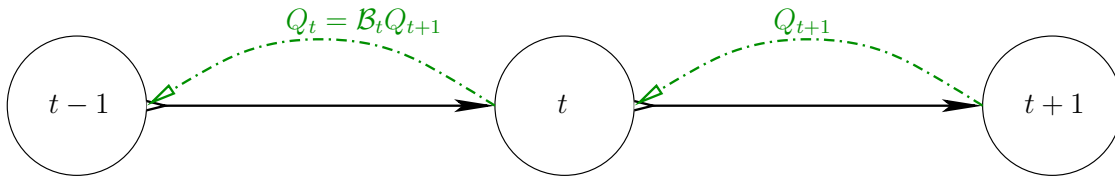


Figura 3.1: Programação dinâmica no estágio t : funções de custo futuro $\bar{Q}$

3.2 Recursão de programação dinâmica

Com as notações acima, a recorrência definida pelas equações de programação dinâmica é dada por:

$$\bar{Q}_t(v_t) = \begin{cases} 0 & , t = T + 1 \\ \mathcal{B}_t(\bar{Q}_{t+1}) & , t \in \{2, \dots, T\} \end{cases} \quad (3.1)$$

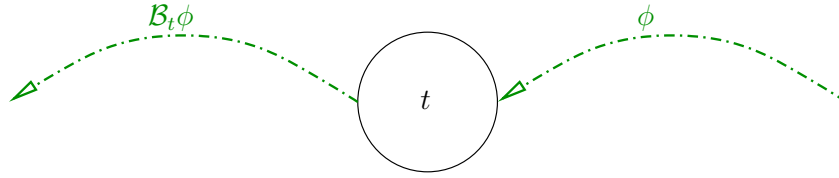


Figura 3.2: Operador de Bellman do estágio t : transformação do custo futuro ϕ em $\mathcal{B}\phi$.

Já no caso da programação dinâmica para horizonte periódico, as equações passam a ser

$$\bar{Q}_t(v_t) = \begin{cases} \mathcal{B}_t(\bar{Q}_2) & , t = p + 1 \\ \mathcal{B}_t(\bar{Q}_{t+1}) & , t \in \{2, \dots, p\} \end{cases} \quad (3.2)$$

onde p é o período da sazonalidade. Esta notação permite focalizar ainda mais no fato que as funções de custo futuro de um período são interdependentes, já que agora $\bar{Q}_2$ irá entrar no cálculo de $\bar{Q}_{p+1}$.

Assim, observamos uma diferença fundamental entre os problemas de horizonte finito e os problemas periódicos. Em teoria, seria possível resolver problemas de horizonte finito de forma *direta* a partir da recorrência, calculando $\bar{Q}_T$ a partir de $\bar{Q}_{T+1} = 0$, e daí $\bar{Q}_{T-1}$ a partir de $\bar{Q}_T$, e assim sucessivamente até obter $\bar{Q}_2$ e finalmente resolver o problema de primeiro estágio. Por outro lado, o problema periódico não possui um “último estágio” do qual começar a recorrência, e portanto o próprio método de solução é diferente.

Vamos primeiramente considerar o caso em que o período é 1, para simplificar a notação. O diagrama de estágios, análogo ao que vimos na seção 2, está ilustrado na figura 3.3. Nele, vemos o primeiro estágio, que é determinístico, e leva ao segundo estágio, que é estocástico. Este segundo estágio corresponde a uma função de custo futuro $\bar{Q}_2$, que é usada, como no caso de horizonte finito, para representar o custo futuro no problema do primeiro estágio. Além disso, vemos a periodicidade

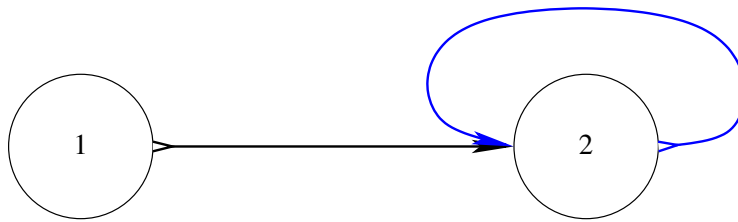


Figura 3.3: Problema 1 periódico: uma única função de custo futuro, em $t = 2$.

que aparece a partir do segundo estágio: daí em diante, o problema que será resolvido será sempre o mesmo, variando apenas a condição inicial, que corresponde ao volume no final de cada etapa do planejamento.

Portanto, neste caso buscamos uma única função de custo futuro $\bar{Q}$, que satisfaz à equação

$$\bar{Q}(v) = (\mathcal{B}\bar{Q})(v). \quad (3.3)$$

Este tipo de equação é chamado de *equação de ponto fixo*: buscamos um função que, após ser transformada pelo operador de Bellman $\mathcal{B}$, dá a mesma função. Continuando a analogia, vemos que este modelo corresponde à transição de funções de custo futuro ilustrada na figura 3.4. Daqui, podemos ver que apenas o problema de segundo estágio é necessário para calcularmos a função de custo futuro $\bar{Q}_2$.

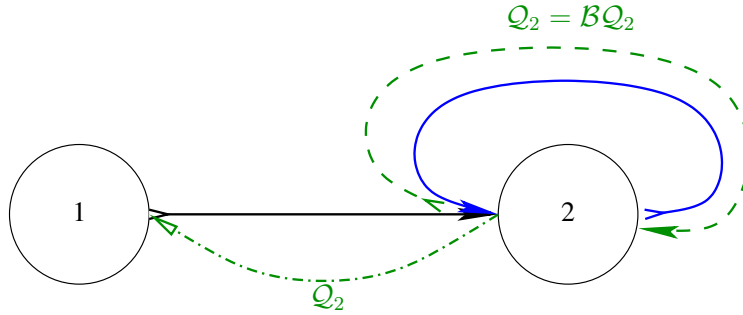


Figura 3.4: Problema 1 periódico: a função de custo futuro satisfaz $\bar{Q} = \mathcal{B}\bar{Q}$.

Em geral, este tipo de equação pode ser resolvido pela construção de uma sequência Q^n que satisfaz

$$Q^{n+1} = \mathcal{B}Q^n,$$

desde que a transformação $\mathcal{B}$ seja uma *contração*. Veremos que, quando $0 < \beta < 1$, de fato $\mathcal{B}$ é uma contração, e com isso o limite $Q^\infty = \lim_{n \rightarrow \infty} Q^n$ satisfaz a equação (3.3). Esta técnica é conhecida como *iteração de ponto fixo*, porque desejamos encontrar um valor fixado pelo operador $\mathcal{B}$.

Intuitivamente, a construção da sequência de funções de custo futuro Q^n pode ser pensada como a aproximação do custo de infinitos estágios a partir do custo de um estágio, depois de dois estágios, e assim por diante, como se incluíssemos sucessivamente mais estágios no problema de otimização. Mais ainda, como vimos anteriormente, se os custos imediatos são limitados por C , o custo total é sempre menor do que ou igual a $\frac{C}{1-\beta}$, e portanto sabemos que as funções Q^n não vão ficar arbitrariamente grandes.

Vejam agora porque a sequência de funções Q^n converge. Como anunciado anteriormente, a condição fundamental é que o operador $\mathcal{B}$ seja uma contração, ou seja, que exista uma constante $L < 1$ tal que, para quaisquer funções f e g ,

$$d(\mathcal{B}f, \mathcal{B}g) \leq Ld(f, g), \quad (3.4)$$

onde d indica a distância entre estas funções. No nosso caso, a constante L será exatamente β .

Vejam: seja x_f uma solução ótima para o problema que calcula $(\mathcal{B}f)(v)$, e respectivamente x_g para o problema de $(\mathcal{B}g)(v)$. Como a condição inicial é a mesma, a solução x_f é *viável* para o problema de otimização $(\mathcal{B}g)(v)$. Substituindo na função objetivo, isto dá

$$(\mathcal{B}g)(v) \leq c^\top g_f + c_{df}^\top df_f + \beta g(v_f).$$

Por outro lado,

$$(\mathcal{B}f)(v) = c^\top g_f + c_{df}^\top df_f + \beta f(v_f),$$

e portanto

$$(\mathcal{B}g)(v) - (\mathcal{B}f)(v) \leq \beta[g(v_f) - f(v_f)].$$

Analogamente, substituindo x_g como solução viável para o problema de $(\mathcal{B}f)(v)$, obtemos

$$(\mathcal{B}f)(v) - (\mathcal{B}g)(v) \leq \beta[f(v_g) - g(v_g)].$$

Com estas duas desigualdades, vemos que, se a distância entre f e g sempre for menor do que D , ao calcularmos $\mathcal{B}f$ e $\mathcal{B}g$ a distância será sempre menor do que βD .

Daí, podemos analisar a sequência Q^n diretamente. Seja D a distância entre Q^0 e $Q^1 = \mathcal{B}Q^0$, as duas primeiras funções de custo futuro construídas. A distância entre as funções Q^1 e Q^2 será, no máximo:

$$d(Q^1, Q^2) = d(\mathcal{B}Q^0, \mathcal{B}Q^1) \leq Ld(Q^0, Q^1) = LD.$$

Da mesma forma, teremos

$$d(Q^n, Q^{n+1}) \leq L^n D,$$

e como $L < 1$, isso mostra que a distância entre funções sucessivas converge. Mais ainda, pela desigualdade triangular, a distância entre Q^n e Q^{n+m} será no máximo a soma das distâncias sucessivas, ou seja, no máximo

$$d(Q^n, Q^{n+m}) \leq d(Q^n, Q^{n+1}) + \dots + d(Q^{n+m-1}, Q^{n+m}) \leq L^n(D + LD + \dots + L^{m-1}D) \leq L^n \frac{D}{1-L}$$

que também tende a zero, e mostra que a sequência Q^n de fato converge para um limite Q^∞ . Como $Q^{n+1} = \mathcal{B}Q^n$, para todo n , tomando o limite temos

$$Q^\infty = \mathcal{B}Q^\infty,$$

o que mostra que, de fato, Q^∞ é solução da equação (3.3).

3.2.1 Caso geral: período p

Vejamos agora como generalizar as ideias apresentadas para o caso em que o período é maior do que um. Agora, teremos p funções de custo futuro, $\overline{Q}_2$ até $\overline{Q}_{1+p}$, ligadas entre si por diferentes operadores de Bellman $\mathcal{B}_t$, correspondentes aos problemas dos estágios 2 até $1+p$.

Da mesma forma que anteriormente, cada um dos operadores de Bellman $\mathcal{B}_t$ é uma contração com constante β , já que este resultado depende apenas da estrutura do problema de otimização. Assim, podemos adaptar a construção da sequência Q^n para a construção de uma sequência de p funções de custo futuro

$$(Q_2^n, Q_3^n, \dots, Q_{1+p}^n).$$

A recorrência que define esta sequência é também a aplicação dos operadores de Bellman, correspondentes a cada caso. Lembrando que Q_2 é construída a partir de Q_3 , e Q_{1+p} depende de Q_2 , teremos:

$$(Q_2^{n+1}, Q_3^{n+1}, \dots, Q_{1+p}^{n+1}) = (\mathcal{B}_2 Q_3^n, \mathcal{B}_3 Q_4^n, \dots, \mathcal{B}_{1+p} Q_2^n). \quad (3.5)$$

Aqui, observamos que a dependência cíclica das funções de custo futuro é realizada pela relação da última coordenada com a primeira: $Q_{1+p}^{n+1} = \mathcal{B}_{1+p} Q_2^n$.

Daí, podemos continuar exatamente como no caso de período um. Coletivamente, as funções $(Q_2^n, Q_3^n, \dots, Q_{1+p}^n)$ fazem parte de uma contração conforme n aumenta, e convergem para o ponto fixo

$$(Q_2^\infty, Q_3^\infty, \dots, Q_{1+p}^\infty).$$

Este limite dá as p verdadeiras funções de custo futuro $\overline{Q}_2, \overline{Q}_3, \dots, \overline{Q}_{1+p}$ do problema periódico.

Finalmente, no caso em que o problema é pré-periódico, com τ estágios antes de começar a parte periódica de período p , este processo pode ser adaptado para construir as funções de custo futuro $\overline{Q}_{\tau+1}$ até $\overline{Q}_{\tau+p}$, que só dependem umas das outras. Em seguida, as funções de custo futuro $\overline{Q}_t$, para $2 \leq t \leq \tau$, podem ser calculadas normalmente, já que para estas basta dispor da função de custo $\overline{Q}_{\tau+1}$, que substitui a condição terminal $\overline{Q}_{T+1} = 0$ no problema de horizonte τ .

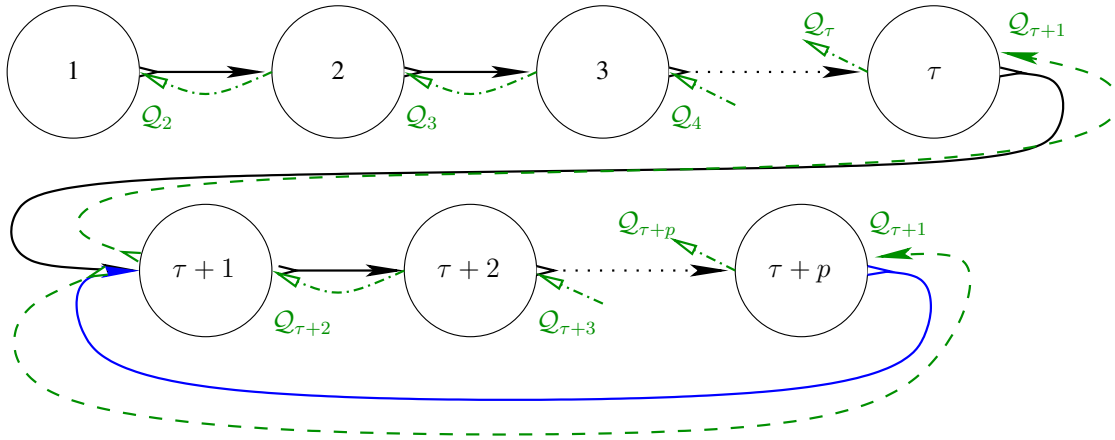


Figura 3.5: Funções de custo futuro e recorrência no caso pré-periódico.

3.3 Modificações na PDDE

Como vimos, a relação entre as funções de custo futuro em problemas periódicos é muito similar à que temos em problemas de horizonte finito: a função de custo futuro no tempo t será calculada a partir da função de custo futuro no tempo $t + 1$. No caso de horizonte finito, esta construção poderia, a princípio, ser feita de forma completa através de programação dinâmica, começando em T e voltando até o primeiro estágio. Já no caso de horizonte periódico, como utilizaremos apenas as p primeiras funções, o cálculo de $\bar{Q}_{1+p}$, que depende de $\bar{Q}_{2+p}$, na verdade irá se referir a $\bar{Q}_2$, o que forma um ciclo de dependências.

Na seção anterior, vimos que este ciclo não é um problema, ao contrário: é justamente ao percorrer este ciclo, a princípio infinitas vezes, que obteremos as funções de custo futuro $\bar{Q}_t$ para t entre 2 e $1 + p$. Mais ainda, na prática, é inviável a construção exata das funções $\bar{Q}_t$, pois o cálculo delas a partir de $\bar{Q}_{t+1}$ exigiria efetuar o equivalente a um passo *backward* para todos os estados possíveis — ou, pelo menos, um passo para cada multiplicador dual, o que ainda é um número muito grande.

No caso de horizonte finito, o algoritmo de PDDE usa um método, também iterativo, para aproximar as funções de custo futuro, calculando *cortes* para $\bar{Q}_t$ a partir de uma aproximação Q_{t+1} da real função $\bar{Q}_{t+1}$. Em geral, estas aproximações são indexadas pela iteração do algoritmo, e Q_{t+1}^k denota a aproximação de $\bar{Q}_{t+1}$ vigente na k -ésima iteração da PDDE. Esta etapa de cálculo dos cortes é conhecida como *etapa backward*, e a etapa que determina quais os pontos serão utilizados para calcular os cortes, é denominada *etapa forward*.

Vejamos, a seguir, como estas etapas são modificadas no caso de horizonte periódico.

3.3.1 Convergência e etapa *Forward*

Para motivar as modificações aportadas à etapa *forward*, vamos analisar um critério de convergência para o algoritmo de PDDE. Em teoria, uma forma de garantir a convergência da PDDE consiste em comparar a cota inferior, que é obtida no problema de primeiro estágio, com uma estimativa (em geral, via Monte Carlo) do custo da política que seria implementada dadas as funções de custo futuro naquele momento do algoritmo. Esta estimativa normalmente é obtida através da simulação de um grande número de cenários, e fornece intervalos de confiança para a cota superior z_{sup} .

No caso do horizonte infinito, também é possível obter uma estimativa para o custo de operação, porém devemos lidar com a limitação de simular apenas um número finito de estágios. Se

calculássemos o custo *exato* da política T estágios à frente, cometeríamos um erro no máximo de

$$\sum_{t=T+1}^{\infty} \beta^{t-1} C = \beta^T \frac{C}{1-\beta}$$

utilizando a hipótese de que qualquer estágio terá custo imediato no máximo C , que já utilizamos anteriormente.

Ora, na prática iremos apenas estimar o custo da política por Monte Carlo. Assim, ao calcularmos o custo de N cenários, ao longo de T estágios cada, obtemos que o custo da política será, com 95% de probabilidade, no máximo

$$\mu + 1.96\sigma + \beta^T \frac{C}{1-\beta},$$

onde μ é a média e σ o desvio padrão amostral do custo dos N cenários. Aqui, vemos a presença de *dois* termos de erro: o primeiro, onde aparece o desvio padrão, corresponde ao erro amostral, decorrente da estimativa de Monte Carlo. O segundo, contendo β^T , corresponde a um erro sistemático, devido ao truncamento do horizonte no tempo T . Assim, vemos que, para reduzir o impacto deste último termo, devemos simular um número de estágios suficientemente grande, para que o fator β^T seja pequeno o suficiente.

Ora, se precisamos simular T estágios para estabelecer que as políticas estão suficientemente boas, é natural escolher pontos na etapa *forward* que correspondam, também, a avançar T estágios à frente. Assim, no caso de horizonte periódico, a etapa *forward* deve ser modificada para gerar não apenas uma sequência de estados para cada função de custo futuro, mas para cada estágio do problema até T , que é maior do que o período p .

Assim, cada cenário da etapa *forward* conterá T realizações da incerteza, e determinará T estados finais $v_2, v_3, \dots, v_{T+1}$. Aqui, vale ressaltar como estes estados serão calculados. Naturalmente, a etapa *forward* procederá naturalmente do primeiro estágio até o último estágio antes do ciclo. Ao final deste último estágio $\tau + p$, o estado final $v_{\tau+p+1}$ será utilizado como estado inicial para o problema do estágio $\tau + 1$ (que também é o problema do estágio $\tau + p + 1$), e daí calculamos novamente os estados finais, seguindo os estágios sucessivos. A cada vez que chegamos no problema do estágio $\tau + p$, o estado final $\tau + mp + 1$ (ao final do m -ésimo ciclo) será utilizado novamente no estágio $\tau + 1$, até que esgotemos o comprimento T dos cenários. Esta construção está ilustrada na figura 3.6, onde indicamos os estados finais de cada estágio no início da flecha que vai para o estágio seguinte.

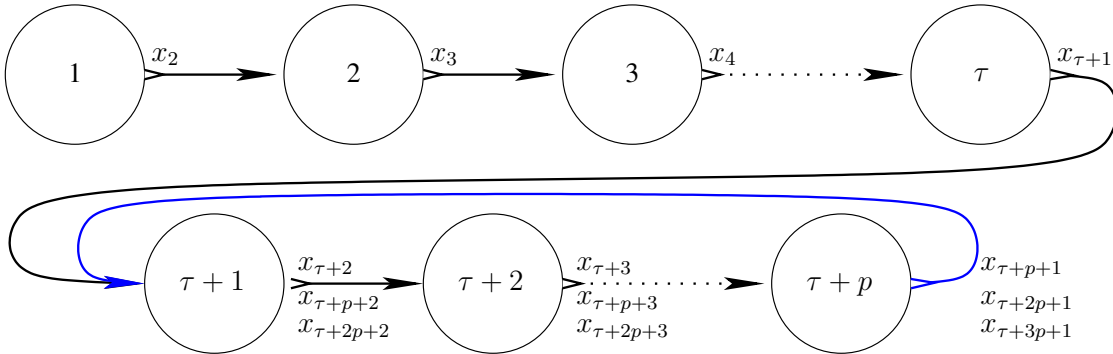


Figura 3.6: Construção de estados durante a etapa *forward*.

Veremos na próxima seção como estes estados serão utilizados na etapa *backward*.

Para simplificar, em geral escolhemos T como τ mais um múltiplo de p , para realizar um número inteiro de ciclos, mas esta escolha não é absolutamente necessária.

Mais ainda, na prática, é possível escolher cenários *forward* de comprimento variável. Isto permite evitar que os cenários tenham um comprimento muito grande, realizando passos *forward* curtos para tentar aprimorar rapidamente as funções de custo futuro, sem deixar de realizar passos mais longos para garantir a convergência. Com isto, não é necessário fixar, a priori, o comprimento máximo dos cenários *forward*, o que dá mais flexibilidade para o algoritmo.

3.3.2 Construção de cortes e etapa *Backward*

Em modelos periódicos, a etapa *backward* da PDDE deve ser alterada por duas razões. Primeiro, para representar a dependência cíclica das funções de custo futuro, dada a natureza do problema. Segundo, para lidar com os estados dados pelos cenários *forward*, que irão produzir mais de um estado por estágio.

Dado um estado v_{t+1} , final do estágio t , resolvemos o problema do estágio $t + 1$, para todos os cenários possíveis, e calculamos o corte resultante da mesma maneira que no caso da PDDE em horizonte finito. A única diferença é que, ao calcular um corte para $\mathcal{Q}_{\tau+p}$, o problema do estágio $\tau + p + 1$ corresponde, pela periodicidade, ao problema do estágio $\tau + 1$. Portanto, iremos resolver o problema do estágio $\tau + 1$, mas com a condição inicial sendo dada por $v_{\tau+p+1}$, e não por $v_{\tau+1}$ (que é o estado final vindo do estágio τ).

Com isto, garantimos que a periodicidade do problema será capturada pelas funções de custo futuro, pois o cálculo de todas elas será ciclicamente dependente das demais, como ilustrado na figura 3.5.

Enfim, resta a questão da ordem, dentro do algoritmo, com a qual os estados calculados pela etapa *forward* serão utilizados para a construção de cortes. Em teoria, estes poderiam ser usados em qualquer ordem, mas na prática é mais razoável realizar a atualização das funções de custo futuro da mesma forma que na PDDE em horizonte finito: do último estágio para o primeiro. Aqui, isso significa “desenrolar” o ciclo do último estado final x_{T+1} até o primeiro estado final x_2 , atualizando as funções de custo futuro a cada estado considerado. Isto permite utilizar, ao longo da própria etapa *backward*, os cortes já calculados e correspondentes a estágios seguintes.

Mais ainda do que no caso da PDDE tradicional, problemas de horizonte periódico necessitam de algoritmos para *seleção de cortes*. Isto pode ser compreendido ao analisar as primeiras iterações: na primeira vez, se começamos de um custo igual a zero, o custo que a parte periódica poderá considerar corresponde a apenas p estágios. Já na segunda iteração, após ter calculado cortes para todos estágios, o custo final já não será mais nulo, e a parte periódica pode acessar informação de 2 períodos a frente². Com isto, praticamente todos os cortes realizados na primeira iteração estarão obsoletos, e isto continuará até que tenhamos realizado suficientes iterações.

Para se ter uma ideia desta magnitude, a estimativa de custo para um fator de desconto correspondente a uma taxa de juros de 12% em um período corresponde a 9.333 vezes o custo do primeiro período. Além disso, devido ao fator de desconto, o custo de 5 períodos completos é de 4.038 vezes o custo de um período. Assim, se a primeira iteração realizasse 5 ciclos na parte periódica, ainda assim teríamos considerado menos da metade dos custos. Portanto, ao realizar mais iterações, e convergir para as reais funções de custo futuro, estes cortes provavelmente estarão bem abaixo dos cortes mais recentes.

²na verdade, de pelo menos 2 períodos; possivelmente mais se a etapa *forward* tiver realizado mais de um ciclo

3.4 Compartilhamento de cortes

Depois desta exposição, é possível apresentar uma nova interpretação do algoritmo de PDDE para problemas de horizonte periódico. De fato, ao dizer que as funções de custo futuro dos estágios t e $t + p$ são iguais, e ao utilizar os estados finais de ambos para gerar cortes para a mesma função, estamos efetivamente *compartilhando* os cortes, não apenas entre cenários diferentes do mesmo instante t , mas também entre instantes de tempo diferentes, e separados pelo período p .

Por um lado, este uso de cortes do tempo $t + p$ no tempo t permite incorporar uma maior quantidade de informação em uma mesma função de custo futuro, e com isto esperamos uma convergência mais rápida do algoritmo de PDDE. Por outro lado, o uso de um corte do tempo t no tempo $t + p$ irá se traduzir em um aumento do custo no tempo $t + p$, o que é esperado devido à periodicidade: os custos futuros, no tempo t ou no tempo $t + p$, devem ser os mesmos, se considerados em valores respectivamente dos tempos t e $t + p$.

Assim, uma implementação ainda mais simples do algoritmo de horizonte periódico pode considerar um problema com horizonte finito, e realizar o compartilhamento de cortes entre estágios correspondentes. Entretanto, esta implementação é limitada por não explorar corretamente todo o espaço de estados, o que seria necessário, pelo menos em teoria, para estabelecer a convergência do método.

4 Regime estacionário

Nas seções anteriores, apresentamos os modelos de otimização em horizonte periódico e pré-periódico, descrevendo os casos em que estes podem ser utilizados, e as modificações a serem aportadas no algoritmo de PDDE para resolver tais problemas quando os ruídos são *stagewise independent*.

Nesta seção, iremos analisar não o algoritmo de solução, mas as soluções propriamente ditas de problemas de horizonte periódico. Assim, nesta seção, supomos que o problema de horizonte periódico está resolvido, o que significa que as funções de custo futuro $\bar{Q}_t$ já foram determinadas, e observaremos como as variáveis de decisão e de estado evoluem ao longo do tempo. Com efeito, veremos que em muitos casos estas variáveis convergem para um “regime estacionário”, típico de problemas periódicos.

Os resultados aqui apresentados podem ser transferidos para o caso de problemas pré-periódicos, ao considerar apenas os p estágios que compõem um período do problema.

4.1 Definições

Primeiramente, vamos fixar a notação para os cenários e as correspondentes decisões tomadas ao longo do tempo, para evitar ambiguidades.

Definição 4.1 (Processo estocástico). Um processo estocástico a tempo discreto é uma sequência de variáveis aleatórias Ξ_t , para $t = 1, 2, \dots$. Em geral, estes processos são denotados por $(\Xi_t)_{t=1}^\infty$, ou, quando for claro pelo contexto, apenas (Ξ_t) .

Neste relatório, consideraremos apenas processos a tempo discreto, e usaremos apenas a terminologia “processo estocástico”.

Definição 4.2 (Cenários). Seja $(\Xi_t)_{t=1}^\infty$ um processo estocástico. Um *cenário de comprimento T* é dado por uma realização das T primeiras variáveis aleatórias $\Xi_1, \dots, \Xi_T$. Habitualmente, os cenários são denotados com uma variável minúscula, e o seu comprimento é indicado pelos colchetes, para diferenciar da realização em um tempo t particular:

$$\xi_{[T]} = (\xi_1, \xi_2, \dots, \xi_T).$$

Em otimização estocástica, os cenários são comumente representados em uma árvore, indicando a passagem do tempo e a relação entre os cenários.

Definição 4.3 (Árvore de cenários). Se, para cada tempo t , as realizações do processo estocástico (Ξ_t) são finitas, então é possível construir uma *árvore de cenários*, onde os cenários $\xi_{[t]}$ são os nós, e as arestas ligam um cenário $\xi_{[t]}$ de comprimento t a cenários de comprimento $t + 1$ cujas t primeiras coordenadas coincidem com as realizações de $\xi_{[t]}$.

A versão extensiva de problemas de otimização estocástica associa um vetor de decisões x_t para cada cenário $\xi_{[t]}$, respeitando as restrições de não-antecipatividade [SDR09]. Assim, se consideramos que os cenários $\xi_{[t]}$ são realizações de um processo estocástico, então também podemos considerar as decisões $x_t(\xi_{[t]})$ como realizações do processo estocástico que indica qual decisão será tomada a cada instante. Em particular, o estado final de cada estágio, que no problema de planejamento da operação é v_{t+1} , também será função do cenário atual, com todo o seu histórico, e também pode ser visto como um processo estocástico.

Estes processos serão de uma natureza especial para o caso de problemas de otimização onde a incerteza (Ξ_t) é *stagewise independent*. Nestas condições, podemos descrever o problema através de uma recorrência de programação dinâmica, e determinar as funções de custo futuro a partir da recorrência (2.7):

$$Q_t(x_t, \xi_t) = \min_{(x_{t+1}, y_t)} c^\top y_t + \beta \bar{Q}_{t+1}(x_{t+1}) \quad (4.1)$$

$$\text{s.a. } A_t x_{t+1} + B_t x_t + W_t y_t = d_t,$$

$$\bar{Q}_{t+1}(x_{t+1}) = \rho_t [Q_{t+1}(x_{t+1}, \xi_{t+1})] \quad (4.2)$$

Esta formulação mostra que as decisões x_{t+1} e y_t dependem apenas do estado inicial x_t e de uma realização particular ξ_t da incerteza no tempo t .

Este tipo de dependência é dito *Markoviano*, como detalhado no apêndice B, em particular nas definições B.3 e B.41. Assim, como a determinação do valor de x_{t+1} depende apenas do estado atual x_t , e de uma variável aleatória, independente, ξ_t , o processo (X_t) , associando as decisões a cada um dos cenários, será de fato uma cadeia de Markov.

Já o caso das variáveis de decisão locais y_t é um pouco diferente. De fato, estas também dependem de x_t , e não do valor de y_{t-1} , para que a sequência (Y_t) seja uma cadeia de Markov. Mas, ao incluirmos os valores de x_t , obtemos uma cadeia de Markov estendida

$$(X_t, Y_t)_{t=1}^\infty$$

pois agora incorporamos a informação de estado necessária para que Y_t dependa apenas da realização anterior do processo.

Precisamos de uma notação para descrever a evolução das decisões ao longo do tempo, dada uma distribuição inicial $\mathbb{P}$ para o estado x_t . Isto é importante para podermos identificar quando uma distribuição será estacionária. Ao tratarmos as decisões (e estados) como variáveis aleatórias, passa a ser mais importante estudar a *distribuição* das decisões do que a correspondência de uma decisão particular com o cenário associado a ela.

Note que, se o estado x_t segue uma distribuição de probabilidade, o estado x_{t+1} depende de *duas distribuições*: $\mathbb{P}$, a distribuição de x_t , e Ξ_t , que representa a incerteza observada no tempo t . Para isto, vejamos como definir uma função de transição para a cadeia de Markov.

Definição 4.4. Seja Ξ_t a distribuição da incerteza no tempo t . A função de transição $K_t(x_t, x_{t+1})$ dá a probabilidade de chegar no estado x_{t+1} a partir do estado x_t . Esta função é determinada a partir de Ξ_t , da função de custo futuro $\bar{Q}_{t+1}$ e da recorrência de programação dinâmica.

$$K_t(x_t, x_{t+1}) = \sum_{\substack{\text{cenários } \xi_t \text{ cujo estado final é } x_{t+1} \\ \text{para o estado inicial } x_t.}} \mathbb{P}(\Xi_t = \xi_t)$$

Vemos, pela própria definição, que esta função de transição satisfaz as condições dadas em B.38.

A partir desta função de transição, vamos introduzir uma notação para a evolução das distribuições de probabilidade ao longo do tempo.

Definição 4.5. A distribuição de probabilidades do estado X no estágio $t + 1$ é calculada a partir de

$$\mathbb{P}_{t+1}(x_{t+1}) = \sum_{x_t} \mathbb{P}_t(x_t) K(x_t, x_{t+1}).$$

Esta operação será denotada por $\mathbb{P}_{t+1} = \mathcal{F}_t(\mathbb{P}_t)$, onde o $\mathcal{F}$ é mnemônico para *forward*.

Com isto, podemos finalmente introduzir a definição de estado estacionário para problemas de otimização estocástica com período 1. Esta definição captura a ideia que, se no tempo t a distribuição de X_t for dada por $\mathbb{P}$, então no tempo seguinte a distribuição de X_{t+1} será a mesma.

Definição 4.6 (Estado estacionário). Em um problema de otimização de período 1, dizemos que a distribuição de probabilidade $\mathbb{P}$ é um estado estacionário para o problema se

$$\mathbb{P}_1 \sim \mathcal{F}_1(\mathbb{P}_1).$$

Enfim, podemos generalizar a definição acima para problemas de período p :

Definição 4.7 (Regime periódico). Em um problema de otimização de período p , dizemos que as distribuições de probabilidade

$$(\mathbb{P}_1, \mathbb{P}_2, \dots, \mathbb{P}_p)$$

formam um regime periódico para o problema se

$$\mathbb{P}_2 \sim \mathcal{F}_1(\mathbb{P}_1) \quad \mathbb{P}_3 \sim \mathcal{F}_2(\mathbb{P}_2) \quad \mathbb{P}_1 \sim \mathcal{F}_p(\mathbb{P}_p). \quad (4.3)$$

4.2 Existência

Se o espaço de estados de um problema de otimização estocástica fosse finito, ele sempre possuiria pelo menos um regime periódico, como demonstrado em B.11. Para espaços mais gerais, é necessário impor outras condições.

No nosso caso, temos um espaço de estados contínuo e limitado, dado pelos limites dos reservatórios (e das vazões no caso de modelos autorregressivos). Mais ainda, é razoável assumir que para o problema do planejamento da operação também observaremos um regime periódico, já que, mesmo quando o problema é resolvido com horizonte finito, é possível observar o estabelecimento de uma periodicidade, por exemplo para as trajetórias dos armazenamentos.

4.3 Convergência

Para além da existência do regime periódico, a teoria de Cadeias de Markov indica que, sob certas hipóteses, este regime periódico irá se estabelecer, independente da condição inicial da cadeia. Por exemplo, para cadeias em espaços de estado finitos, este resultado é dado pelo teorema B.14, que pede que a matriz de transição seja *regular*. Para espaços de estado enumeráveis, o teorema B.37 também dá convergência, sob condições um pouco mais restritivas para evitar que o estado “desapareça no infinito”.

A convergência para o regime periódico é importante por duas razões. Primeiramente, ela garante que o processo de otimização irá estabilizar ao gerar cenários *forward*, pois ao seguir a cadeia de Markov os estados irão se aproximar cada vez mais da distribuição estacionária correspondente. Assim, esta convergência também permite acelerar a convergência do algoritmo de planos cortantes descrito na seção 3, pois não será necessário gerar trajetórias arbitrariamente longas para obter cortes que representem de forma satisfatória as funções de custo futuro.

Em segundo lugar, a velocidade de convergência para o regime estacionário pode ser interpretada como uma capacidade de o sistema se recuperar face a eventos extremos. De fato, se no tempo $t = 1$ temos uma distribuição concentrada em um ponto, que são os armazenamentos observados, quanto mais rápido o sistema convergir para o estado estacionário, mais rápido a influência do estado presente irá desaparecer para considerar a tomada de decisões, ainda sob incerteza, dos estágios futuros.

Parte II

Implementação e Resultados

5 Implementação

5.1 Desenvolvimento do protótipo

O protótipo computacional desenvolvido ao longo das duas primeiras fases do projeto “*Metodologias de Programação Estocástica Inteira*”, estabelecido entre o ONS e a COPPETec/Instituto de Matemática da UFRJ, contratos GMC-CT-093/17 e PEM-CS-200/18, continuou a ser desenvolvido para realizar os experimentos numéricos desta terceira fase. Além das modificações inerentes ao tratamento de problemas periódicos em horizonte infinito, também realizamos nesta etapa a atualização do código para a versão 1.6 da linguagem *Julia*, além de atualizar as diversas bibliotecas por ele usadas.

5.1.1 Adaptação para horizonte periódico

A modificação mais importante do programa consiste em incluir uma nova variável de configuração, correspondente à duração do período, em estágios, e o seu uso ao longo de toda a parte de construção dos problemas de otimização. Isto permite realizar testes em problemas de período arbitrário, como iremos ver na seção 6.

Além disto, também adicionamos uma outra variável de configuração, correspondente ao número de estágios a serem simulados. Anteriormente, com problemas de horizonte finito, era natural realizar tantas simulações como estágios do problema. Agora, é fundamental poder simular um horizonte maior do que um período, para poder observar se houve o estabelecimento ou não de um regime estacionário. Assim, podemos por exemplo simular um horizonte de 20 estágios em um modelo estacionário, ou seja, de período 1 e com apenas um estágio.

5.1.2 Migração do código para *Julia* 1.6

A versão 1.6 da linguagem *Julia* foi lançada em 21 de Março de 2021, e se tornou a nova versão de suporte de longo prazo. Com ela, diversas melhorias das versões anteriores foram consolidadas, em particular quanto à gestão de pacotes e de compilação de código [BBD⁺21]. Em paralelo, novas versões das bibliotecas *JuMP* e *SDDP.jl* também foram lançadas, o que justificava uma atualização do código global.

Assim, uma parte inicial do esforço foi dedicada a adaptar nossa base de código para as novas versões dos pacotes. Com isto, ganhamos velocidade com as novas versões da biblioteca *JuMP*, e também puramente pelos ganhos de performance da linguagem *Julia* ao longo destes anos.

Entretanto, dada uma mudança substancial na forma como o pacote *SDDP.jl* realiza a fase *backward*, as modificações que havíamos introduzido para o cálculo de cortes de Benders reforçados não puderam ser adaptadas à versão 0.4.4. De fato, elas não seriam possíveis nem na versão 0.1.0, de 2020. Entretanto, os cortes de Benders reforçados são apenas necessários quando se trata de problemas com variáveis inteiras, o que não ocorre nos testes que vamos realizar. Em última análise, a versão anterior (0.1.0) do pacote *DisjHTPlan* ainda está disponível no repositório BitBucket para integrantes do convênio, e usa as versões antigas das bibliotecas *SDDP.jl* e *JuMP*.

5.1.3 Outras modificações e melhorias

Além das adaptações para horizonte infinito e periódico, e às novas versões da linguagem e das bibliotecas, também realizamos outras melhorias no protótipo.

Em primeiro lugar, passamos a utilizar a construção `@expression` para construir expressões auxiliares, em vez de adicionar variáveis e restrições ao modelo. Isto permite recuperar a mesma informação, e construir um modelo totalmente equivalente, sem introduzir mais variáveis, o que pode prejudicar a performance do `solver` ao calcular a solução ótima do problema de otimização.

Também introduzimos uma nova variável de configuração, correspondente ao arquivo de demanda de energia a cada estágio, separando-a do arquivo com as configurações do sistema. Como é de praxe executar testes com uma demanda levemente superior à histórica, para observar efeitos de estresse no sistema, adaptamos o código e o arquivo `.ini` de parâmetros para ler um arquivo a mais. Isto resulta em mais flexibilidade ao usuário para realizar testes do modelo, e também do sistema.

Mais ainda, uniformizamos o acesso aos dados para a realização de gráficos, utilizando a estrutura de `DataFrames`.

Por fim, adaptamos o protótipo para utilizar, por default, a biblioteca `GLPK` [con20] como *solver* de problemas de programação linear. Esta é uma biblioteca de código aberto, o que permite maior acesso ao protótipo por parte do ONS. Em sequência a esta substituição, também adicionamos dois novos *solvers* lineares aos habilitados ao protótipo: o `Clp`, do projeto COIN-OR [con22], e o `HiGHS` [HH18].

5.2 Ambiente de programação

Os casos de prova de conceito, utilizados como exemplo na seção 6, foram executados com o seguinte ambiente de programação:

- Hardware
 - CPU: Intel (R) Celeron(R) CPU 3865U (2 cores, 2 threads e clock de 1.80 GHz)
 - RAM: 4Gb RAM, DDR4, 2133 Mhz
- Ubuntu Linux 18.04.5 LTS
- Linguagem de Programação: Julia 1.6.1
- Solver: GLPK 5.0
- Versão dos pacotes Julia utilizados:

Pacote	Versão	commit (SHA-1)
Desenvolvido para o convênio:		
DisjHTPlan	0.2.99	5a5da962
Dependências DisjHTPlan:		
ConfParser	0.1.2	88353bc9
JuMP	0.22.3	4076af6c
NPZ	0.4.1	15e1cf62
PyPlot	2.10.0	d330b81b
Requires	1.3.0	ae029012
SDDP	0.4.4	f4570300
Random		9a3f8284
Statistics		10745b16
Outros pacotes Julia:		
CSV	0.10.4	336ed68f
DataFrames	1.3.4	a93c6f00
GLPK	0.15.3	60bf3e95
JSON	0.21.3	682c06a0
OSQP	0.7.0	ab2f91bb
PyPlot	2.10.0	d330b81b

Já os casos mais demandantes, estudados na seção 7 foram executados em um *container docker* dedicado:

- Hardware
 - CPU: Intel(R) Core(TM) i7-12700 (10 cores, 20 threads e clock de 2.10 GHz)
 - RAM: 12b RAM, DDR4, 2133 Mhz
- Ubuntu Linux 22.04.5 LTS
- Linguagem de Programação: Julia 1.6.6
- Solver: GLPK 5.0 e HiGHS 1.4
- Versão dos pacotes Julia utilizados:

Pacote	Versão	commit (SHA-1)
Desenvolvido para o convênio:		
DisjHTPlan	0.2.99	5a5da962
Dependências DisjHTPlan:		
ConfParser	0.1.2	88353bc9
JuMP	1.1.1	4076af6c
NPZ	0.4.2	15e1cf62
PyPlot	2.10.0	d330b81b
Requires	1.3.0	ae029012
SDDP	0.4.6	f4570300
Random		9a3f8284
Statistics		10745b16
Outros pacotes Julia:		
CSV	0.10.4	336ed68f
DataFrames	1.3.4	a93c6f00
GLPK	1.0.1	60bf3e95
HiGHS	1.1.4	87dc4568
JSON	0.21.3	682c06a0
OSQP	0.8.0	ab2f91bb
PyPlot	2.10.0	d330b81b

5.3 Contribuições à comunidade Julia

Ao realizar estudos de casos para problemas de otimização de grande porte, com um grande número de variáveis de decisão a cada estágio, também acabamos estressando mais os *solvers* usados para resolver cada PL durante a PDDE.

Assim, especialmente após adicionar um grande número de cortes, pudemos observar que os *solvers* tinham dificuldade ao resolver os PLs, resultando às vezes em erros incompatíveis com a formulação. Por exemplo, às vezes o *solver* terminava acusando inviabilidade, o que é impossível visto que há recurso relativamente completo. Em outros casos, o *solver* terminava sem calcular os multiplicadores duais por instabilidade numérica.

Nestas circunstâncias, especialmente devido ao uso do algoritmo do simplexo nas variáveis duais, pode ocorrer que a base usada pelo *solver* acumule erro numérico progressivamente, por exemplo ao resolver diversas vezes um PL similar durante a etapa *backwards*. Assim, mesmo que, matematicamente, o problema esteja bem posto, o *solver* irá utilizar uma condição inicial que pode ter tanto erro numérico que ele acusa algum erro. A solução possível para este tipo de situação é reinicializar o problema, descartando as condições iniciais pré-calculadas. Este tipo de operação é computacionalmente custoso, visto que a grande vantagem do algoritmo do simplexo dual é que partir de uma solução anterior para um problema similar reduz drasticamente o número de iterações necessárias para a convergência.

Uma parte destas condições já era testada no código da biblioteca `SDDP.jl`, mas não todas, o que gerou diversos erros que interromperam o processo da PDDE. Estes erros aconteceram tanto com o *solver* GLPK como com o HiGHS. Após identificar mais precisamente outros casos que também poderiam impactar a solução dos problemas via PDDE, realizamos uma modificação no código do pacote, para uso nos últimos casos de estudo que fizemos. Naturalmente, também disponibilizamos esta modificação para que ela fosse incorporada no código oficial do pacote, o que já aconteceu: <https://github.com/odow/SDDP.jl/pull/531>. Até a próxima versão do pacote `SDDP.jl` (prova-

velmente a 0.4.7), esta estará disponível apenas para quem usar a versão de desenvolvimento

Com isto, também podemos reduzir o tempo médio de cálculo, seja para a determinação da política (a PDDE propriamente dita) seja para o cálculo da simulação final. De fato, em geral o GLPK resolve os PLs em um tempo da ordem de poucos milissegundos - muitas vezes por já possuir uma boa solução inicial. Assim, reduzimos o tempo máximo por PL a 60 segundos, o que é amplamente suficiente para resolver os problemas numericamente estáveis, e rápido o suficiente para interromper um problema em que o *solver* esteja apresentando dificuldades.

6 Prova de conceito

Nesta seção, apresentamos os resultados da aplicação da metodologia de horizonte infinito em diversos problemas de otimização estocástica. Primeiro, iremos descrever os casos selecionados, e, a seguir, apresentar as análises em cada grupo.

6.1 Casos de estudo

Os problemas que destacamos para esta prova de conceito são todos referentes ao planejamento da operação energética de longo prazo. Por simplicidade, iremos considerar a operação do SIN como um único mercado, ou seja, ignoraremos as restrições de transmissão entre os diferentes mercados do país. Com isto, podemos considerar apenas um Reservatório Equivalente de Energia (REE), o que simplifica a análise da função de custo futuro, pois esta dependerá de apenas uma variável de estado. O caso com 4 REEs, mais próximo do modelo atual, será abordado na seção 7.

No estudo que se segue, dividimos os casos em três grupos:

1. Problemas com estágio de um ano: estes problemas são estacionários, pois a periodicidade da incerteza é também de um ano. Aqui, os objetivos principais são: validar a metodologia, auferir o estabelecimento de um regime estacionário, e iniciar a análise dos resultados no caso mais simples possível.
2. Problemas com estágio de um semestre: estes problemas terão período dois, e com isto já podemos observar o impacto da sazonalidade na solução. O objetivo principal da comparação dos dois casos apresentados é analisar o impacto de uma sazonalidade mais ou menos marcada, tanto na oscilação das soluções como no tempo para atingir o regime estacionário. Além disso, observaremos as diferenças para a solução do regime anual, advindas da introdução da sazonalidade.
3. Problemas com estágio de um mês: aqui, nos aproximamos do modelo de planejamento da operação, e também aproveitaremos para comparar a solução de um problema de 120 estágios com um periódico de 12 estágios.

6.2 Regime anual

Os problemas de regime anual são os mais simples, e por eles começaremos as análises e gráficos. Com isto, esperamos que os casos subsequentes, com mais detalhes, sejam mais facilmente compreendidos. De forma esquemática, os problemas anuais têm apenas um estágio periódico, como podemos ver na figura 6.1.

Resumimos, na tabela 1, as principais características do modelo utilizado. Notamos, principalmente, que a capacidade de geração hidráulica é superior à demanda, não se constituindo uma restrição ativa. Mais ainda, a energia afluyente média, somada à geração térmica mínima, não é

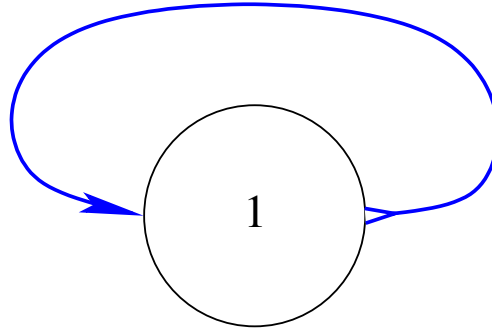


Figura 6.1: Representação esquemática do estágio periódico no caso anual

suficiente para atender a demanda. Por fim, a menor energia afluyente, somada à geração térmica máxima, está abaixo da demanda, da ordem de 45.000 MWmês.

Parâmetro	Valor (MWmês)
Energia Armazenada máxima	291.579,1
Geração Hidráulica máxima	945.422,4
Geração Térmica máxima	279.876,0
Geração Térmica mínima	76.639,6
Demanda de energia	800.951,0
Energia Afluyente média	707.212,0
Energia Afluyente máxima	1007.745,5
Energia Afluyente mínima	477.102,7

Tabela 1: Principais características do modelo anual

Também apresentamos na figura 6.2 a curva de custo variável unitário (CVU) em função da geração térmica despachada, ou seja, para além da geração térmica mínima (“inflexibilidade”).

Na figura 6.3 podemos ver os gráficos de 500 séries simuladas para o regime anual, iniciando cada ano em Novembro, que é o início da estação úmida. Cada série tem a duração de 20 anos, com o objetivo de detectar, de forma clara, o estabelecimento - ou não - de um regime estacionário. Em quase todos gráficos, a unidade é o MWmês; como um estágio corresponde a um ano, o valor indicado é a geração (ou deficit, vertimento, etc.) daquele ano, medida em MWmês. A exceção é o custo de operação, dado em Reais. Vale notar que este é o custo da operação sem adicionar o custo de geração da inflexibilidade térmica.

O que primeiro salta à vista é que há uma oscilação bastante grande de um ano ao seguinte para, essencialmente, todas as variáveis. Isto é reflexo do regime de afluentes, que pode ser bastante diferente de um ano para o seguinte, como pode ser visto na figura 6.3g. Notamos que estas grandes variações também ocorrem na própria série histórica de vazões anuais.

Em seguida, podemos ver que os casos extremos, de deficit, na figura 6.3a, e de vertimento, na figura 6.3b, são relativamente raros. Isto é bem claro para o deficit, onde há apenas umas poucas curvas que não estão em zero. No caso de vertimento, podemos notar que há de fato menos curvas que saem de zero do que as que vemos nos outros gráficos, ao comparar o nível de cinza dos gráficos. Esta análise ficará muito mais simples ao analisarmos os gráficos de quantis, na figura a seguir.

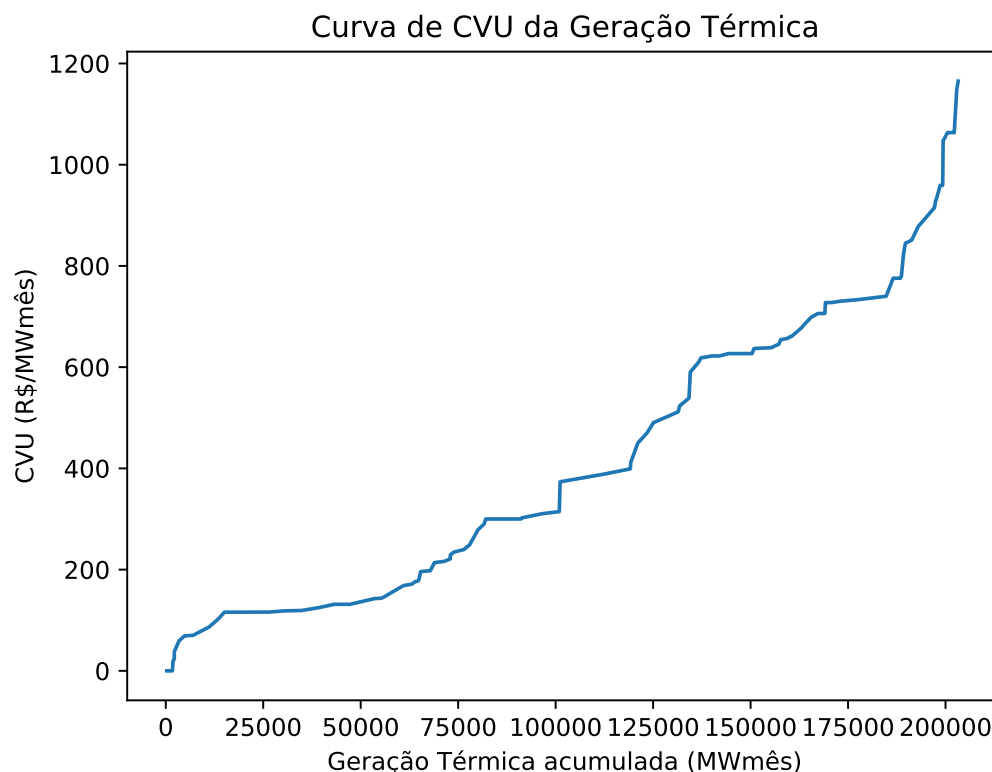


Figura 6.2: CVU da geração térmica no caso anual

Antes disto, vejamos o que acontece com as variáveis de geração térmica e hidráulica. Podemos notar que estas se comportam de forma complementar, o que é esperado visto que, a menos de deficit, a soma da geração térmica com hidráulica deve ser igual à demanda anual. Também notamos que, de forma geral, a geração hidráulica se concentra em torno do máximo da geração anual, enquanto que a geração térmica, no mínimo. Notemos que neste gráfico não incluímos a geração devido à inflexibilidade das usinas térmicas.

O custo de operação segue, essencialmente, o custo das usinas térmicas, já que a ocorrência de deficit é bastante rara. Entretanto, ao comparar com o gráfico de geração térmica, notamos que ele é mais extremo, o que se explica pelo custo marginal crescente quanto maior a geração térmica.

Uma outra métrica relevante é dada pela energia armazenada ao final de cada estágio. Aqui, podemos observar excursões ao longo de toda a capacidade de armazenamento dos reservatórios, ao longo do período destacado. Também notamos que há séries que permanecem por mais de um período “totalmente vazios”, assim como séries que se mantêm “totalmente cheios” por mais de um estágio. Deste gráfico, entretanto, não é imediato observar uma mudança entre os primeiros e últimos estágios, o que, mais uma vez, ficará mais evidente ao observarmos os gráficos de quantis e violinos, a seguir.

Enfim, incluímos as séries de energia afluyente, para ilustrar dois pontos: primeiramente, que não são sempre os mesmos cenários a cada estágio, o que gera uma certa flutuação ao longo do tempo, naturalmente. Em seguida, para destacar a variabilidade ao longo do tempo, onde anos muito secos são seguidos de anos muito úmidos, o que também reflete na grande variabilidade anual de todas as outras séries aqui apresentadas.

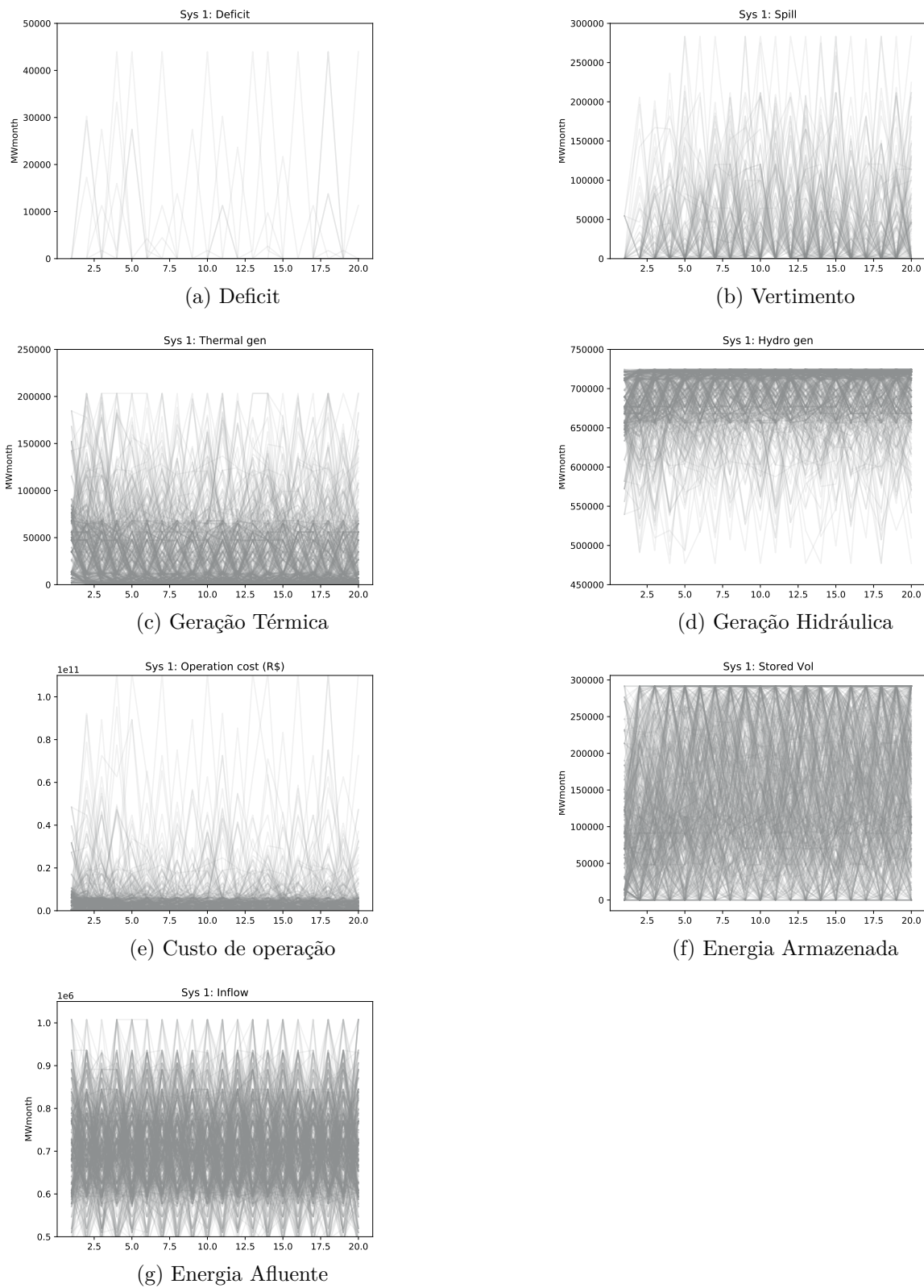


Figura 6.3: Resultado de simulação de 500 séries para Regime anual

Agora, passamos aos gráficos de quantis, na figura 6.4. Nestas figuras, estão representados os quantis de 5%, 10%, ..., 90% e 95% para cada uma das variáveis de interesse exploradas nos gráficos de séries da figura 6.3. Estes quantis devem ser entendidos “a cada estágio”, ou seja, os percentis representados em cada estágio referem-se à distribuição de probabilidade dos valores observados naquele estágio. Com isto, iremos analisar estas distribuições, em vez de nos focarmos nas trajetórias correspondentes a cada cenário: assim, poderemos melhor observar a convergência ao longo dos estágios.

Aqui, podemos já notar de forma mais quantitativa algumas das observações que fizemos anteriormente. Em primeiro lugar, o gráfico do deficit deixa claro que, em nenhum ano, houve probabilidade maior do que 5% para déficit, já que até o quantil 95 o valor era zero. Ademais, vemos que o vertimento também é bastante limitado, de forma geral. Apenas estão não-nulas as curvas dos quantis de 95% e 90%, em valores relativamente baixos se comparados à demanda anual de energia, que é da ordem de 800 GW-mês, e também são compatíveis com a flutuação da energia afluenta, que oscila entre 540 e 1000 GW-mês nas séries históricas.

Outro fenômeno que podemos observar é uma tendência, entre o início e o final do período, de estabilizar os quantis em patamares diferentes dos iniciais, indicando que o sistema está buscando uma configuração de equilíbrio diferente da que saiu. Isto é bastante visível no caso do vertimento, onde o aumento do vertimento ao longo dos anos indica que, a princípio, o armazenamento estava abaixo do equilíbrio. Mas também podemos notar esta tendência na geração térmica e no custo de operação que, ambos, diminuem conforme a simulação avança, indicando que no início houve um maior esforço térmico para preservar os volumes dos reservatórios. De forma oposta, a geração hidráulica e a energia armazenada tendem a aumentar, com a energia armazenada no limite do reservatório em mais de 10% dos cenários a partir do quarto ano. Isto pode ser visto pela coincidência de duas curvas de quantis, a mais alta de 95% e a seguinte, de 90%, no quarto ano.

Não podemos deixar de notar, entretanto, que ainda há, pelo menos, 5% dos cenários em que o reservatório fica completamente deplecionado ao final do ano, o que indica que ainda há anos com uma significativa escassez de afluições.

Enfim, notamos no gráfico dos quantis para as afluições que, novamente, mesmo que estas sejam amostradas uniformemente, em diferentes anos a distribuição empírica não é exatamente a mesma. Como, ademais, os valores da distribuição são discretos, podemos ver que há mudanças bruscas entre os quantis amostrais. Isto mais uma vez indica que não devemos esperar um gráfico perfeitamente constante para garantir que foi atingido o regime estacionário. Pelo contrário, apenas pela flutuação amostral, como neste caso, é possível ver alterações nas distribuições, neste caso evidenciadas pelos quantis.

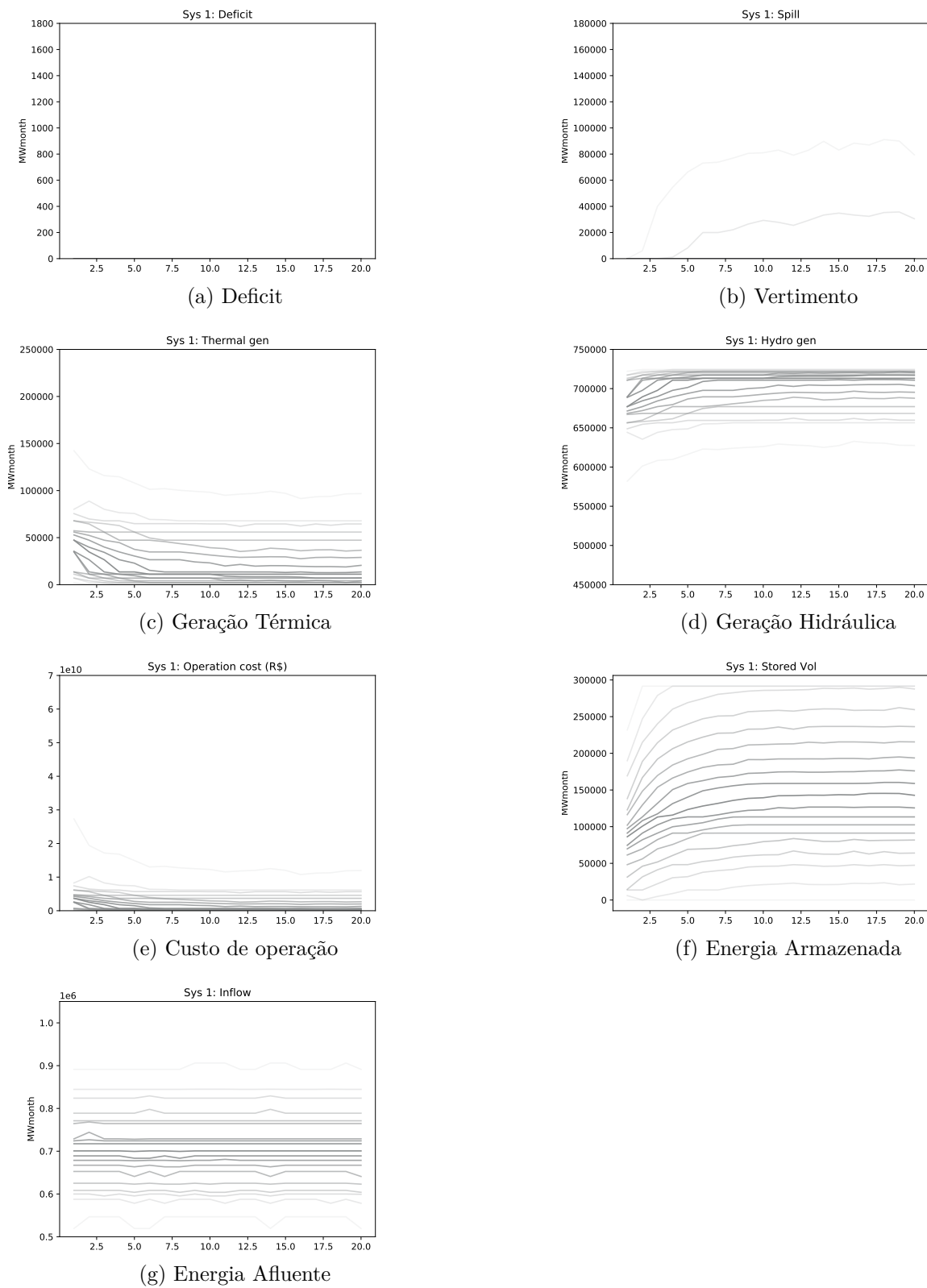


Figura 6.4: Quantis em passos de 5% para Regime anual

Finalmente, retomamos a análise através dos gráficos em violinos, para os 12 primeiros estágios, ou seja, os 12 primeiros anos da simulação, apresentados na figura 6.5. Não mostramos todos os 20 estágios porque, para estes gráficos, os violinos ficariam muito finos, comprometendo nossa visualização. As larguras dos violinos indicam, para cada variável, a densidade de probabilidade estimada através da simulação dos cenários. Além disso, incluímos os valores da média e da mediana, com traços horizontais respectivamente em azul e preto, para auxiliar a análise destes gráficos.

Aqui, vemos de forma ainda mais nítida a transição entre o estado inicial e o comportamento a partir dos últimos anos de simulação da série. Isto é perfeitamente visível no caso da geração térmica e hidráulica, cuja média e medianas vão se dirigindo ao valor final das mesmas e se estabilizam a partir já do sexto ano da simulação. O mesmo se passa, na respectiva escala, para o caso do custo de operação.

Já a energia armazenada e o vertimento têm um comportamento ligeiramente diferente dos três acima. Em particular, podemos notar que a tendência de aumento de ambos só se esgota ao final de 10 anos de simulação. Mesmo que esta seja bem mais rápida no início e, em especial, nos três primeiros anos, ela ainda continua muito depois de as gerações térmica e hidráulica se haverem estabilizado. Vale notar, aqui, que é a energia armazenada, e não as decisões de geração, que é a variável de estado fundamental do processo. Assim, a principal variável para a cadeia de Markov só estabiliza, pelo menos nesta análise gráfica, depois de várias das “variáveis dependentes” haverem estabilizado. O vertimento, entretanto (e naturalmente) acompanha a tendência da energia armazenada; podemos dizer que é basicamente o vertimento que está concluindo o “ajuste” da energia armazenada ao seu valor estável, quando as decisões de geração já estão estabilizadas desde muito antes.

No caso particular do planejamento da operação, e do modelo de programação linear multi-estágio e estocástico que empregamos, uma explicação para isto vem de, justamente, as decisões serem tomadas utilizando modelos lineares, que tendem a selecionar vértices da região viável. Isto se traduz, aqui, para decisões de ligar ou desligar, totalmente, a geração térmica de uma usina, e, correspondentemente, diminuir ou aumentar a geração hidráulica. Assim, mesmo que as decisões de geração já se tenham estabilizado, a energia armazenada, que não está sujeita a estas descontinuidades, ainda não terá chegado à estabilidade.

Por fim, o gráfico do déficit apresenta um comportamento um pouco surpreendente, a princípio, pois é nulo no primeiro ano, aumenta no segundo, para em seguida ir, lentamente, diminuindo. O fato de não ocorrer déficit inicialmente se deve às condições iniciais, que, mesmo que desfavoráveis (já que o sistema tende a aumentar o armazenamento com relação ao valor inicial), não são tão extremas a ponto de gerar déficit. Ao contrário, ao longo dos primeiros anos já é possível ter uma sequência suficientemente crítica que leve ao deplecionamento dos reservatórios e ao consequente déficit. Entretanto, como, em média, o sistema tende a ficar com maior armazenamento, o montante de déficit também irá se reduzindo ao longo dos anos da simulação.

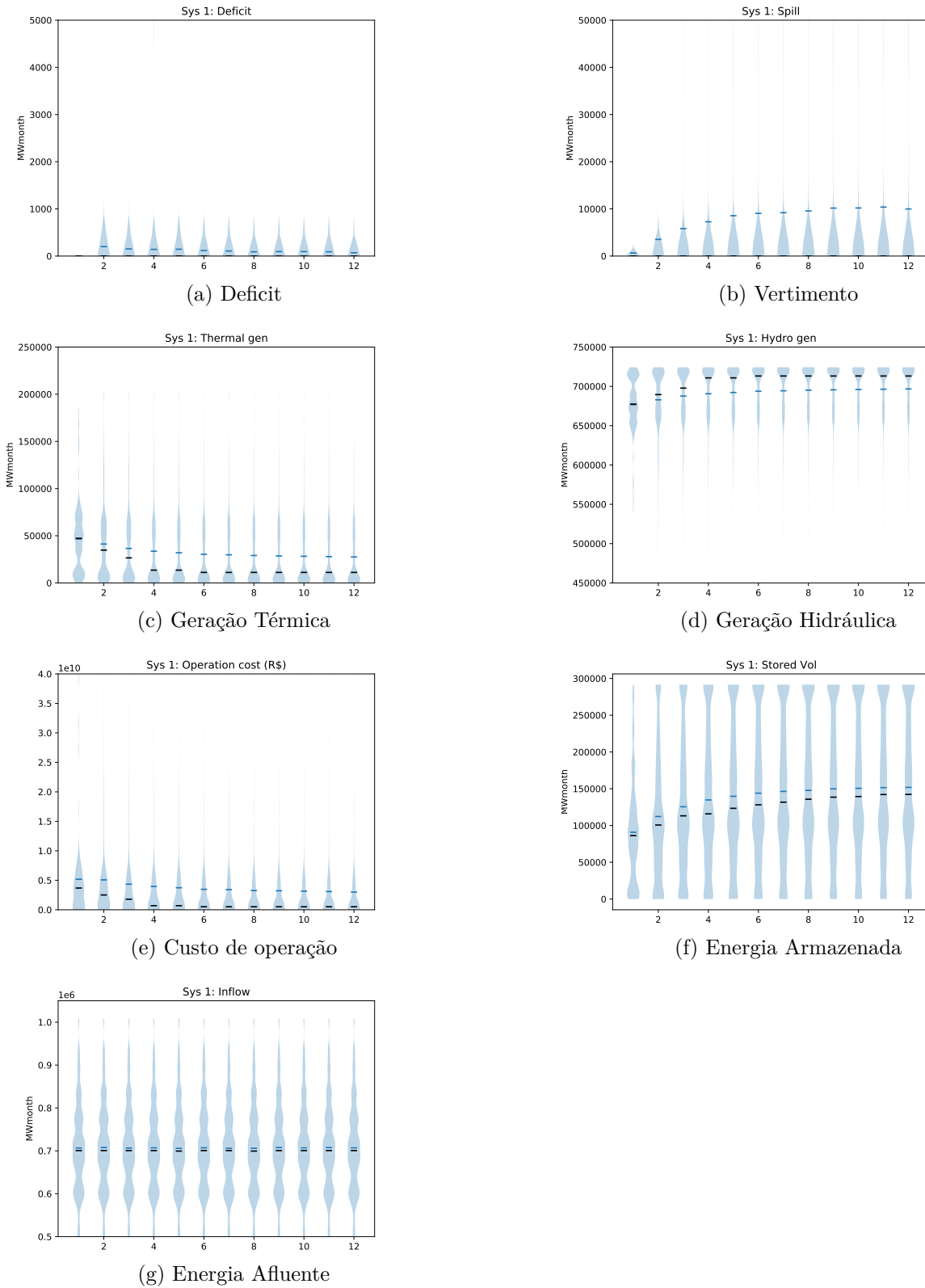


Figura 6.5: Histogramas em violinos dos 12 primeiros estágios para Regime anual

6.2.1 Análise do estado estacionário

Para avaliar de forma mais quantitativa a convergência para o regime estacionário, calculamos a distância entre as distribuições a cada estágio. Como realizamos 20.000 simulações, para garantir a independência das amostras, utilizamos mil cenários diferentes para cada estágio para calcular estas distâncias. Estas distâncias podem ser vistas na figura 6.6, onde cada quadrado corresponde à distância entre as distribuições de energia armazenada dos estágios entre linhas e colunas. Naturalmente, as distâncias ao longo da diagonal são sempre zero.

Distância entre CDFs de Energia Armazenada para diferentes estágios

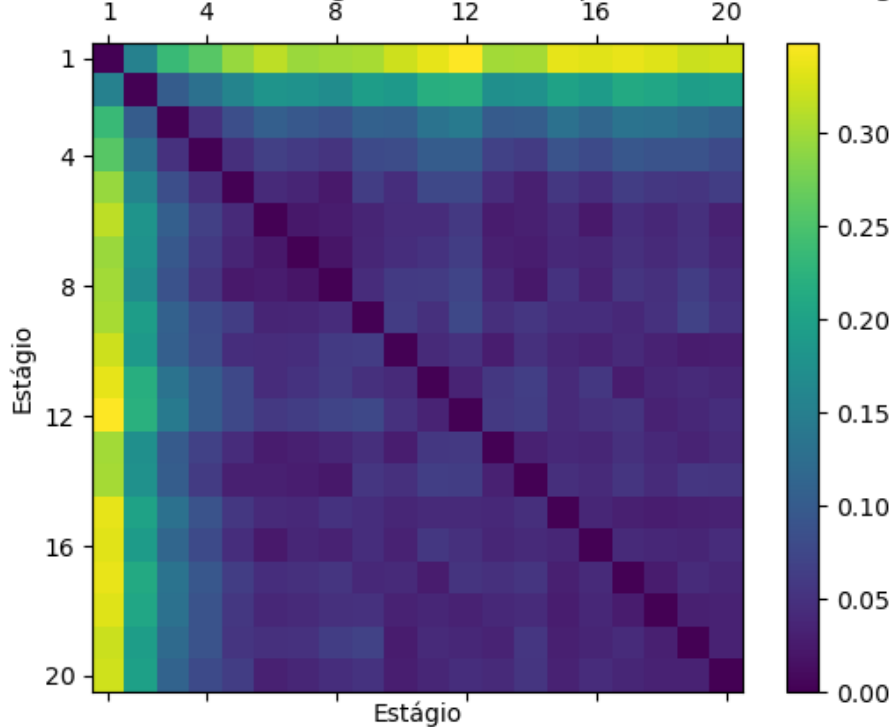


Figura 6.6: Distâncias entre distribuições de energia armazenada ao longo de 20 estágios no regime anual

Nesta figura, também podemos notar que a distribuição da energia armazenada ao final do primeiro estágio é significativamente diferente das outras, dando uma distância de mais de 0.3 entre as curvas de distribuição acumulada para a distância entre vários outros estágios. Em seguida, observando com mais detalhes, vemos que os estágios de dois a cinco ainda estão a uma certa distância dos últimos estágios, o que mostra que, de fato, ainda não está tão perto do estado estacionário, como pudemos observar nos gráficos de violinos da figura 6.3f. Por outro lado, também vemos uma clara tendência à redução das distâncias entre as distribuições conforme os estágios avançam, bastante notável pelas cores mais escuras no canto inferior direito do gráfico. Esta redução também é natural, já que esperamos que, conforme as simulações avançam, o sistema saia do estado inicial, particular, e convirja para o estado estacionário da cadeia de Markov.

Para compararmos alguns valores específicos, destacamos os estágios 1, 5, 10, 15 e 20. Na tabela 2, reportamos as distâncias entre as respectivas CDFs empíricas, obtidas através de amostras

Estágios		Distância
1	5	0.228
1	10	0.314
1	15	0.308
1	20	0.299
5	10	0.101
5	15	0.096
5	20	0.087
10	15	0.017
10	20	0.040
15	20	0.040

Tabela 2: Algumas distâncias entre CDFs de energia armazenada

independentes de tamanho 1000 para cada uma: Aqui, vemos que de fato a distância entre a CDF da energia armazenada do primeiro estágio e os seguintes é grande. Também notamos que a distância entre a CDF do quinto estágio e os outros é praticamente um terço das anteriores, mas ainda bem maior do que as distâncias que vemos entre as CDFs dos estágios finais.

Por fim, utilizamos as mesmas séries, separadas por cenários da mesma forma que para o cálculo da distância, para realizar um teste estatístico de proximidade das distribuições. Utilizamos o teste de Kolmogorov-Smirnov, que é adaptado para amostras de distribuições contínuas. Como a distribuição da energia armazenada é predominantemente contínua, salvo pela acumulação nos limites superior e inferior, este teste já nos dá uma primeira avaliação estatística do estabelecimento ou não do regime estacionário. Este teste utiliza a distância entre as CDFs, como vista no gráfico da figura 6.6, e, em função do tamanho das amostras dá a probabilidade, supondo que as amostras utilizadas são oriundas de distribuições contínuas, de que as distribuições subjacentes sejam iguais. Esta probabilidade é chamada de *p-valor* do teste de Kolmogorov-Smirnov.

Os p-valores calculados para as energias armazenadas ao longo dos 20 estágios podem ser vistos na figura 6.7. Nesta figura, vemos que de fato até o quarto estágio, as distribuições acumuladas ao final de cada estágio ainda estão bastante distantes entre si, e também das distribuições de estágios mais avançados da simulação. Em seguida, vemos um “bloco” formado pelos estágios de 5 a 8, onde a maior parte das distribuições está bem próxima entre si, mas claramente diferente dos primeiros estágios, e também não tão próxima assim das distribuições dos estágios seguintes. Há algumas exceções, como os estágios 13 e 14, que já dão um p-valor mais alto para serem equivalentes aos estágios de 5 a 8, mas de forma geral é razoável interpretar este gráfico como indicando que o estado estacionário só será atingido por volta do 10 ano, o que corrobora as análises feitas pelas médias e medianas da figura 6.3f.

Novamente, observemos os valores específicos dados pelos p-valores para os estágios 1, 5, 10, 15 e 20, na tabela 3, arredondados para 3 casas significativas: É notável a transição entre o primeiro estágio e os seguintes, onde a probabilidade de as amostras empíricas virem de uma mesma distribuição é baixíssima, para o quinto estágio, onde a probabilidade ainda é bem baixa, para os estágios finais, onde a probabilidade passa a ser bastante grande, o que confirma a nossa análise.

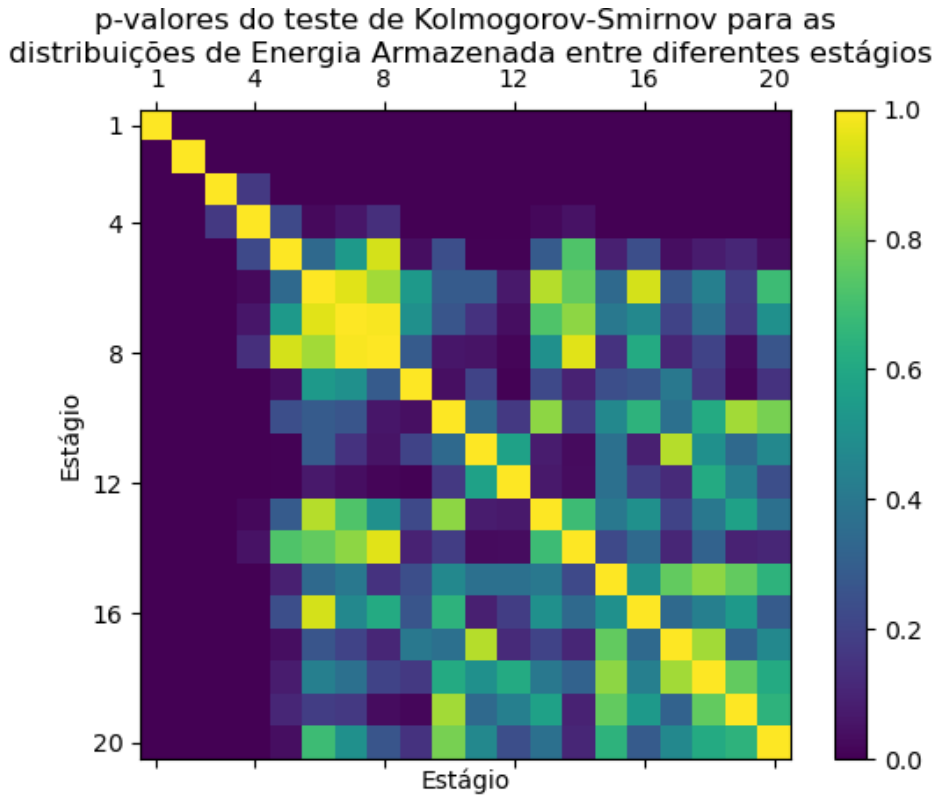


Figura 6.7: p-valores para o teste de Kolmogorov-Smirnov de 2 amostras, aplicado às distribuições de energia armazenada, simuladas ao longo de 20 estágios no regime anual

Estágios		p -valor
1	5	$3.4393 \cdot 10^{-23}$
1	10	$5.9008 \cdot 10^{-44}$
1	15	$2.7946 \cdot 10^{-42}$
1	20	$7.8477 \cdot 10^{-40}$
5	10	0.000073359
5	15	0.00019698
5	20	0.0010265
10	15	0.99872
10	20	0.40063
15	20	0.40063

Tabela 3: p-valores para o teste de Kolmogorov-Smirnov de similaridade entre CDFs de energia armazenada

6.2.2 Comparando o estabelecimento do regime estacionário

Após as análises apresentadas na seção anterior, é natural nos indagarmos se, ao iniciar de estados diferentes do que foi utilizado para aquele estudo, também obteríamos resultados semelhantes. Para isto, iremos analisar dois casos, um onde a energia armazenada inicial é nula, e um segundo onde o reservatório equivalente está cheio.

Nestes casos, alteramos o valor da energia armazenada inicial desde a etapa de cálculo da política pela PDDE. Naturalmente, se houvéssemos utilizado a política calculada a partir do exemplo anterior, e mudássemos a energia inicial apenas para a simulação de cenários, estaríamos observando o comportamento do sistema dinâmico determinado, essencialmente, pela função de custo futuro que lá fora calculada. Aqui, vamos um passo além: não apenas daremos uma outra condição inicial para as simulações, mas também para o algoritmo que constrói a função de custo futuro correspondente ao modelo periódico anual. Assim, ao observarmos os resultados, devemos ter em mente que a convergência para um mesmo regime implica, não apenas que o sistema de fato “esquece” o estado inicial, mas que a função de custo futuro calculada é, essencialmente, independente do estado inicial.

Começando pelo primeiro caso, com energia inicial zero no início da estação úmida, vemos na figura 6.8 as séries anuais geradas, a partir das mesmas séries de energias afluentes utilizadas para a figura 6.3. Aqui, o mais notável é quão similar os gráficos das séries são. Tanto nas grandes oscilações entre um ano e o seguinte, como na forma geral das séries, ambos os casos são muito semelhantes.

Entretanto, não podemos deixar de notar que há uma maior concentração de séries com alta geração térmica (e, eventualmente, de déficit) no início do período de simulação. Isto é de fato esperado, já que, ao iniciar com o reservatório vazio, o sistema terá que fazer um esforço ainda maior para recuperar o armazenamento e chegar ao nível ótimo para a operação em regime estacionário. Naturalmente, este aumento também leva a um maior aumento dos custos, que pode ser visto também em um maior número de séries altas no início do período; e a um menor número de eventos de vertimento no início da simulação.

Para podermos melhor comparar os regimes obtidos ao final de cada simulação, observemos na figura 6.9 como se comportam os quantis para esta configuração com energia inicial nula. Aqui, notamos que há uma probabilidade de déficit de mais de 5% no primeiro ano, o que não ocorreu no caso original, mas que este é o único ano para o qual esta probabilidade é tão grande: nos anos seguintes, todas as cuvas de quantis já estão em zero. Novamente, observamos o aumento do vertimento ao longo das simulações, chegando a patamares muito similares, porém demorando mais tempo para que os respectivos quantis 95% e 90% deixem de ser nulos.

O mesmo se pode dizer dos outros gráficos, de forma geral: começam de situações mais extremas do que no caso de estudo anterior, mas dão a impressão de convergirem, ao final do período de simulação, para um estado muito similar ao observado anteriormente. Assim, mesmo que demorando um pouco mais de tempo, temos a impressão de que o sistema possui um único estado estacionário, que é *atrator* do processo estocástico: independentemente das condições iniciais, tanto as decisões correntes (geração, déficit, etc.), como o estado (energia armazenada) convergem para uma mesma distribuição estacionária.

Finalmente, vejamos novamente o resultado deste caso através dos gráficos de violinos, na figura 6.10. A maior parte das observações já efetuadas anteriormente pode ser confrontada a estes gráficos.

Assim, por exemplo, o maior deficit no primeiro ano está bem nítido, e podemos ver que a distribuição de deficit no segundo ano também está maior do que anteriormente, reflexo de um estado ainda de baixo armazenamento.

Da mesma forma, vemos novamente uma geração térmica (e consequente custo operacional) bem mais importantes nos primeiros anos, mas que, igualmente, tende a reduzir conforme o passar das simulações. Ademais, podemos notar que os violinos para a geração térmica dos anos finais (digamos, a partir do sexto) já estão praticamente indistinguíveis, independentemente da energia armazenada inicial com que começamos a simulação.

Enfim, os violinos para a energia armazenada continuam sugerindo o mesmo fenômeno: a convergência para um regime estacionário. Aqui, é ainda mais visível o quão mais rápida é esta convergência nos estágios iniciais, quando o sistema se encontra bastante longe do estado estacionário. Naturalmente, este esforço se traduz por um alto custo de operação nestes mesmos primeiros anos, mas também por uma menor ocorrência de vertimentos, já que, em geral, o reservatório estará mais vazio do que no caso anterior. Isto mais uma vez reforça a dinâmica que leva à convergência ao estado estacionário, neste caso: primeiramente, o deficit e a geração térmica, que têm custos bem-definidos e em saltos, são usados para compensar a baixa energia armazenada. Uma vez que este esteja dentro de uma região estável, as proporções na geração hidrotérmica praticamente não se alteram, mas apenas o vertimento irá ser responsável por, paulatinamente, terminar de equilibrar a energia armazenada ao nível estacionário.

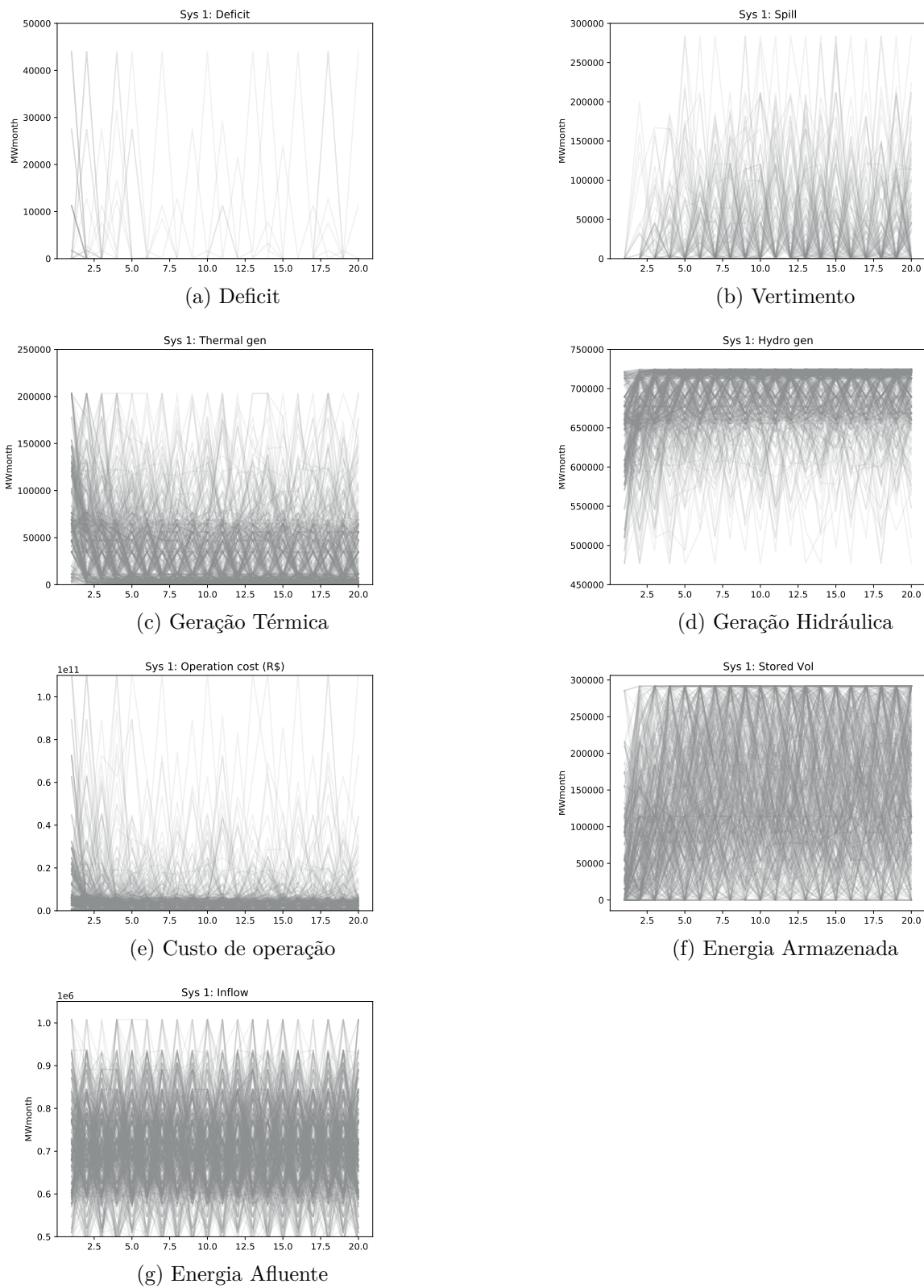


Figura 6.8: Resultado de simulação de 500 séries para Regime anual saindo do zero

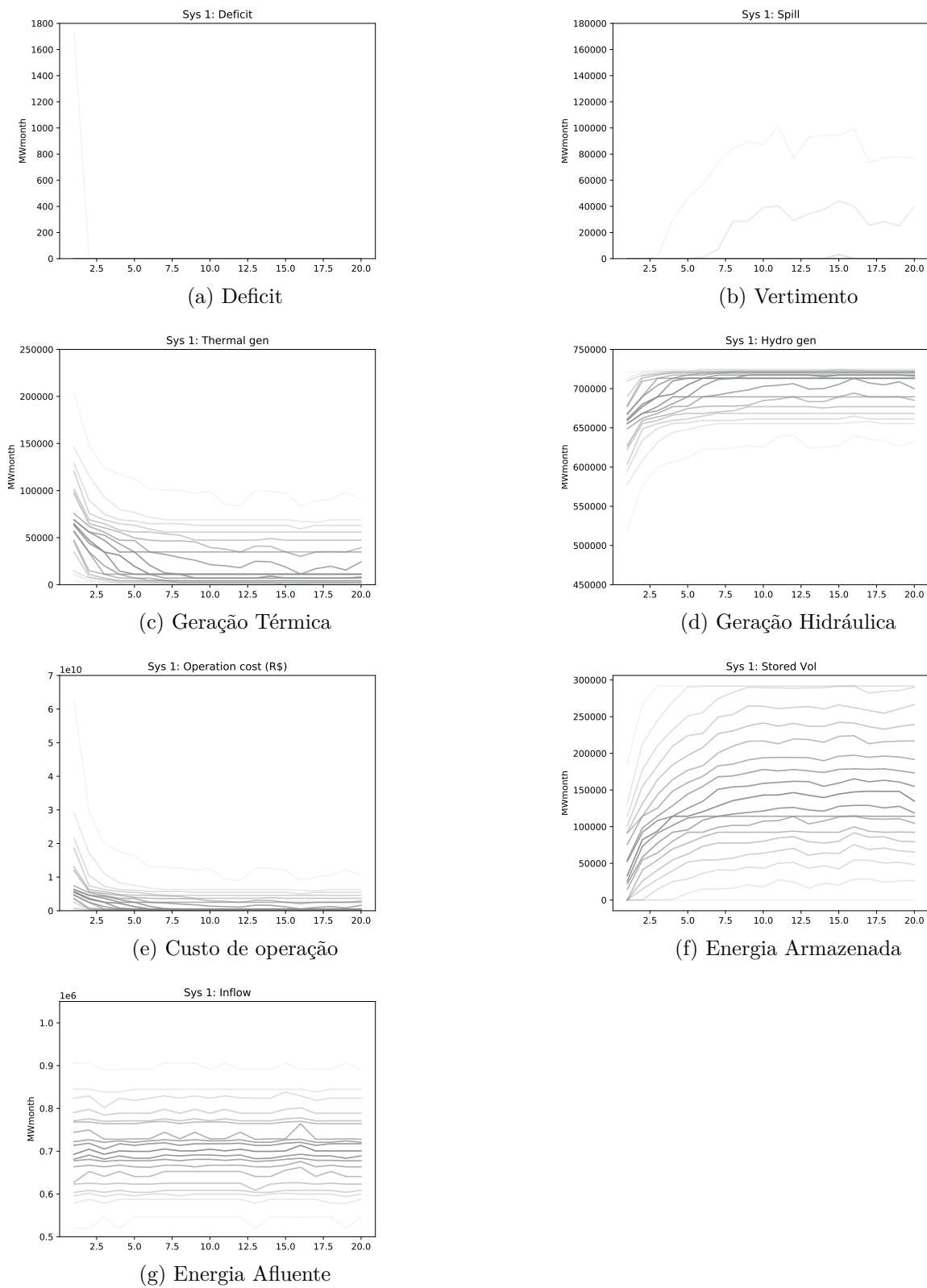


Figura 6.9: Quantis em passos de 5% para Regime anual saindo do zero

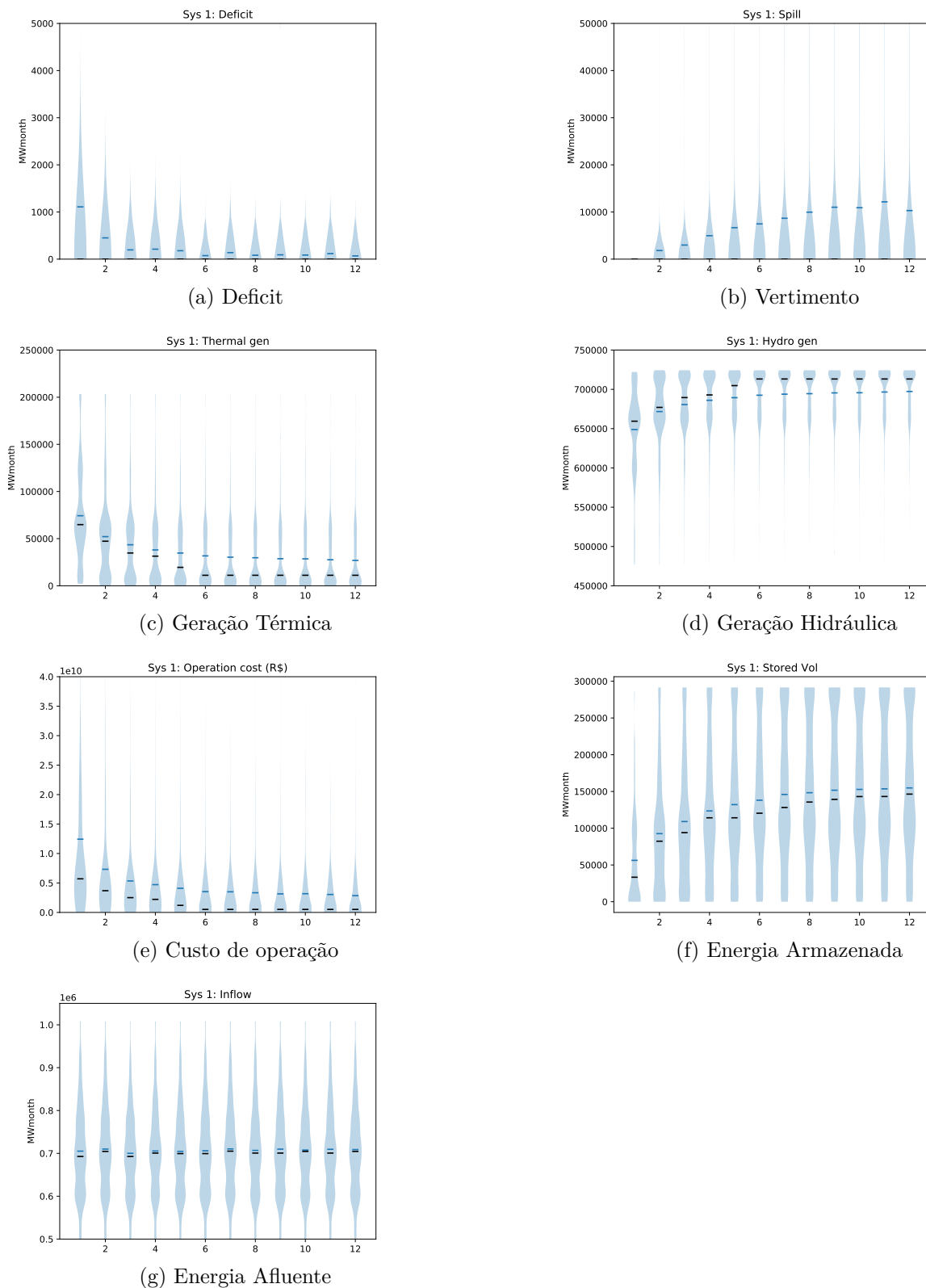


Figura 6.10: Histogramas em violinos dos 12 primeiros estágios para Regime anual saindo do zero

Continuando o estudo do caso anual, também ilustramos, a seguir, a operação no caso de partirmos de um reservatório com o armazenamento completo. Por simplicidade, apresentamos na figura 6.11 apenas os gráficos de violinos.

Como esperado, o comportamento é o oposto do que vimos nos dois casos anteriores, e o sistema tenderá a reduzir a energia armazenada, visto que, nesta situação, é muito comum ocorrer vertimento. Naturalmente, em cenários de grande afluência, o sistema continuará privilegiando uma grande geração hidráulica. Assim, a redução da energia armazenada média ao longo das simulações ocorre nos eventuais cenários de menor afluência, onde o sistema irá dispende a reserva de armazenamento para não incorrer em um custo de operação tão grande. Podemos notar, por exemplo, o aumento bem-definido da mediana de geração térmica, que acontece entre o oitavo e nono anos de simulação.

Também notamos que a convergência da distribuição do armazenamento ocorre de forma muito similar ao que observamos anteriormente. Num primeiro momento (ao longo dos dois primeiros anos) é bem rápida; a seguir, perde um pouco de velocidade (até o sexto ano); e enfim passa a ser bastante suave (do sétimo ano em diante).

Enfim, para concluir o estudo do caso anual, também apresentamos, na figura 6.12 as funções de custo futuro calculadas para cada um dos casos que vimos. Aqui, notamos que, de fato, as funções calculadas estão muito próximas, praticamente sobrepostas no gráfico. Isto demonstra que, de fato, a metodologia calcula os custos do regime periódico de forma consistente, pois esta função, em teoria, deve ser independente das condições iniciais. Mais ainda, isso também corrobora a nossa observação de que os regimes estacionários obtidos são similares: como a operação de cada estágio depende da energia armazenada inicial daquele estágio e da função de custo futuro, ao passar por vários estágios sucessivos é bastante razoável que o impacto da condição inicial desapareça, e subsista a função de custo futuro, que é sempre a mesma a ser aplicada.

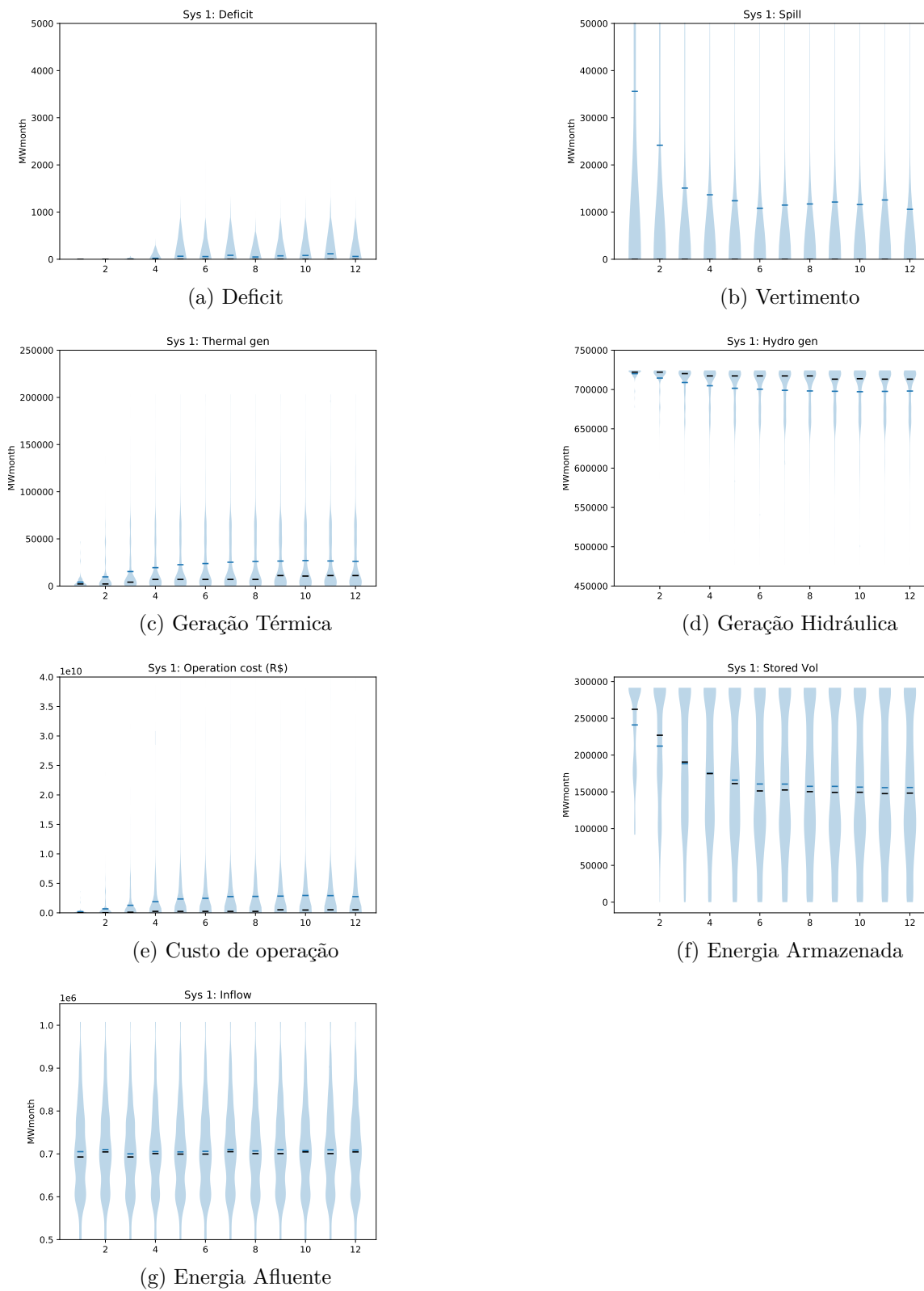


Figura 6.11: Histogramas em violinos dos 12 primeiros estágios para Regime anual começando cheio

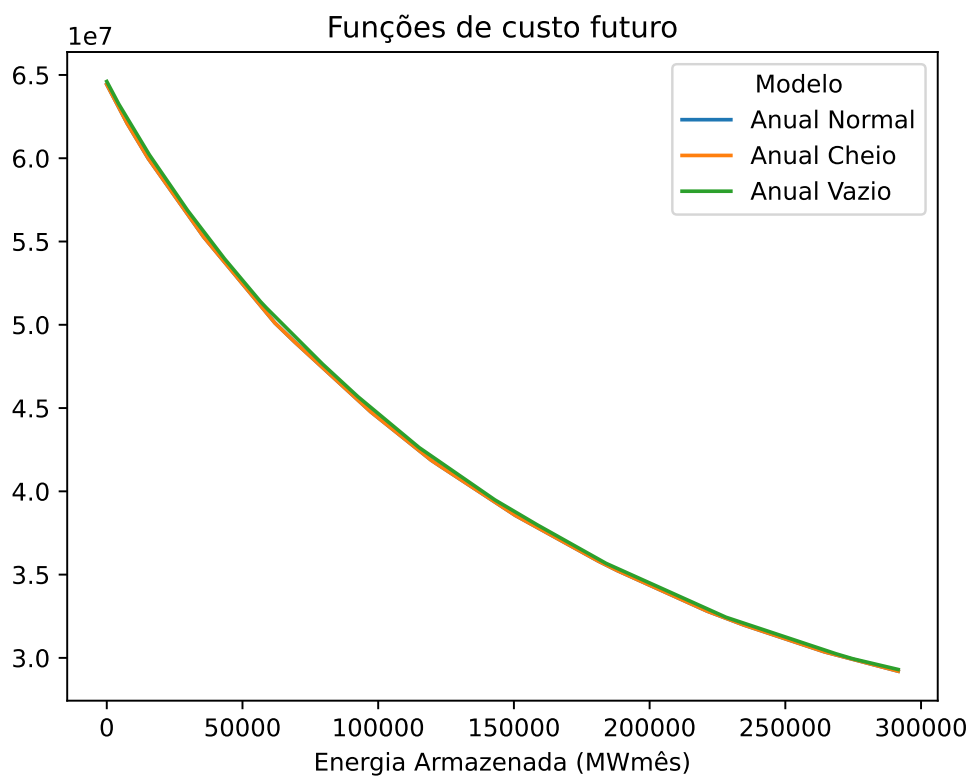


Figura 6.12: Funções de custo futuro no regime anual, calculadas a partir de condições iniciais diferentes para a PDDE.

6.3 Comparação dos regimes semestrais

Vejam, agora, dois problemas periódicos, mas não estacionários, para observar o efeito da *sazonalidade* em problemas de horizonte infinito, e, em particular, no regime limite que será atingido em cada um deles. Para tanto, adaptamos o modelo do SIN para considerar estágios *semestrais*. Com isto, podemos ver uma variação na energia afluyente, e também na demanda, devido à sazonalidade dos regimes de aflluência. Analogamente ao caso anual, ilustramos na figura 6.13 os dois estágios e a inclusão do *loop* do segundo de volta ao primeiro, que permite realizar simulações *forward* e, através do qual também passamos os cortes gerados ao resolver o problema do primeiro estágio para o segundo.

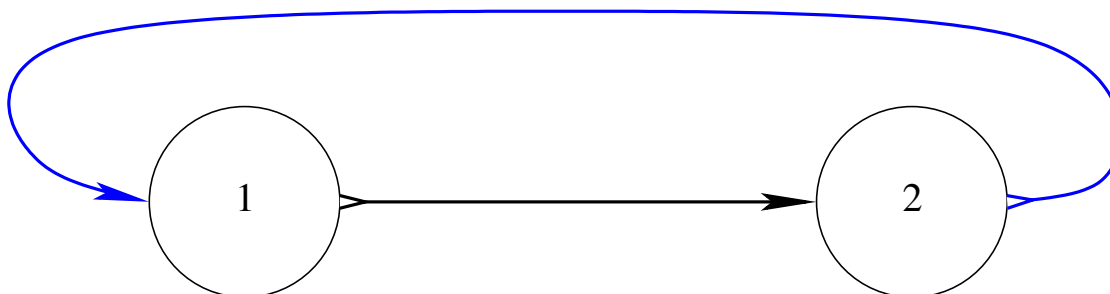


Figura 6.13: Representação esquemática de estágios no caso semestral

O primeiro modelo considera semestres de estações úmidas e secas, respectivamente, com um estágio correspondente aos meses de novembro a abril, e o outro, aos meses de maio a outubro. O segundo modelo usa estágios “mistos”, com o primeiro de fevereiro a julho, incluindo o final da estação chuvosa e o início da estação seca, e o segundo de agosto a janeiro, com o final da estação seca e o início da estação chuvosa, onde a sazonalidade das aflluências é bem menos marcada.

Como na seção anterior, apresentamos na tabela 4, as principais características do modelo utilizado. O tamanho dos reservatórios continua o mesmo, mas as capacidades de geração de energia, a cada estágio, estão divididas à metade, pois correspondem à mesma *potência* instalada.

Parâmetro	Valores (MWmês)			
	Nov/Maio Sazonalidade forte		Fev/Agosto Estações mistas	
	Úmida	Seca	“Úmida”	“Seca”
Energia Armazenada máxima	291.579,1			
Geração Hidráulica máxima	472.711,2			
Geração Térmica máxima	139.938,0			
Geração Térmica mínima	38.319,8			
Demanda de energia	412.802	388.149	403.878	397.073
Energia Afluyente média	460.354,8	246.857,2	404.175,9	304.729,1
Energia Afluyente máxima	660.287,2	372.042,6	575.516,4	448.948,0
Energia Afluyente mínima	286.311,9	153.115,2	247.088,5	202.330,0

Tabela 4: Principais características dos modelos semestrais

Também apresentamos na figura 6.14 a curva de custo variável em função da geração térmica despachada, ou seja, para além da geração térmica mínima (“inflexibilidade”). Ela é análoga à

curva do caso anual, apenas considerando a geração ao longo de um semestre, em vez de um ano, e portanto sua escala de energia é metade daquela.

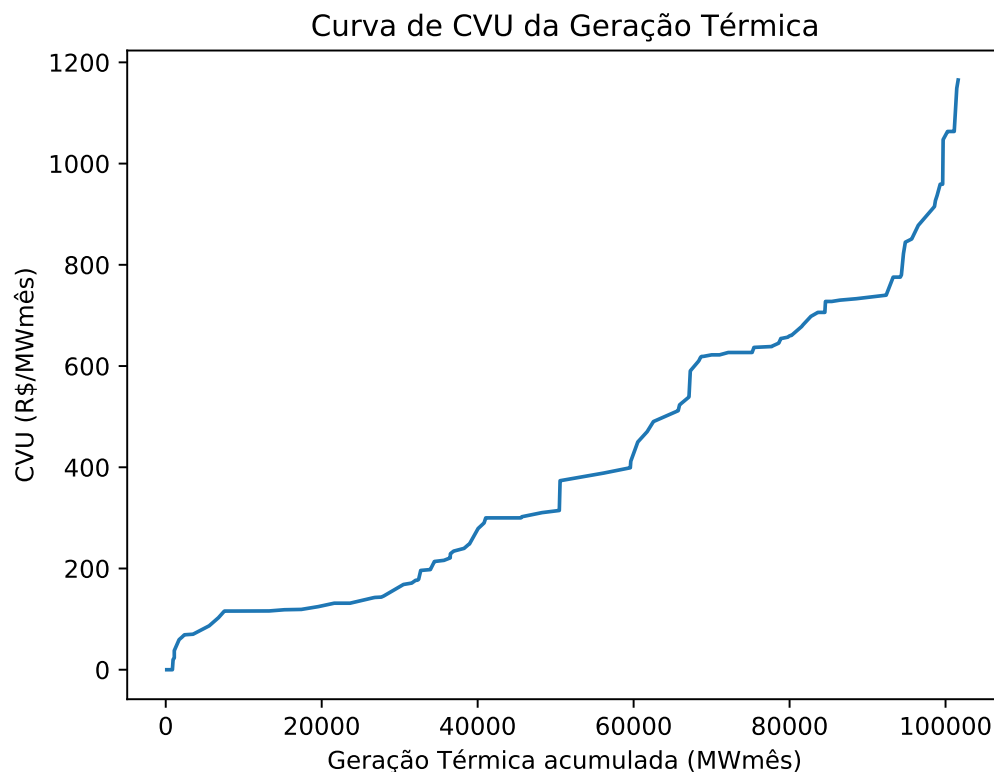


Figura 6.14: CVU da geração térmica no caso semestral

Vale notar que, como o sistema possui uma capacidade de armazenamento, a operação ideal é manter a geração térmica o mais constante possível, enquanto que a geração hidráulica compensa a variação da afluência e da demanda entre os estágios sucessivos. Naturalmente, quanto mais forte for a variação entre os estágios, e, em particular, quanto mais forte for a sazonalidade, mais difícil será acomodar estas flutuações sem alterar a geração térmica. Assim, este estudo também permite observar o quanto o armazenamento do sistema é capaz de suportar a variação sazonal da incerteza.

A seguir, à esquerda iremos sempre mostrar os gráficos do regime novembro/maio, e, à direita, o gráfico correspondente para o regime fevereiro/agosto.

Começamos observando a diferença entre as séries de afluências, entre cada um dos regimes. Os gráficos de quantis e violinos da figura 6.15 têm todos o comportamento esperado, e são representativos das amostras de cenários escolhidos para a simulação. Desta forma, podemos ver a periodicidade da incerteza, alternando os semestres do ano, que são especialmente marcadas pelas estações no caso novembro/maio. No gráfico de quantis, esta diferença se nota pelas maiores inclinações no regime novembro/maio, e também pela maior amplitude entre mínimas e máximas. Já no gráfico de violinos, observamos que, mesmo que o regime fevereiro/agosto ainda tenha uma sazonalidade clara, esta é muito menos extrema do que no outro caso: a média (e a mediana) de afluências na estação úmida é maior do que o máximo na estação seca, e, simetricamente, a média de afluências na estação seca é menor do que a menor afluência histórica na estação úmida.

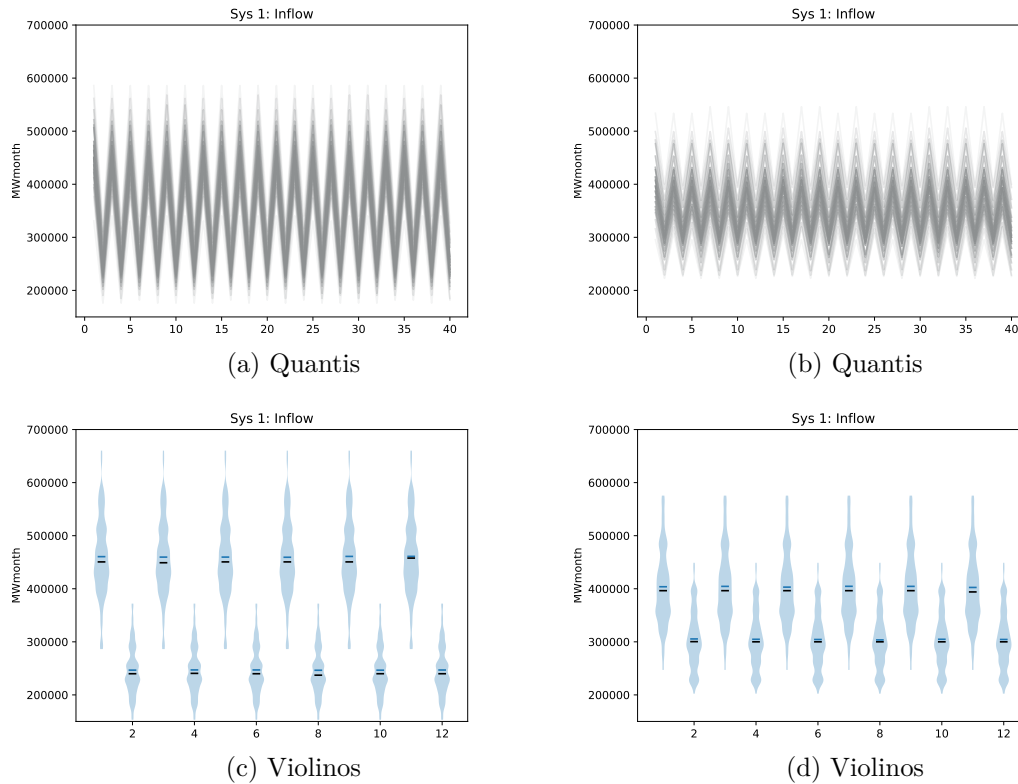


Figura 6.15: Comparação de regimes semestrais: Energia Afluente

Também notamos que as distribuições não são exatamente as mesmas, o que é natural visto que os cenários utilizados para as simulações são amostrados independentemente.

Novamente, como no caso anual, o deficit, para a maior parte dos cenários, será zero. Em ambos os regimes, os gráficos de quantis da figura 6.16 indicam que, para qualquer estágio, dentre todos os cenários amostrados, há menos de 5% de probabilidade de deficit. Além disso, os gráficos de violinos mostram que, no caso do regime fevereiro-agosto, observamos que há eventos de deficit tanto nos semestres “úmidos” como nos de “seca”, ao passo que o regime novembro-maio só apresenta deficit nos meses de seca.

Já o vertimento tem o padrão semestral inverso, como era de se esperar. Vemos na figura 6.17 que, no regime novembro-maio, temos mais vertimento nos estágios correspondentes aos períodos de cheia, e praticamente não há vertimento nos períodos húmidos, representando sempre menos do que 5% probabilidade de vertimento. O regime fevereiro-agosto apresenta menos vertimento na estação mais chuvosa, e mais vertimento na estação de menor afluência, se comparado ao regime novembro-maio, mas ainda tem menos de 5% de probabilidade de vertimento na estação seca. Já para as estações úmidas, no regime mais extremo há mais de 15% de probabilidade de vertimento, ao passo que no regime menos extremo esta passa de 10%, mas não chega a 15%.

Observamos, também nos gráficos em violinos, que há uma tendência de crescimento do vertimento ao longo dos estágios de simulação. Isso sugere que, em ambos os casos, os reservatórios estão também em média com maior energia armazenada, ao longo dos estágios da simulação.

Para a geração térmica, observamos mais uma vez na figura 6.18 que o comportamento do regime fevereiro-agosto tem oscilações com menos amplitude do que as vistas no regime novembro-maio.

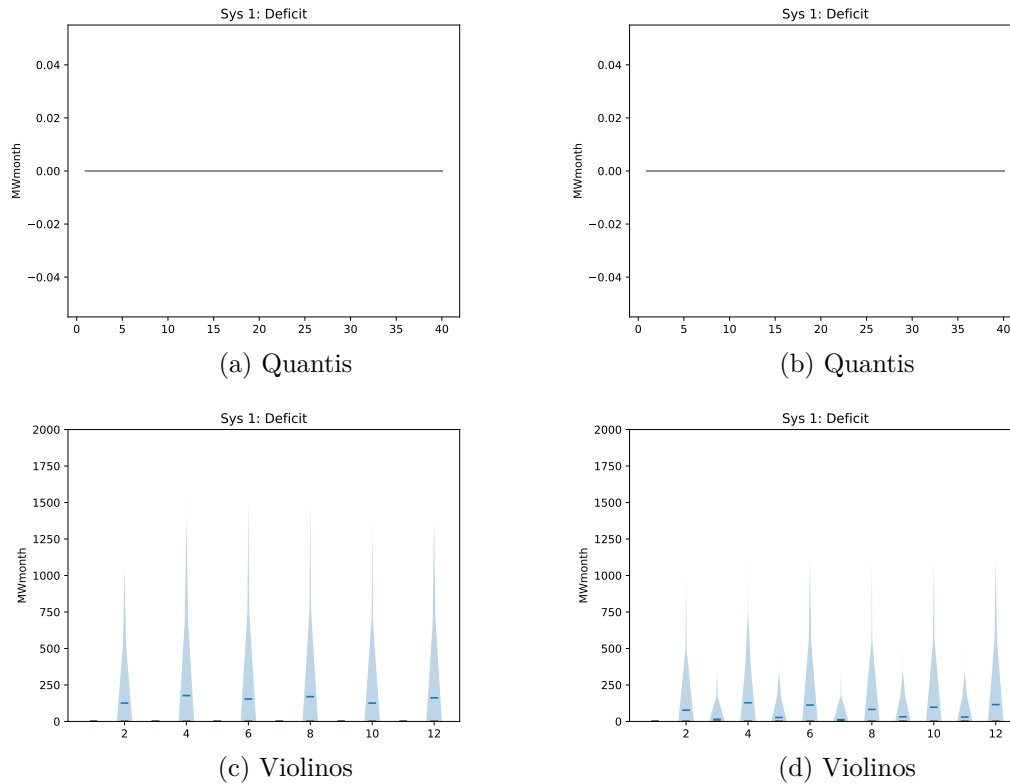


Figura 6.16: Comparação de regimes semestrais: Deficit

Também parece, para esta variável, que a convergência para o regime estacionário é mais rápida no regime fevereiro-agosto do que no regime novembro-maio, pois as medianas (traços em preto) atingem valores mais baixos bem mais rápido. Aqui, observamos um fenômeno bastante interessante: as medianas de geração térmica na estação seca do regime novembro/maio, correspondentes aos estágios pares como visto na figura 6.15, tendem a ser mais baixas do que as medianas para a estação úmida; as médias, entretanto, tendem a ser mais altas. Isso mostra que a distribuição da geração térmica é bastante assimétrica.

Vale notar que o regime esperado para a geração térmica seria de mantê-la, tanto quanto possível, constante ao longo do tempo, utilizando a geração hidráulica para compensar as flutuações. Isso se deve à convexidade dos custos de geração térmica, pois é mais barato gerar a média ao longo de um período do que ter estágios com maior e menor geração térmica. Assim sendo, vemos que o regime com uma sazonalidade menos forte tem bem menos variações sazonais na geração térmica do que o regime novembro/maio, onde podemos notar, tanto pelos quantis como pelos violinos, que há uma maior oscilação periódica de geração.

Também notamos uma particularidade da geração térmica no gráfico de quantis, onde há várias curvas de quantis que coincidem. Isto fica evidente quando duas curvas de quantis que estavam diferentes num estágio passam a ficar iguais no estágio seguinte, ou, ao contrário, se separam de um estágio anterior. Isso se deve à geração térmica ter “saltos” nos custos, o que leva a uma maior concentração de cenários para os quais a geração térmica será a mesma.

Para a geração hidráulica, apresentada na figura 6.19, notamos que em ambos os regimes a periodicidade da geração hidroelétrica é muito bem definida, o que pode ser observado tanto no gráfico de quantis como no gráfico de violinos. Dada a discussão anterior sobre a geração térmica,

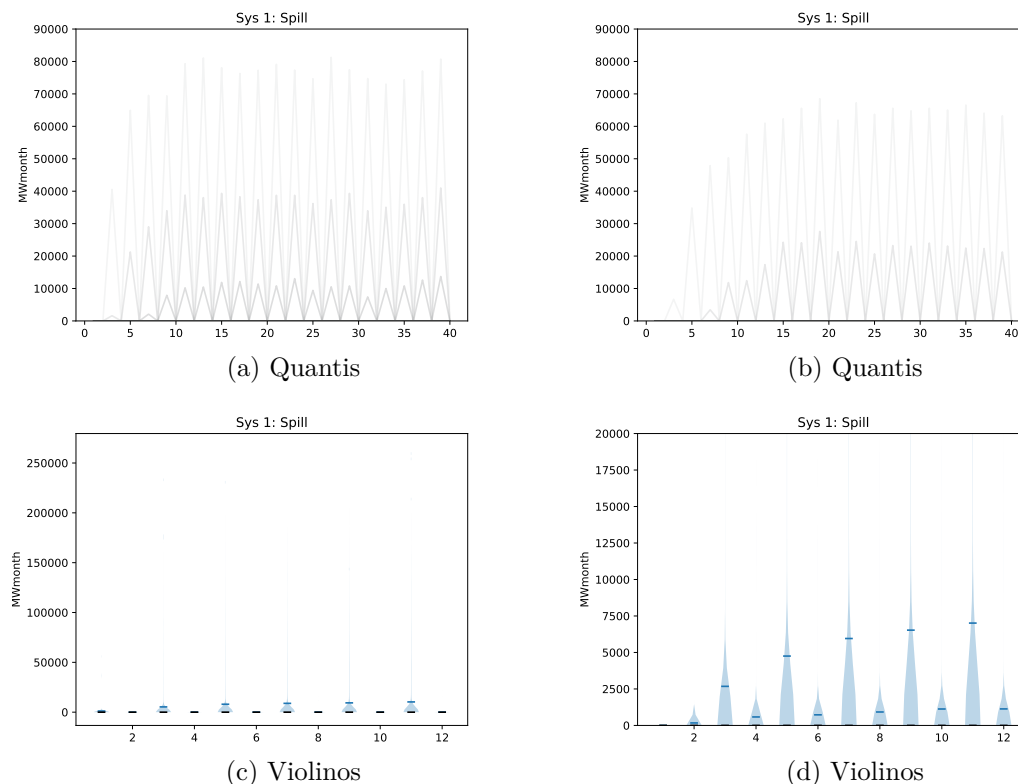


Figura 6.17: Comparação de regimes semestrais: Vertimento

este resultado é bastante natural. Vale notar que a uma parte significativa da diferença pode ser explicada pela variação da demanda, da ordem de 24 mil MWmês no caso novembro/maio, e apenas da ordem de 6 mil MWmês no caso fevereiro/agosto.

Entretanto, também é notório que a oscilação da geração hidroelétrica é maior no regime novembro/maio. Por exemplo, os primeiros quantis são menores nesse primeiro regime, e os últimos quantis são mais altos, refletindo uma amplitude maior. No gráfico de violinos essa diferença na oscilação também é bastante evidente: no regime de fevereiro/agosto as médias e medianas estão muito mais próximas do que no regime novembro/maio.

Ao concluir a análise de déficit e vertimento, por um lado, e geração térmica e hidráulica, por outro, notamos que o regime fevereiro/agosto apresenta uma sazonalidade menos marcada nas variáveis de geração do que para as variáveis de déficit e vertimento. Isto pode ser explicado porque estas últimas são marcadores de eventos extremos, que, naturalmente, são mais destacados do que os eventos “regulares” de geração térmica e hidráulica.

Na figura 6.20, vemos os gráficos de quantis e violinos para o custo de operação. Como já observado no caso anual, o custo da operação em cada estágio é praticamente determinado pela geração térmica e o montante de deficit. Como há menos de 5% de probabilidade de deficit, não vemos nos quantis o impacto do deficit; este impacto é visível nos pontos mais altos dos violinos, que às vezes atingem e ultrapassam a marca de 1.25×10^{10} reais.

Assim, muitas das observações feitas para a geração térmica também valem aqui, lembrando que os gráficos parecem distorcidos porque a função de custo para uma dada geração térmica é não-linear. Além disso, como os custos unitários ficam cada vez mais altos quanto maior a geração

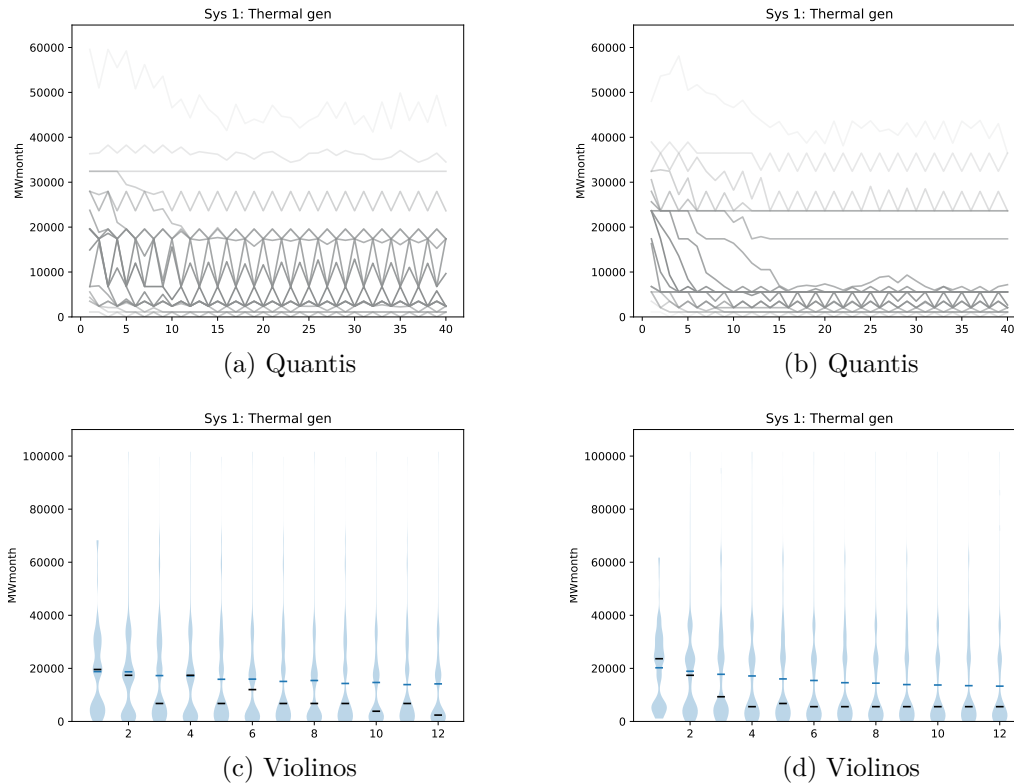


Figura 6.18: Comparação de regimes semestrais: Geração Térmica

térmica, e maiores ainda quando há deficit, os valores extremos daqueles gráficos ficam ainda mais extremos aqui. Por estas razões, por mais que o custo seja um resultado fundamental do estudo, ele não se presta tão bem para a análise do regime estacionário.

Por fim, analisamos a energia armazenada nos reservatórios ao final de cada período, como vistos nos gráficos da figura 6.21. Como é de se esperar, no caso de estudo do regime semestral novembro/maio, onde as estações estão bem marcadas, observamos também uma grande oscilação entre a energia armazenada ao final dos estágios ímpares, correspondentes ao final do período úmido, e a energia armazenada ao final do período seco em outubro. O regime defasado de 3 meses, onde cada semestre inclui porções das estações úmida e seca, apresenta oscilações muito menores.

A energia armazenada é uma das variáveis mais importantes para o estudo do regime estacionário, visto que este é o estado Markoviano do processo de decisões: todas as outras decisões são funções da incerteza e da energia armazenada. Assim, a convergência desta variável implica a convergência das outras, mesmo que em velocidades diferentes. Neste aspecto, podemos observar que ainda há uma mudança visível entre as energias armazenadas do quinto ano (semestres 9 e 10) e do sexto ano (11 e 12), o que sugere que, mesmo com este modelo onde não há restrições de transmissão, o estabelecimento de um regime estacionário leva mais do que 5 anos, o que também foi observado no estudo dos regimes anuais na seção anterior.

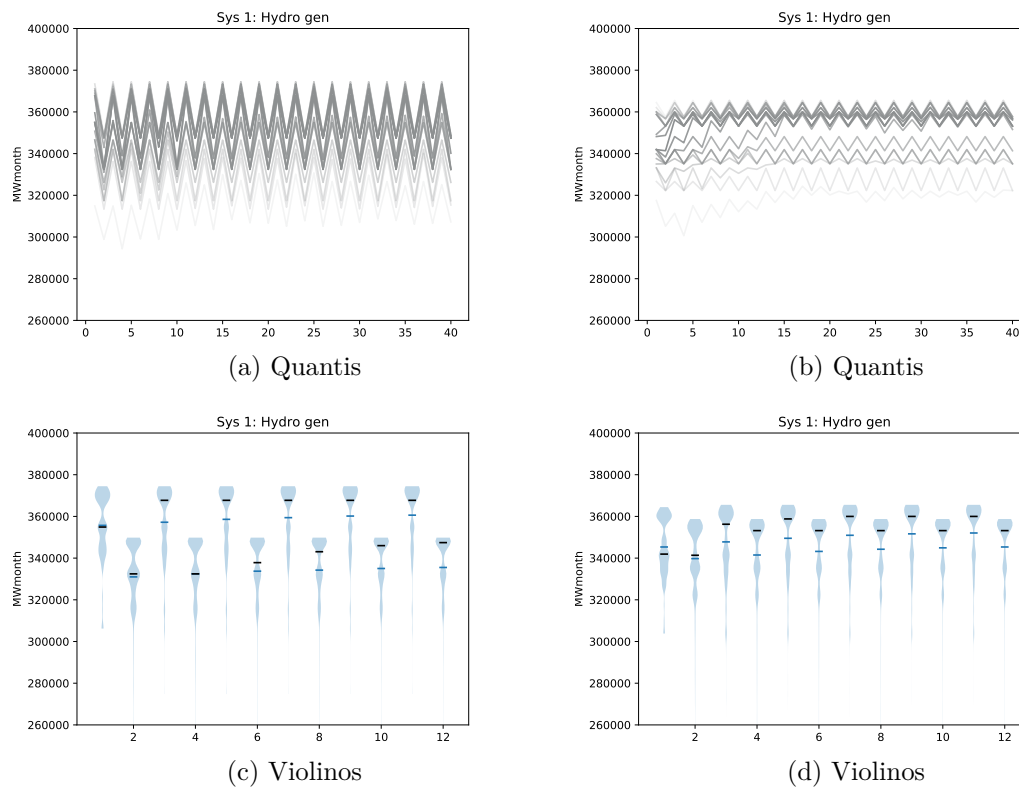


Figura 6.19: Comparação de regimes semestrais: Geração Hidráulica

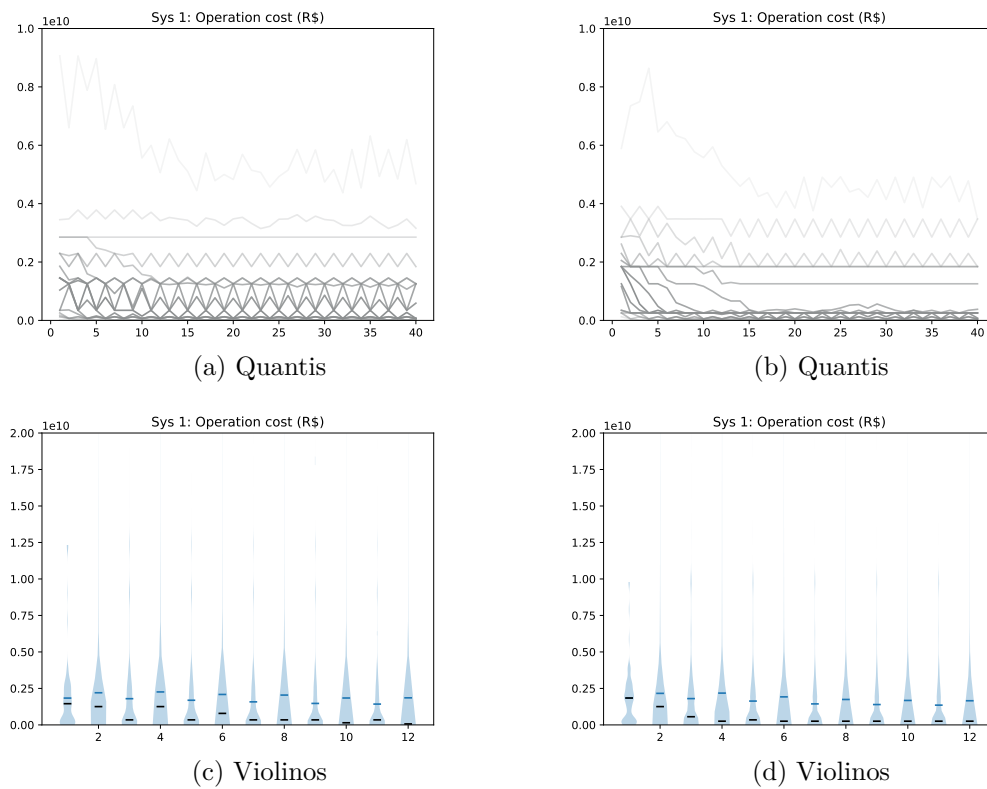


Figura 6.20: Comparação de regimes semestrais: Custo de operação

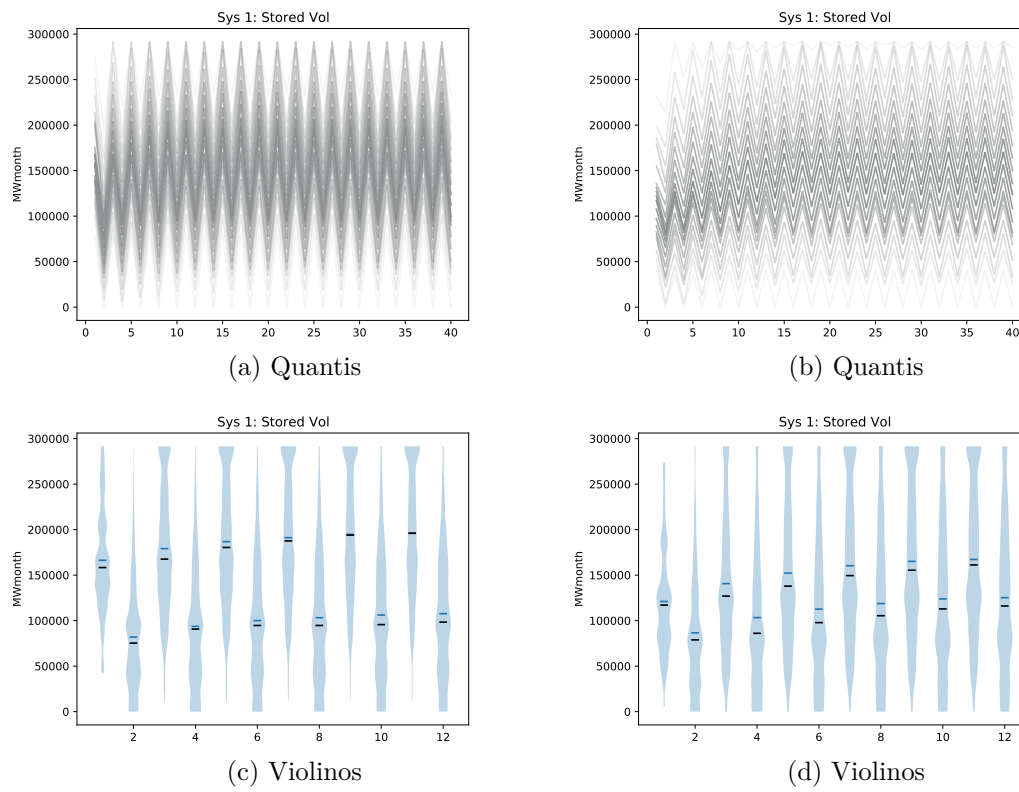


Figura 6.21: Comparação de regimes semestrais: Energia Armazenada

6.4 Casos de estágios mensais

Finalmente, vejamos os resultados para dois casos de estágios mensais, já mais próximos do modelo típico de planejamento da operação. O primeiro modelo será periódico, de período 12, enquanto que o segundo será um modelo com 120 estágios, representando o horizonte de 10 anos comumente utilizado para os estudos de médio prazo.

A tabela 5 apresenta as principais características dos modelos utilizados, agora para estes casos mensais. Para não ficar extremamente repetitivo, incluímos os valores da demanda e afluência para quatro dos 12 meses que constituem um período. A grande mudança que podemos ver neste modelo é que a energia armazenada máxima passa a ser maior do que todos os outros parâmetros de um estágio. Isto reforça o papel fundamental que o armazenamento terá para a solução. Ademais, vemos que a flutuação da demanda média é relativamente modesta entre estágios, estando na ordem de 10%, enquanto que a oscilação de energia afluente é significativa: a média de afluências em fevereiro é quase três vezes maior do que a de agosto.

Parâmetro	Valores (MWmês)			
Energia Armazenada máxima	291.579,1			
Geração Hidráulica máxima	78.785,2			
Geração Térmica máxima	23.323,0			
Geração Térmica mínima	6.386,6			
	Fevereiro	Mai	Agosto	Novembro
Demanda de energia	71.469	65.761	64.210	66.252
Energia Afluente média	95.458,3	56.242,2	33.628,5	44.927,5
Energia Afluente máxima	154.780,4	93.753,3	69.749,3	79.687,5
Energia Afluente mínima	42.287,0	36.000,6	19.708,4	21.761,0

Tabela 5: Principais características dos modelos mensais

Como nos outros casos, mostramos na figura 6.22 a curva de custo variável unitário (CVU).

A seguir, analogamente ao caso semestral, apresentamos lado a lado os gráficos de quantis e violinos para as diferentes variáveis de decisão do problema de planejamento. Os violinos, novamente por questão de visualização, contém apenas os 12 primeiros meses da simulação. À esquerda, teremos os gráficos correspondentes aos resultados obtidos através do modelo periódico, enquanto que, à direita, temos os gráficos correspondentes ao modelo de 120 estágios. Observamos que o primeiro mês de planejamento, e consequentemente apresentado na simulação, é Janeiro.

Os primeiros gráficos, como de hábito, serão os das afluências, que vemos na figura 6.23. Aqui, vemos de forma bastante clara a sazonalidade das mesmas, com grandes afluências no trimestre Janeiro-Fevereiro-Março, também sujeitas à uma grande variância, enquanto os meses de junho a novembro têm afluências bem menores, e com correspondente menor variação.

Em seguida, temos os gráficos de déficit, que continuam atestando menos de 5% de probabilidade de déficit em todos os cenários, o que continua compatível com os resultados anteriores, e também é esperado. Ademais, vemos na figura 6.25, que o vertimento também está bastante reduzido, em ambos os casos, se comparado com os resultados de vertimento seja para o regime anual, seja para o regime semestral. Isto é fácil de explicar, porque agora temos mais oportunidades para alterar a operação do sistema ao longo de uma série, e com isto melhor aproveitar o armazenamento no caso de grande afluência, suavizando a operação.

Entretanto, esta simulação é, em certo sentido, mais favorável do que o observado nas séries históricas, já que sorteamos, de forma independente, as vazões de cada mês. Ora, em geral há

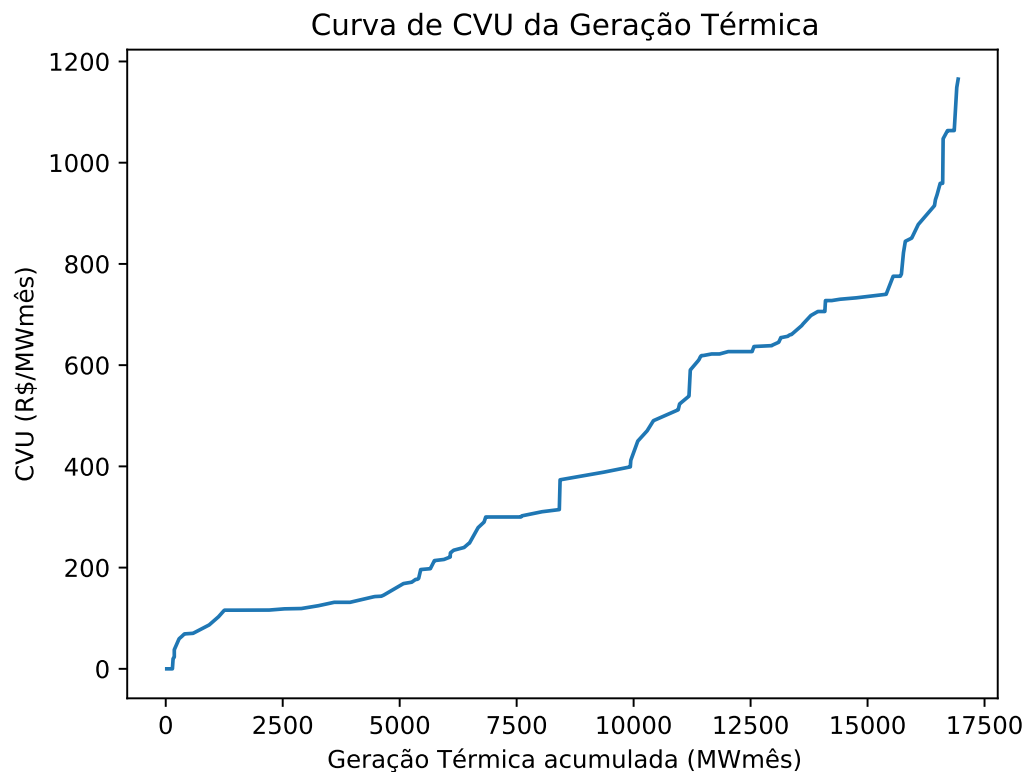


Figura 6.22: CVU da geração térmica no caso mensal

correlação entre as vazões de meses próximos, com anos de cheias e anos de seca mais pronunciadas, o que não é reproduzido pela geração de séries independentes. Isto sugere que, pelo menos para o caso de estágios mensais, seria importante utilizar um modelo mais acurado de geração de séries para a energia afluyente, de forma a obter séries mais extremas do que o modelo de independência. Isto exige, por sua vez, alterar o modelo de PDDE, já que o mesmo exige que a incerteza de cada estágio seja uma variável aleatória *independente* das incertezas observadas nos estágios anteriores. Uma forma de incorporar previsões mais acuradas de energia afluyente é pelo uso de técnicas de séries temporais, com o uso de modelos autorregressivos. Esta generalização da PDDE será explorada no próximo relatório, pois ela necessita uma alteração tanto do modelo de otimização utilizado como a própria reformulação das incertezas que serão associadas a cada estágio.

Com isto, passemos aos gráficos de geração térmica, que podemos ver na figura 6.26. Aqui, podemos observar as primeiras diferenças significativas entre os modelos periódico (à esquerda) e o de 120 estágios (à direita). Em primeiro lugar, notamos que o modelo periódico também resulta em um comportamento periódico no gráfico de quantis, onde é visível uma estrutura de repetição a cada 12 meses, para diversos quantis. Por outro lado, o gráfico de quantis do modelo de 120 estágios não possui a mesma regularidade, tendo níveis de geração térmica diferentes a cada estágio. A explicação deste fenômeno é relativamente direta: o modelo periódico possui apenas 12 funções de custo futuro, a partir das quais irá calcular a operação; já o modelo de 120 estágios possui uma função de custo futuro, potencialmente diferente, para cada estágio. Assim, mesmo que estas sejam similares entre meses correspondentes de cada ano, as diferenças que estas tenham serão visíveis tanto como pequenas variações das decisões, como, de forma mais notável, pela diminuição da

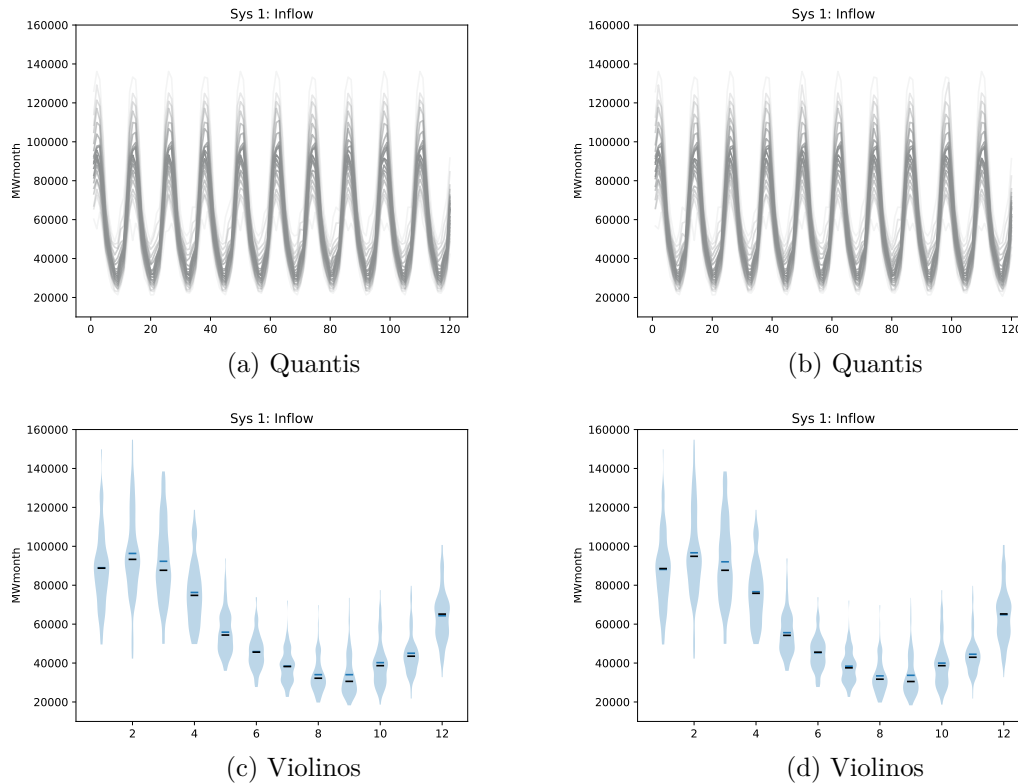


Figura 6.23: Mensal periódico e mensal de 120 estágios: Energia Afluente

periodicidade, aqui notada nos quantis da geração térmica.

Um outro reflexo deste mesmo fenômeno é a surpreendente estabilidade da geração térmica ao longo do primeiro ano, no modelo periódico. Podemos observar, no gráfico de violinos, que tanto a média como a mediana estão muito mais estáveis do que no caso de 120 estágios. Esta maior constância ao longo dos estágios também pode ser observada, diretamente, na forma que tomam os violinos: no modelo de 120 estágios temos a presença de “bolsas” acima de 2000 MWmês, em alturas variadas, o que praticamente não ocorre no modelo periódico. Por fim observamos que, com menos variação da geração térmica, também teremos uma maior estabilidade de custos marginais ao longo do ano, já que estes são, essencialmente, determinados pela usina térmica mais cara despachada.

A segunda diferença, mais esperada, entre os modelos, é o efeito de final de horizonte. Enquanto o modelo periódico não apresenta nenhuma variação visível, o modelo de 120 estágios tem, nos últimos estágios, dois comportamentos que são reflexos do fim de horizonte: o primeiro, uma redução em média da geração térmica, correspondente ao esvaziamento, também em média, dos reservatórios, já que o valor da água fica menor nos últimos estágios. O segundo, em certo sentido oposto, é o aumento do quantil de 95% da geração térmica; ele se explica porque, dado um esvaziamento maior dos reservatórios, a flexibilidade de operação fica reduzida, e uma série mais negativa, que tem pouca probabilidade de ocorrer, torna necessária uma maior geração térmica. Assim, ao se aproximar do final do horizonte, não só a geração térmica tende a ficar menor, em média, como também fica mais dispersa, refletindo uma operação com menos controle do que o normal durante os períodos mais distantes do final do horizonte.

A geração hidráulica, apresentada na figura 6.27, é, naturalmente, complementar à geração térmica, visto que as séries de demanda e de afluência são iguais para os dois casos estudados.

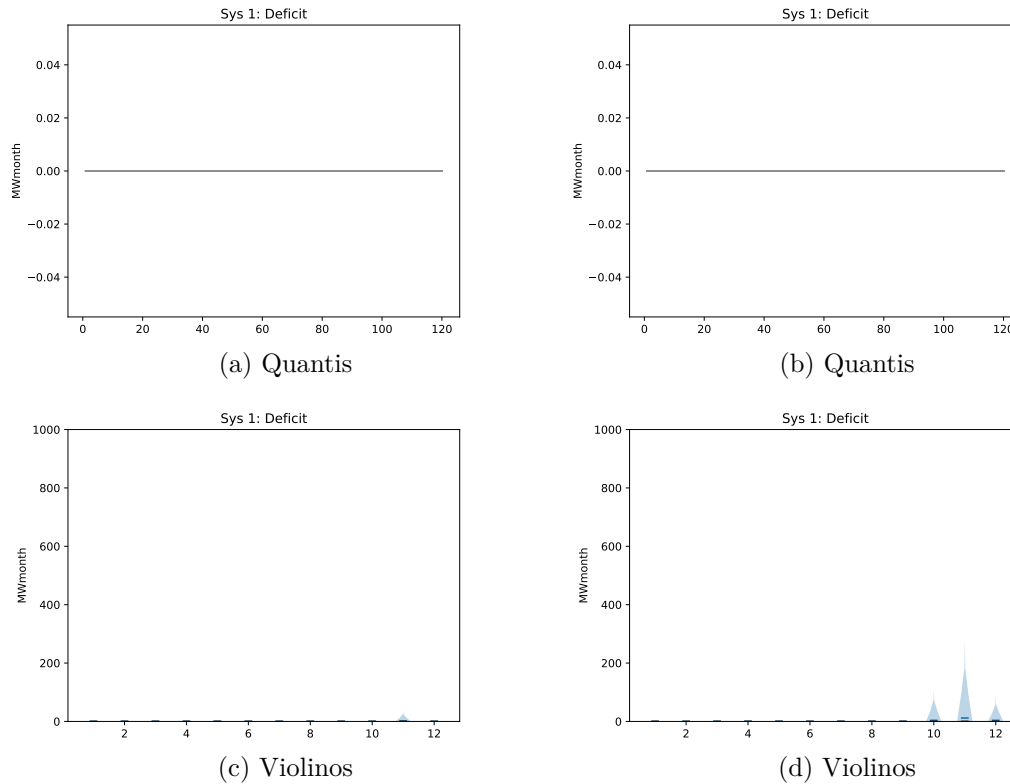


Figura 6.24: Mensal periódico e mensal de 120 estágios: Deficit

Entretanto, devido à escala da geração hidráulica, que é aproximadamente dez vezes maior do que a da geração térmica, a maior parte dos efeitos observados anteriormente se revelam muito mais difíceis de notar.

O que podemos observar destas duas últimas figuras é que as simulações parecem convergir bem mais rápido para o regime estacionário do que nos casos semestral e anual. Com efeito, é praticamente imperceptível a diferença entre gerações hídricas do primeiro para o segundo ano, as maiores diferenças concentrando-se nos quantis mais baixos. A partir daí, é mais difícil ainda afirmar que houve uma mudança entre o segundo e terceiro anos, a partir destes gráficos.

Enfim, passamos aos gráficos de energia armazenada, que vemos na figura 6.29. Nesta, podemos notar que há uma leve tendência de aumento da energia armazenada ao longo dos primeiros anos, especialmente até o final da terceira estação úmida. Também notamos que há uma grande excursão ao longo de cada ano, com quantis chegando próximo do esvaziamento, ao final da estação seca, e com quantis próximos ao máximo de armazenamento, ao final da estação úmida.

De forma geral, ambos os modelos, seja o periódico, seja o de 120 estágios, apresentam comportamentos muito similares para as simulações de energia armazenada. Apenas ao final do período é possível ver uma queda, um pouco mais acentuada, da energia armazenada no modelo de 120 estágios, correspondente ao menor valor da água para estes estágios finais.

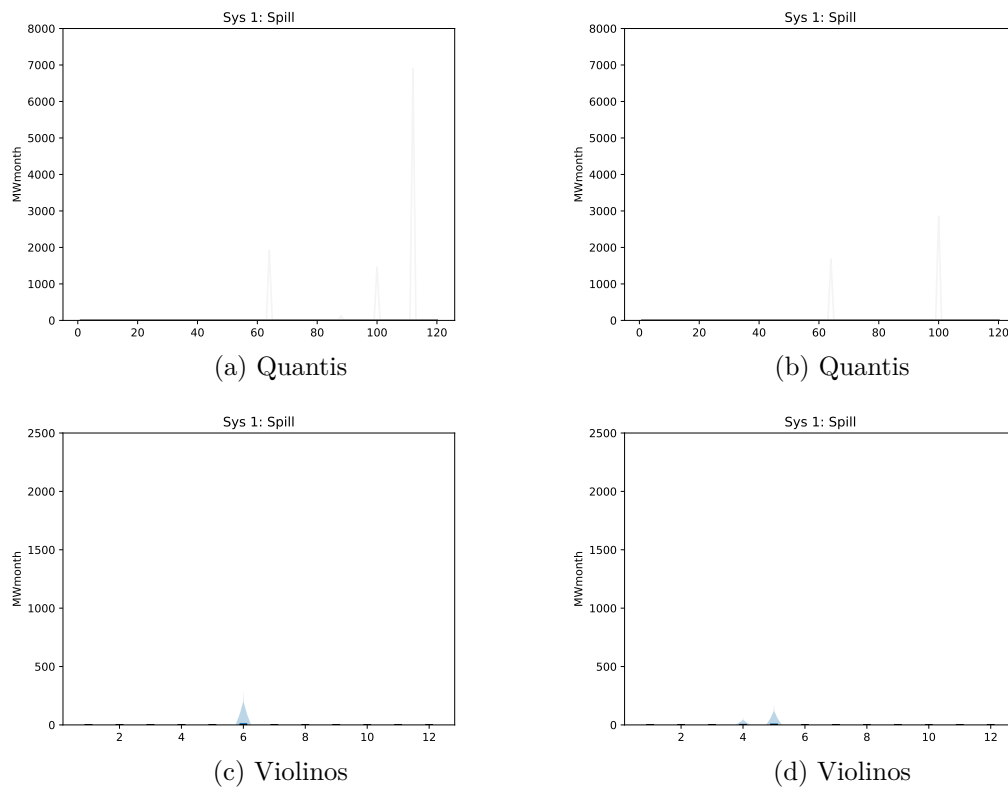


Figura 6.25: Mensal periódico e mensal de 120 estágios: Vertimento

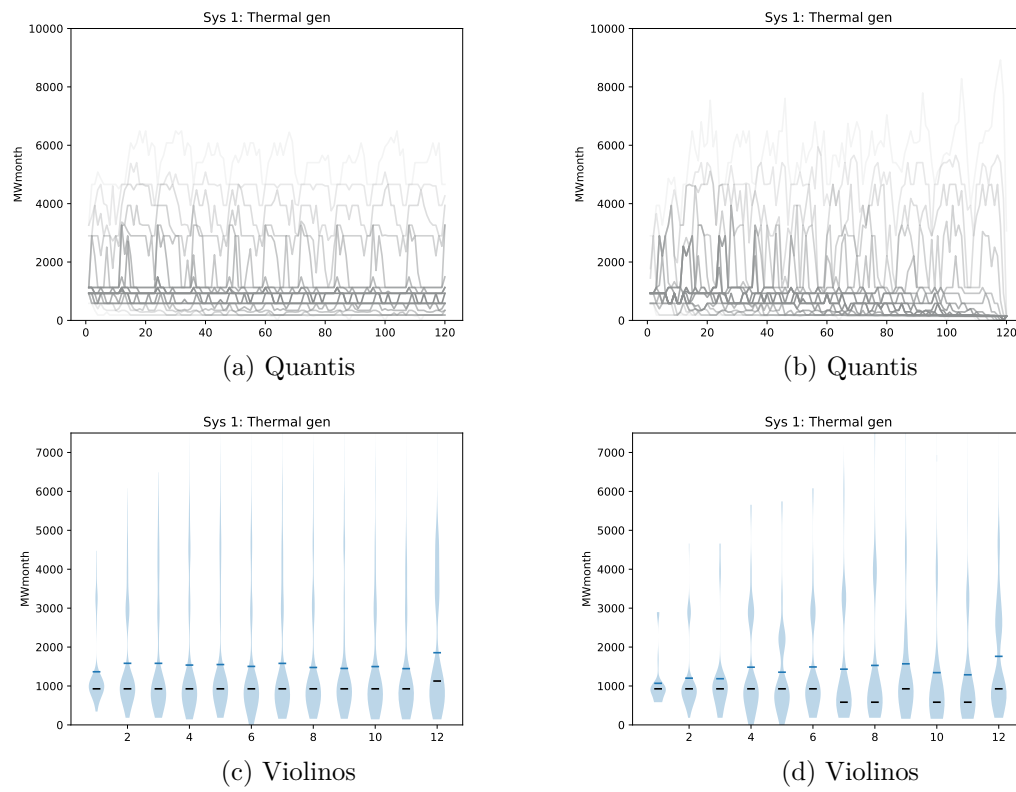


Figura 6.26: Mensal periódico e mensal de 120 estágios: Geração Térmica

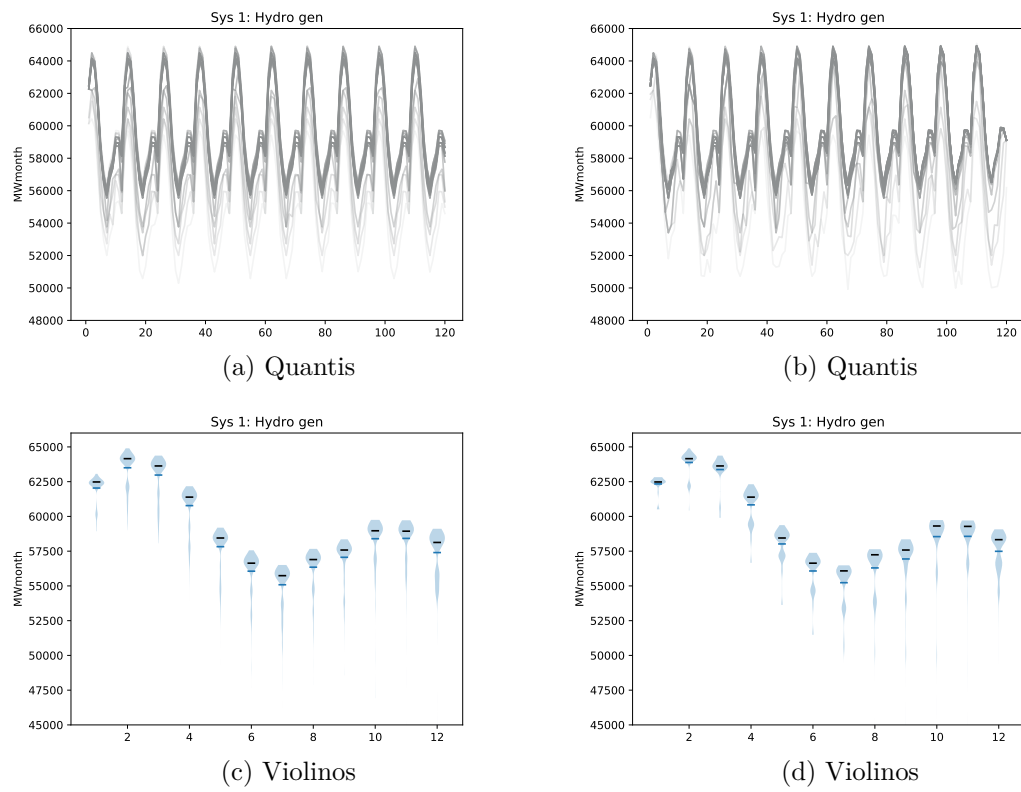


Figura 6.27: Mensal periódico e mensal de 120 estágios: Geração Hidráulica

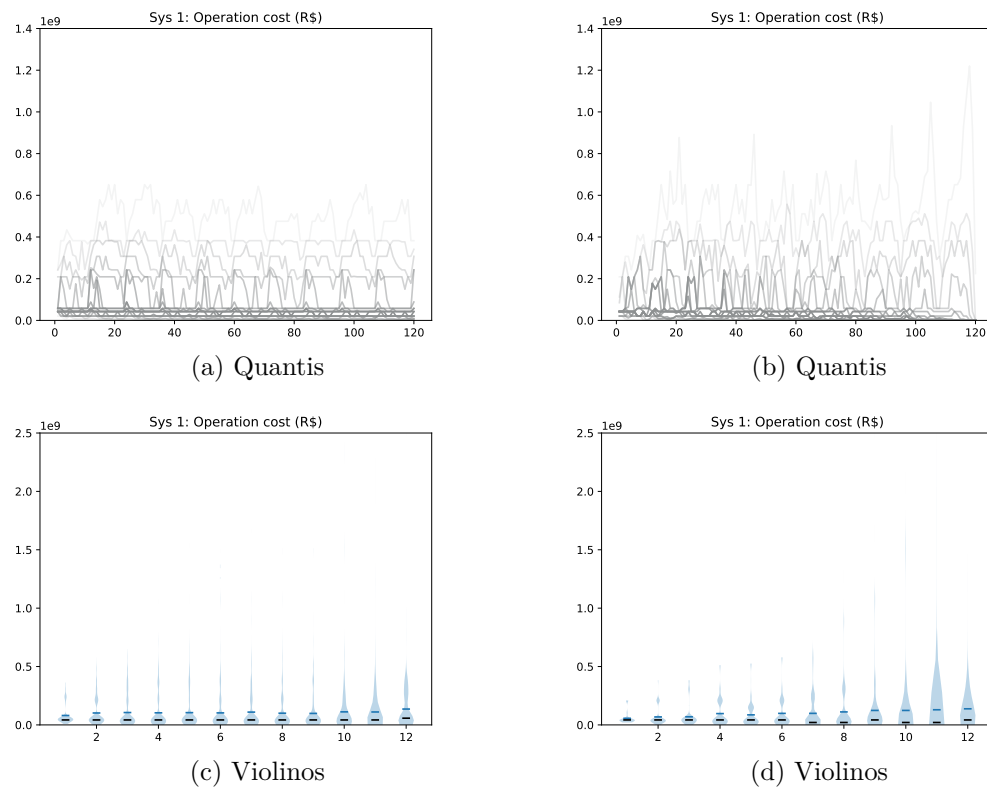


Figura 6.28: Mensal periódico e mensal de 120 estágios: Custo de operação

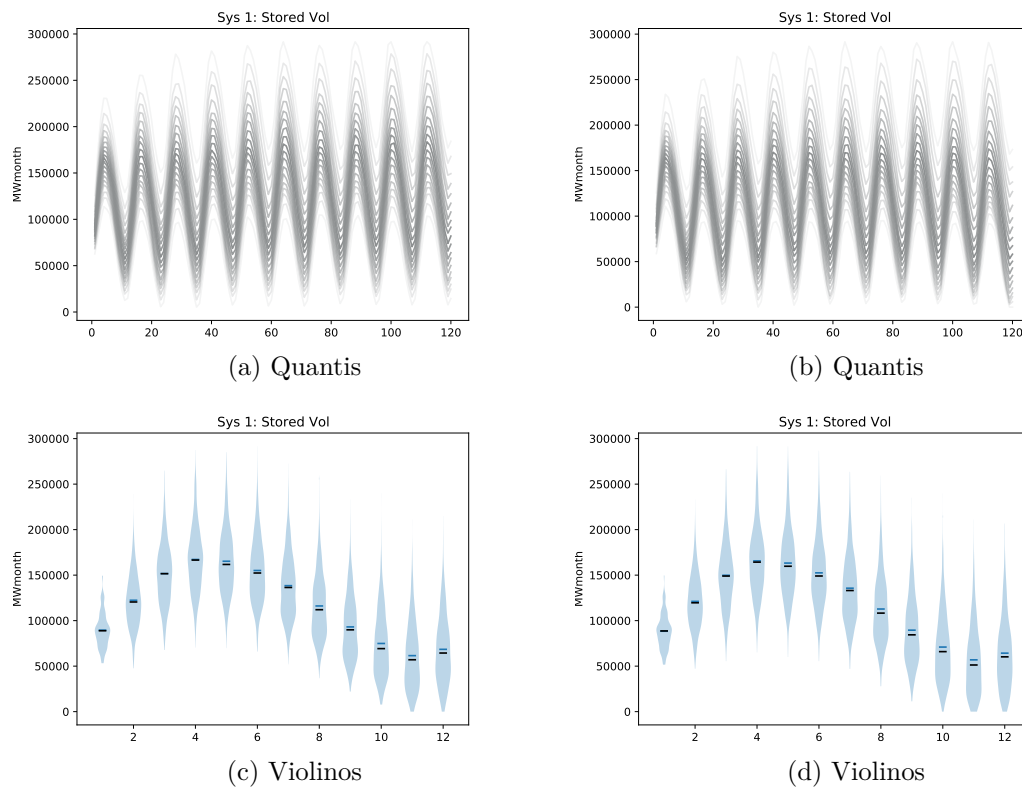


Figura 6.29: Mensal periódico e mensal de 120 estágios: Energia Armazenada

7 Resultados

Nesta seção, continuamos apresentando os resultados da aplicação da metodologia de horizonte infinito em problemas de otimização estocástica para o planejamento da operação. Primeiro, iremos descrever os casos selecionados, e, a seguir, apresentar as análises em cada grupo.

7.1 Casos de estudo

Os problemas que estudamos são referentes ao planejamento da operação energética de médio prazo. Com relação aos já vistos na seção 6, os casos serão mais realistas e, com isto, mais complexos. Assim, se antes consideramos um único mercado — ignorando restrições de transmissão — e um único Reservatório Equivalente de Energia (REE) — ignorando as restrições individuais de capacidade de armazenamento e geração —, aqui iremos tratar de problemas com quatro variáveis de estado, correspondentes aos quatro REEs tradicionais: Sudeste, Sul, Nordeste e Norte. Esta separação também nos permite tratar dos 4 submercados, associados às regiões de cada Reservatório Equivalente, e com isto incorporando também as restrições de transmissão entre as regiões.

Nossas análises se dividem em dois grupos:

1. Comparação do modelo periódico com o não periódico: usaremos um problema neutro a risco, com estágios mensais, e comparamos um problema formulado com horizonte periódico e um problema formulado com 120 estágios. Aqui, o objetivo é avaliar a convergência de ambos os modelos para um regime estacionário, e estudar os efeitos de bordo presentes no caso do modelo de horizonte finito.
2. Modelos com aversão ao risco. O objetivo principal destes casos será averiguar como a inclusão de aversão a risco, através da mudança da função objetivo, substituindo a média dos custos pelo *CVaR aninhado*, irá afetar o estado estacionário e a política de operação do sistema.
3. Modelos periódicos com diferentes horizontes de cálculo da etapa *forward*, para observar o impacto do *comprimento* dos cenários *forward* gerados na qualidade da solução e na aproximação das funções de custo futuro.

7.2 Comparação periódico e não-periódico

Já vimos na seção 6 que o modelo periódico era mais consistente em sua operação ao longo dos meses, tendo uma maior estabilidade do despacho térmico, apesar de os regimes de longo prazo não serem marcadamente diferentes.

A tabela 6 recapitula, por conveniência, as principais características dos modelos utilizados, agora para estes casos mensais. Por simplicidade, apresentamos os valores *agregados nos subsistemas*, para dois meses bastante diferentes: Fevereiro (típico da estação quente e úmida), e Agosto (típico da estação amena e seca). Notamos que a demanda não é significativamente diferente entre os meses, flutuando na ordem de 10%, enquanto que a afluência possui uma variação sazonal bastante significativa, o que reforça o papel fundamental que o armazenamento terá para a solução.

Veremos nesta seção como o sistema se comporta ao levar em conta

1. os limites de armazenamento independentes dos reservatórios;
2. os limites de transmissão entre subsistemas.

Parâmetro	Valores (MWmês)	
Energia Armazenada máxima	291.579,1	
Geração Hidráulica máxima	78.785,2	
Geração Térmica máxima	23.323,0	
Geração Térmica mínima	6.386,6	
	Fevereiro	Agosto
Demanda de energia	71.469	64.210
Energia Afluente média	95.458,3	33.628,5
Energia Afluente máxima	154.780,4	69.749,3
Energia Afluente mínima	42.287,0	19.708,4

Tabela 6: Principais características dos modelos tratados

Começamos, portanto, com as curvas de armazenamento ao longo do tempo, na figura 7.1.

A primeira observação é que ambos os modelos, tanto periódico como não-periódico, têm uma tendência para armazenar mais a partir da condição inicial dada. Isto pode ser observado de forma muito clara no gráfico com as médias anuais da figura 7.1b, que reduzem a variação sazonal do armazenamento, e permitem observar a tendência do sistema. Também podemos notar que o modelo periódico tem uma *maior* tendência ao armazenamento, que desde o *primeiro ano* já se mostra visível na operação.

A outra grande diferença, também esperada, se refere aos efeitos ditos “de bordo”, especialmente na comparação entre o modelo periódico com o não-periódico. De fato, se podemos observar em ambos os modelos uma tendência de armazenamento ao longo dos seis primeiros anos, o modelo não-periódico, por possuir um horizonte finito, também já começa a esvaziar os reservatórios nos últimos anos do período planejado. Aqui, valem duas observações. Em primeiro lugar, o modelo periódico, após uma subida inicial rápida, continua na direção de maior armazenamento bem além do quinto ano, o que sugere que este levará mais do que este tempo para atingir de fato um *regime estacionário*. Em seguida, no modelo de 120 estágios, mesmo apresentando um certo “patamar” nos anos 5, 6 e 7, a curva de energia armazenada tem a forma de uma parábola, que é consistente com não atingir o estado estacionário. Mais ainda, ao levar em consideração os resultados do modelo periódico, que visivelmente não está estável ao final do 7 ano, é bem razoável imaginar que o modelo de horizonte finito também não atingiu, ainda, o regime estacionário.

Vale notar que é comum desprezar alguns anos, tipicamente os últimos cinco, da análise de operação, já que a decisão implícita ali não será tomada, e que estes últimos anos estão representados no modelo apenas para afastar o final do horizonte da operação, e com isto reduzir o impacto deste na decisão presente. Este nosso estudo sugere que este período de cinco anos pode ainda estar muito curto, já que um eventual “regime estacionário”, que seria aquele obtido pelo modelo de horizonte infinito, parece demorar mais para se estabelecer. Ademais, também vemos que a política de armazenamento indicada pelo modelo de horizonte finito é diferente daquela obtida pelo horizonte periódico, apontando para uma menor necessidade de armazenamento.

Vejamos como se comporta a geração térmica para ambos estes regimes, como ilustrado na figura 7.2.

Novamente, os dois modelos dão resultados muito semelhantes ao longo dos 5 primeiros anos, mas uma análise um pouco mais cuidadosa revela diferenças interessantes. A primeira, talvez mais significativa, e que não era tão notável no caso da energia armazenada, é quanto à geração térmica ao longo do primeiro ano, como vemos na figura 7.2c. De fato, podemos notar que o modelo

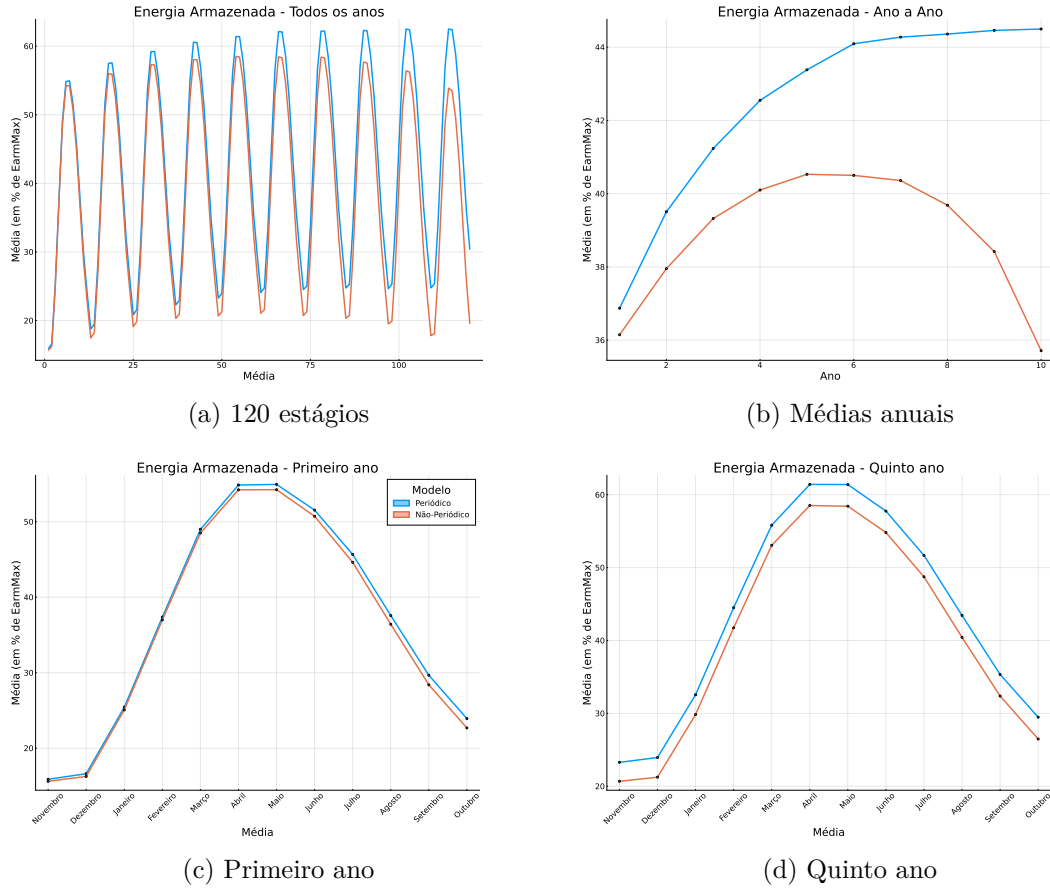


Figura 7.1: Comparativo Energia Armazenada

periódico será muito mais conservador, com um despacho térmico consistentemente maior ao longo de todos os meses do primeiro ano - salvo o último. A diferença de escala entre a geração térmica e a energia armazenada, praticamente de uma ordem de grandeza, faz que agora possamos notar melhor o impacto do modelo periódico na decisão.

Em seguida, notamos que o modelo periódico, de acordo com o que esperávamos dos resultados para a energia armazenada, irá ter um maior despacho térmico, ao longo de todos os anos da operação, vide a figura 7.2b. Isso é bastante claro no primeiro ano, em que há um grande esforço de armazenamento, e também ao longo dos outros anos iniciais, mesmo que a diferença seja menor. Naturalmente, nos últimos anos o modelo de horizonte finito será bastante mais relaxado, com geração térmica significativamente menor, mas esta decisão não será efetivamente implementada.

Enfim, notamos que, após algum tempo de operação, com o sistema já com maior energia armazenada e estabilizado, o modelo periódico será mais capaz de distribuir a geração térmica ao longo dos meses do ano. Assim, notamos na figura 7.2d que a geração térmica varia muito menos no caso periódico do que no caso com 120 estágios. Podemos interpretar este resultado como uma maior dificuldade de o modelo não-periódico capturar corretamente a flutuação dos reservatórios, e que, ao transicionar das estações úmidas para as secas, ele se deixa de ser “confiante” e passa a ser “conservador”, cruzando a geração térmica entre Maio e Junho.

Estas diferenças se refletem no custo de operação, onde podemos comparar as políticas empregadas de acordo com o próprio critério que estabelecemos. Na figura 7.3, apresentamos os gráficos

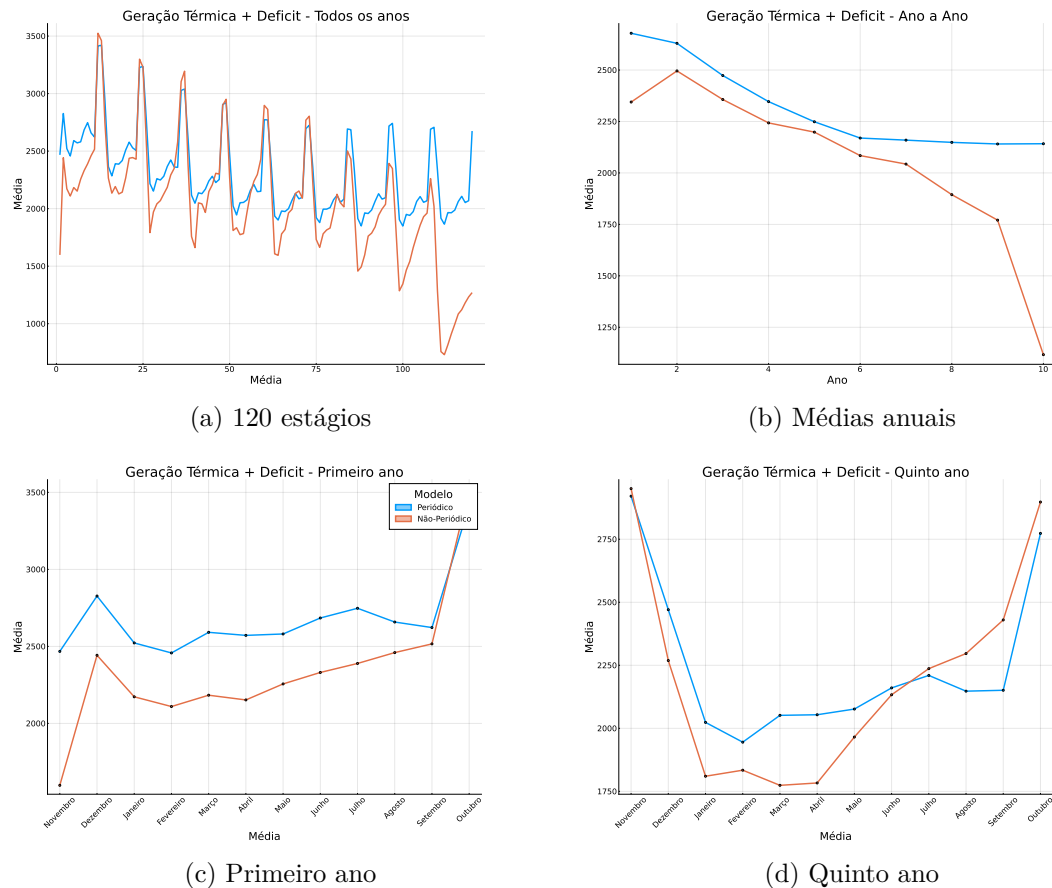


Figura 7.2: Comparativo Geração Térmica

análogos aos que já vimos anteriormente.

O mais notável é que a *política* dada pelo modelo de horizonte periódico tem uma performance muito melhor. Em particular, exceto no primeiro e no último ano, o modelo periódico é mais econômico do que o modelo de horizonte finito. Este resultado é devido a uma melhor distribuição da geração térmica, como vimos nas figuras anteriores, e que resulta em menores variações de custo ao longo dos anos, que vemos nas figuras 7.3c e 7.3d. Aqui, a *não-linearidade* dos custos de geração térmica é evidente, e muito melhor explorada pelo modelo periódico: ao reduzir os picos de geração térmica, o sistema irá incorrer menos vezes em custos elevados, enquanto que um acréscimo da geração térmica nos meses mais úmidos terá um custo total menor. Assim, vemos que, mesmo usando uma *maior* geração térmica ao longo do ano, o modelo de horizonte periódico terá um menor custo total.

Resta analisar os casos dos “anos extremos”, o primeiro e o último. Ambos são relativamente simples de explicar: o primeiro ano é onde haverá o maior esforço para aumento da energia armazenada, e daí o modelo periódico está ainda muito perto da condição inicial, e ainda não terá a oportunidade de compensar geração térmica entre meses, como já acontecerá a partir do segundo ano. Já o último ano é, novamente, um efeito de bordo: o modelo de horizonte finito, por não enxergar valor para o armazenamento no final do horizonte, irá realizar muito menos geração térmica ao esvaziar os reservatórios, e com isto efetuará uma operação bem mais barata. Naturalmente, este último ano “mais barato” nunca será de fato computado a favor do modelo de horizonte infinito, já que a política realmente implementada neste ponto do tempo será recalculada com um horizonte

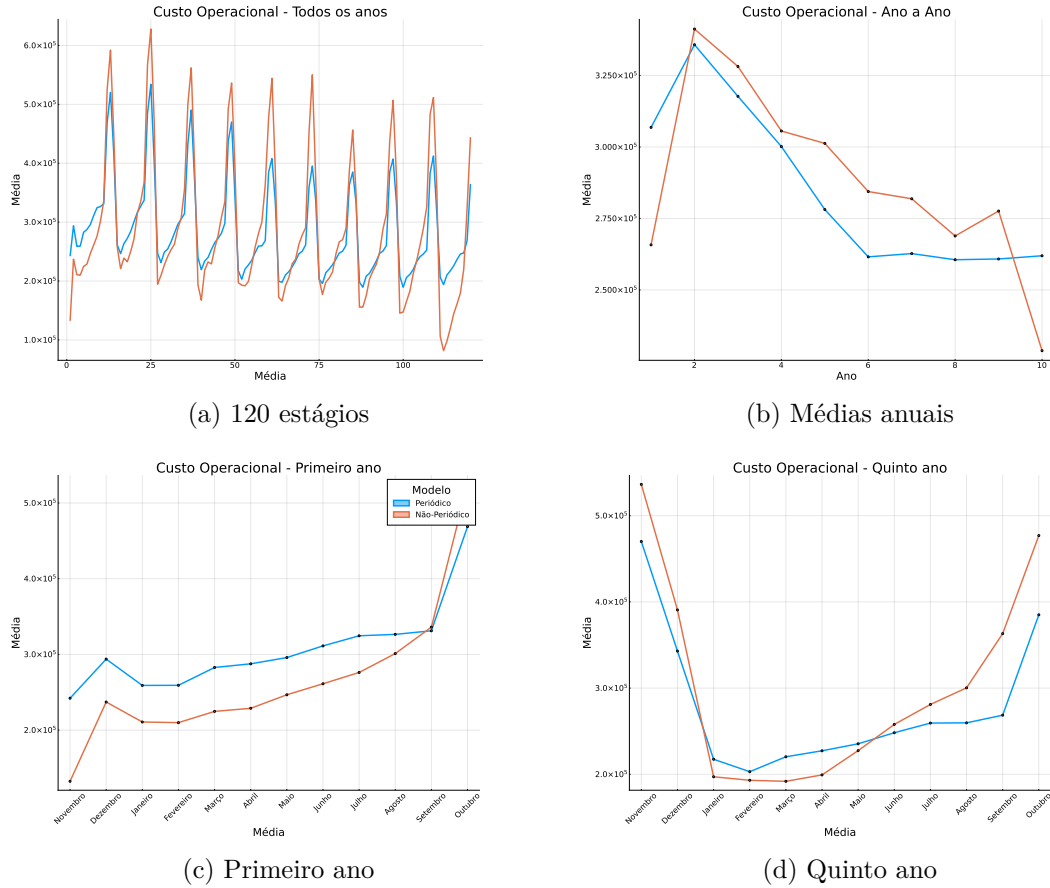


Figura 7.3: Comparativo Custo de operação

maior.

Finalmente, a figura 7.4 mostra a ocorrência de deficit para ambos os modelos.

Mais uma vez, o modelo periódico se destaca em relação ao modelo de horizonte finito. Além de uma redução nítida do déficit — já esperada, a este momento, pelo maior armazenamento nos reservatórios e pela melhor distribuição da geração térmica — notamos que o modelo periódico *sinaliza* a ocorrência de déficit com mais antecedência do que o modelo de horizonte finito. Assim, vemos na figura 7.4c que o déficit deixa de ser nulo no mês de Junho, para o modelo periódico, enquanto o modelo de horizonte finito irá demorar até Setembro para ter deficit. Poder-se-ia esperar que esta decisão fosse baseada em uma menor quantidade de deficit global, mas mesmo no primeiro ano, ao *aguardar* tanto para indicar o deficit, o modelo de horizonte finito acaba caindo numa armadilha e, nos últimos dois meses do período de planejamento, terá tanto corte de carga como o modelo periódico, que o distribuiu ao longo dos meses.

Daí em diante, é uma vitória incontestável do modelo periódico: ele irá incorrer em muito menos deficit ao longo do período de planejamento, com uma sinalização antecipada, menos corte de carga total e, principalmente, menos corte de carga nos meses mais críticos.

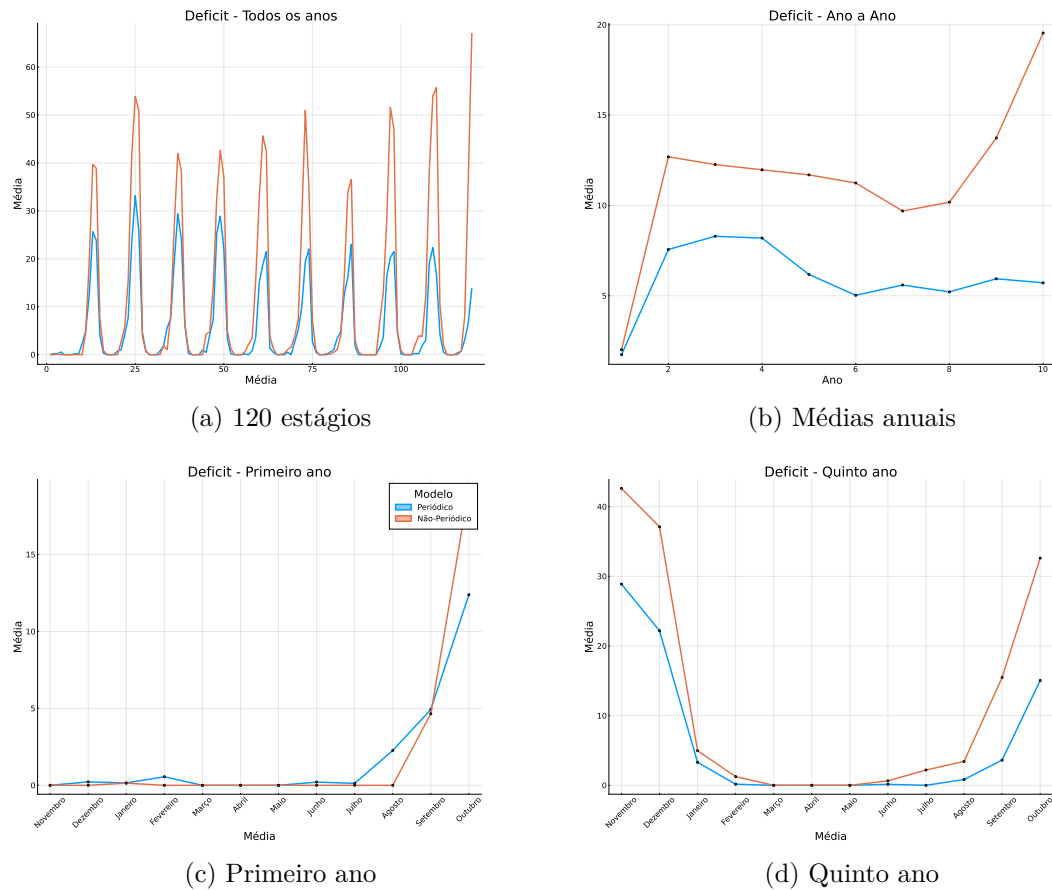


Figura 7.4: Comparativo Deficit

7.3 Comparação com aversão ao risco

Veremos, nesta seção, como a inclusão da aversão ao risco altera a solução dos modelos periódicos. Ao incluir uma medida de risco na função objetivo, estamos privilegiando soluções mais caras, em média, contanto que tenham menos variabilidade dos custos. Assim, também esperamos:

- reduzir os picos de geração térmica e deficit, responsáveis pelos picos de custo;
- um maior armazenamento, em média, para acomodar essas flutuações;
- um maior despacho térmico, também em média.

De fato, tudo isto será observado.

Os parâmetros de aversão ao risco utilizados foram, respectivamente, de 25% para a cauda do CVaR, e 35% para o peso do CVaR na medida de risco; os outros 65% são dados pela média do custo, para efetuar uma política que ainda leva em conta os cenários mais favoráveis.

Começamos pelo custo, que é a quantidade que foi alterada na função objetivo.

Notamos, na figura 7.5, que o desvio padrão, mês a mês, do modelo neutro a risco é muito maior do que o do modelo com CVaR. Em praticamente todos os meses, exceto pelo primeiro ano, o desvio padrão do modelo periódico é maior do que o custo médio do modelo com CVaR. Mais ainda, esta diferença se torna gritante nos meses mais caros, onde o desvio padrão fica mais de

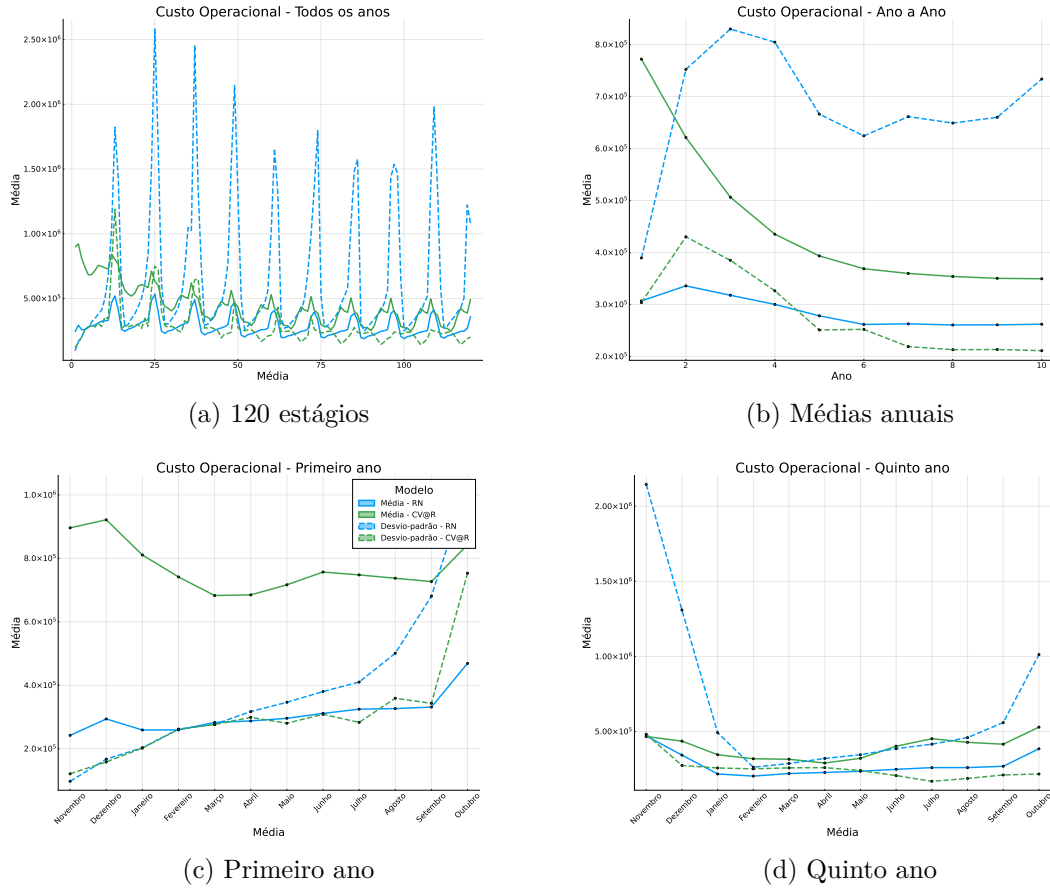


Figura 7.5: Comparativo Custo de operação

três vezes maior do que os custos médios. Por outro lado, o desvio padrão do modelo com CVaR é praticamente sempre menor do que o custo médio correspondente, e às vezes também menor do que o custo do modelo neutro a risco.

Naturalmente, o custo médio para o modelo com CVaR é maior, como podemos ver no gráfico anualizado da figura 7.5b. Se no início este é mais do que o dobro, ao final dos dez anos, que já está mais próximo do regime estacionário, este é apenas em torno de um terço a mais. E, neste ponto, o desvio padrão dos gastos anuais fica nitidamente menor do que os custos médios, tanto no caso avesso ao risco como no caso risco neutro, enquanto que o desvio padrão para a operação neutra a risco é maior do que o dobro do custo médio de operação do caso avesso ao risco.

Esta redução de custos extremos é naturalmente reflexo de uma redução do deficit, como podemos observar na figura 7.6. Aqui, notamos que tanto a média como o desvio padrão do deficit serão reduzidos para a operação avessa a risco. Isto faz bastante sentido, pois a ocorrência de corte de carga já é em si um evento raro, e que justamente é correlacionado aos maiores custos operativos. Assim, ao reduzir a variabilidade dos custos, e, principalmente, a ocorrência de altos custos, o modelo avesso ao risco também resulta em menor ocorrência de deficit.

Passemos à energia armazenada, cujos gráficos podemos ver na figura 7.7,

Aqui, mais uma vez vemos que o modelo de horizonte periódico necessita de um grande tempo de operação para atingir um estado estacionário: é nítido que há ainda aumento da energia armazenada nos reservatórios, para além do sexto ano. Também é interessante notar que há uma menor

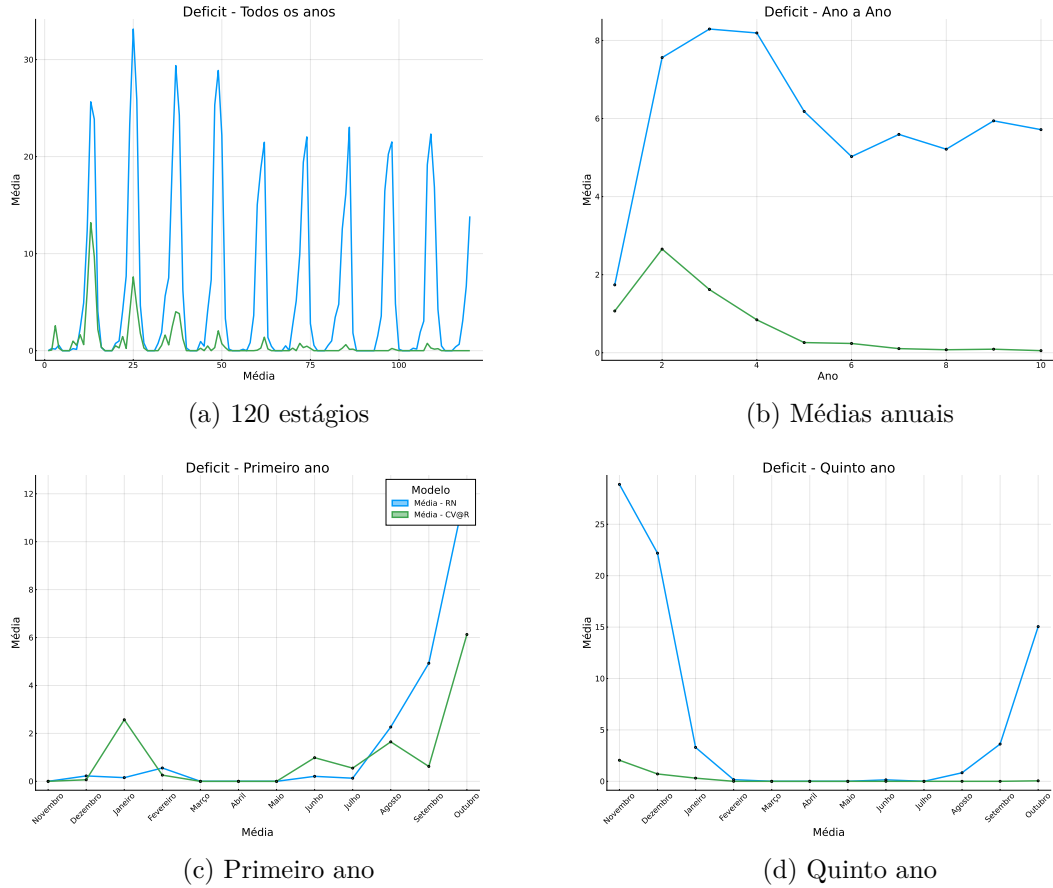
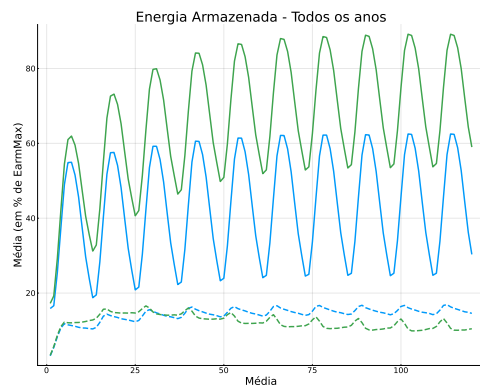


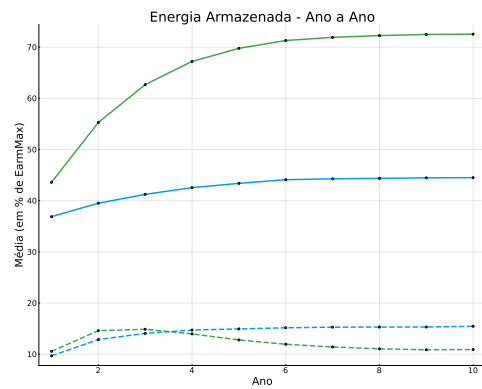
Figura 7.6: Comparativo Deficit

variabilidade desta energia armazenada. Além da própria natureza da operação, vale notar que a energia armazenada está já bem mais próxima do seu máximo (um pouco superior a 290 GWh), e que portanto ela também tem menor capacidade de flutuação. Enfim, e como esperado, a tendência de crescimento da energia armazenada é ainda mais acelerada no modelo avesso ao risco do que no modelo neutro a risco. Isso se reflete até nos gráficos mensais; no primeiro ano, visto na figura 7.7c, é visível o aumento da diferença das energias armazenadas.

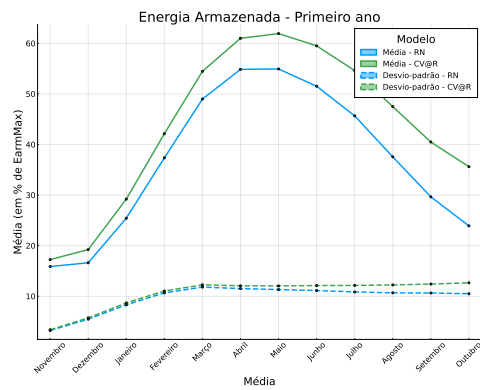
Enfim, podemos ver na figura 7.8 que o modelo avesso ao risco de fato terá uma maior geração térmica consistentemente superior à geração térmica do modelo neutro a risco. O início é muito diferente, já que a condição inicial do sistema está bastante distante do desejado, o que implica em um despacho bem maior de usinas térmicas. Mas, a partir do sexto ano, em que o sistema já está bem mais próximo de um eventual regime estacionário, a diferença é bem menor.



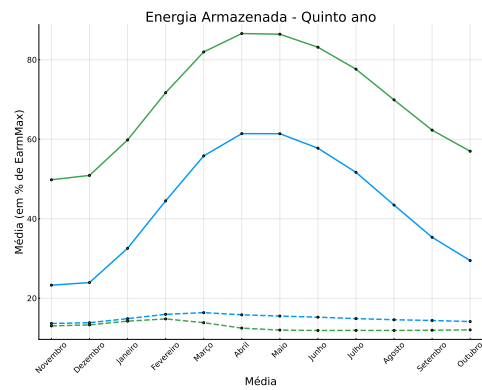
(a) 120 estágios



(b) Médias anuais



(c) Primeiro ano



(d) Quinto ano

Figura 7.7: Comparativo Energia Armazenada

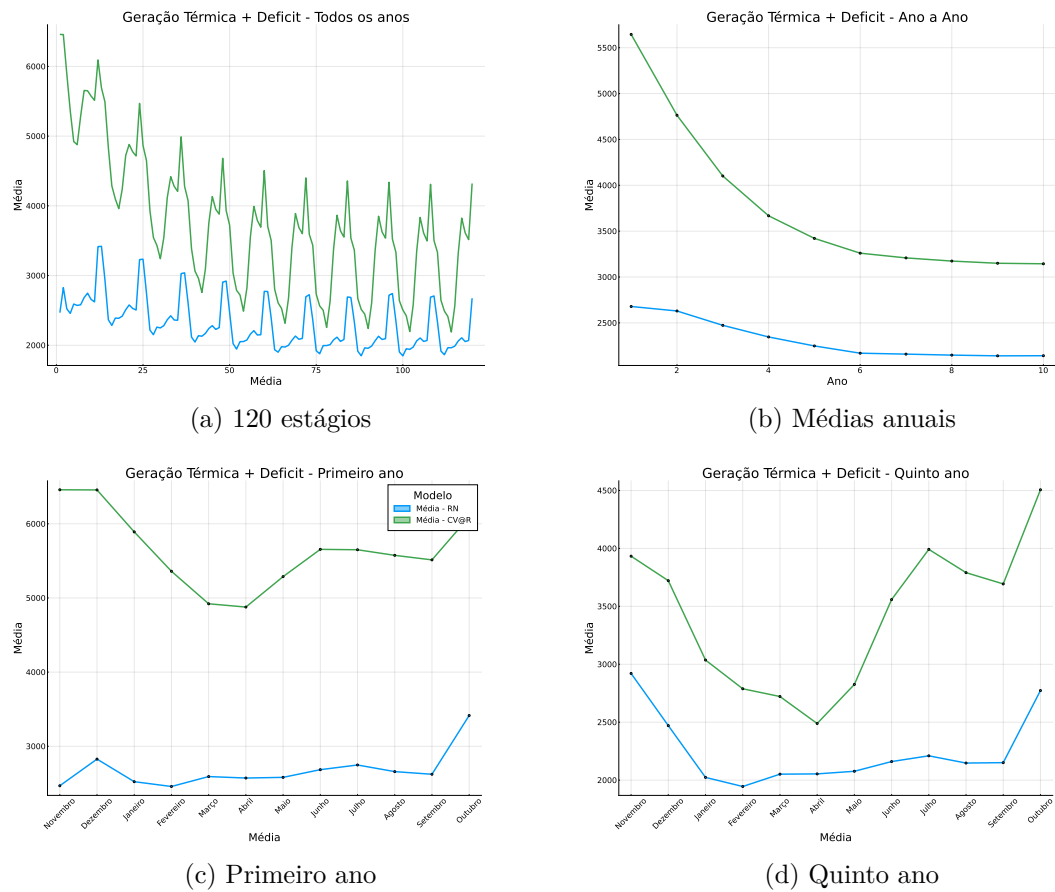


Figura 7.8: Comparativo Geração Térmica

7.4 Comparação entre subsistemas

Nesta seção, veremos como a separação dos subsistemas afeta a operação.

Vamos analisar, primeiramente, a energia armazenada, comparando os modelos periódico e não-periódico. Se, na figura 7.1, o comportamento agregado da energia armazenada nos quatro REEs é bastante suave, vemos que, ao desagregar em subsistemas o modelo de horizonte finito se mostra muito mais instável, como observamos na figura 7.9. Por exemplo, notamos oscilações muito grandes no subsistema 2, correspondente ao REE Sul, que tem um regime de afluições bastante atípico. Estas oscilações não aparecem no modelo periódico, que naturalmente possui uma política muito mais consistente ao longo dos diferentes anos de operação. Uma outra observação interessante é a saturação do subsistema 4, correspondente ao REE Norte, que tem um regime de afluições bem extremo, e com pouca capacidade de armazenamento. Neste, nem é possível notar uma diferença significativa da operação entre os modelos periódico e de horizonte finito, enquanto que nos outros dois reservatórios (1 e 3, respectivamente Sudeste e Sul), que também são os maiores, é nítido que o modelo periódico está promovendo mais armazenamento.

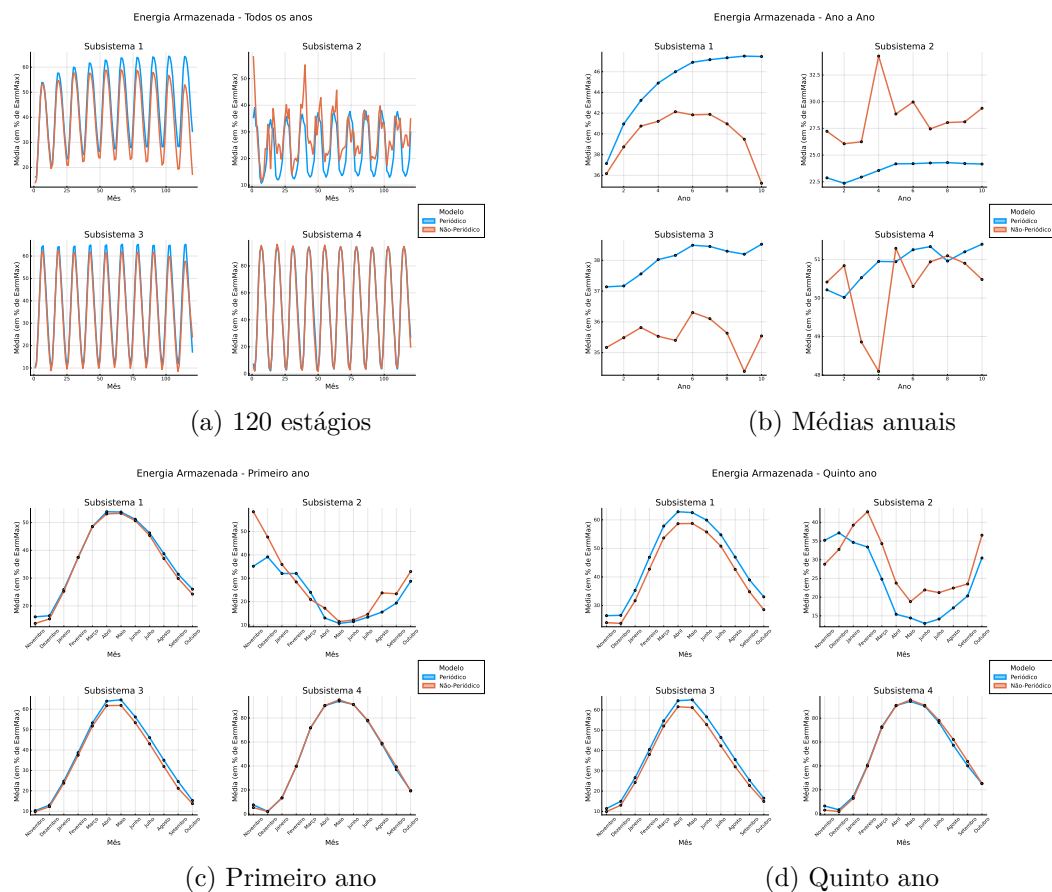


Figura 7.9: Comparativo Energia Armazenada

Enfim, vemos a geração térmica separada por subsistema na figura 7.9, onde também incluímos o caso avesso ao risco. Aqui, o comportamento é muito similar ao agregado: o periódico tem mais geração térmica do que o de horizonte finito, e o caso avesso ao risco mais ainda. Também notamos que, no primeiro ano, o efeito da sazonalidade é praticamente nulo para os subsistemas 1 e 2, que

estão em franca subida de armazenamento, mas que esta sazonalidade vai se estabelecendo até estabilizar mais à frente.

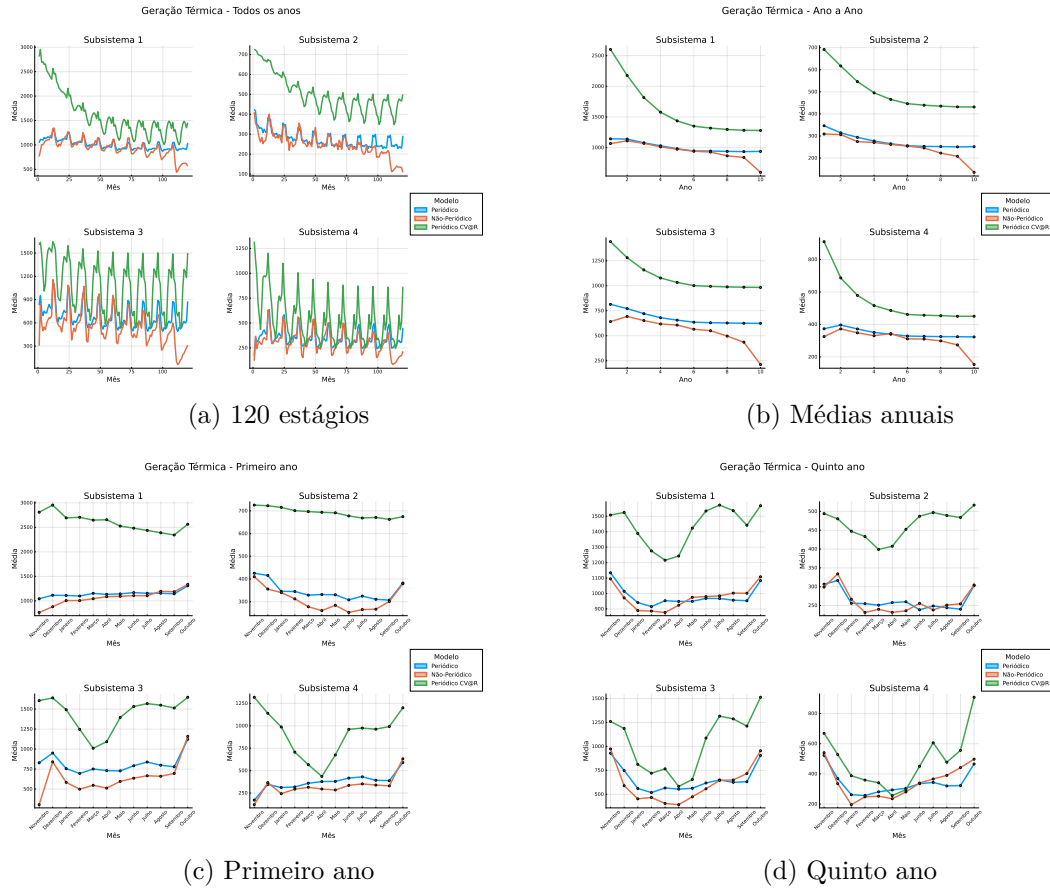


Figura 7.10: Comparativo Geração Térmica

7.5 Construção de cortes e etapa *forward*

Como vimos na seção 3.3.1, é necessário realizar etapas *forward* bastante grandes para garantir a convergência do algoritmo de PDDE. De fato, só assim poderemos garantir que o algoritmo irá visitar todos os possíveis estados e, com isto, realizar os cortes correspondentes que permitiram aproximar corretamente a função de custo futuro para estes estados.

Nesta seção, veremos como a realização de etapas *forward* reduzidas irá impactar negativamente a solução, tanto do ponto de vista do cálculo das funções de custo futuro como através das políticas implementadas. Para tal, destacaremos três casos, a partir do modelo de 4REEs que tratamos nesta seção:

1. O caso base terá *forwards* de tamanho aleatório, e com isto poderá realizar simulações de comprimento bastante grande; ele designado como “**Periódico**” a seguir;
2. O segundo caso terá *forwards* de tamanho 12, visitando portanto apenas uma vez cada estágio; por esta visão limitada dos cenários futuros, chamamos este caso de “**Míope**”;
3. O terceiro caso, designado como “**Semi-Míope**”, terá *forwards* de tamanho 13, e com isto

revisitará o primeiro estágio partindo da condição final do estágio 12, e não apenas a partir da condição inicial do problema.

Vale notar que todos os três modelos usam uma *formulação periódica* do problema de planejamento da operação, e portanto que a função de custo futuro calculada no estágio 1 será usada, após desconto, ao final do estágio 12. A única diferença entre eles será o comprimento das etapas *forward* realizadas a cada iteração, que resultarão em diferentes locais para o cálculo dos cortes.

Para que a comparação seja justa, todos os modelos realizaram o *mesmo número de cortes ao todo*. Isto também corresponde a, essencialmente, o mesmo tempo de solução para cada um dos casos, já que a maior parte do tempo é passada na solução dos PLs.

Vejam, para começar, o comportamento da energia armazenada para cada um destes métodos de cálculo das FCFs, na figura 7.11.

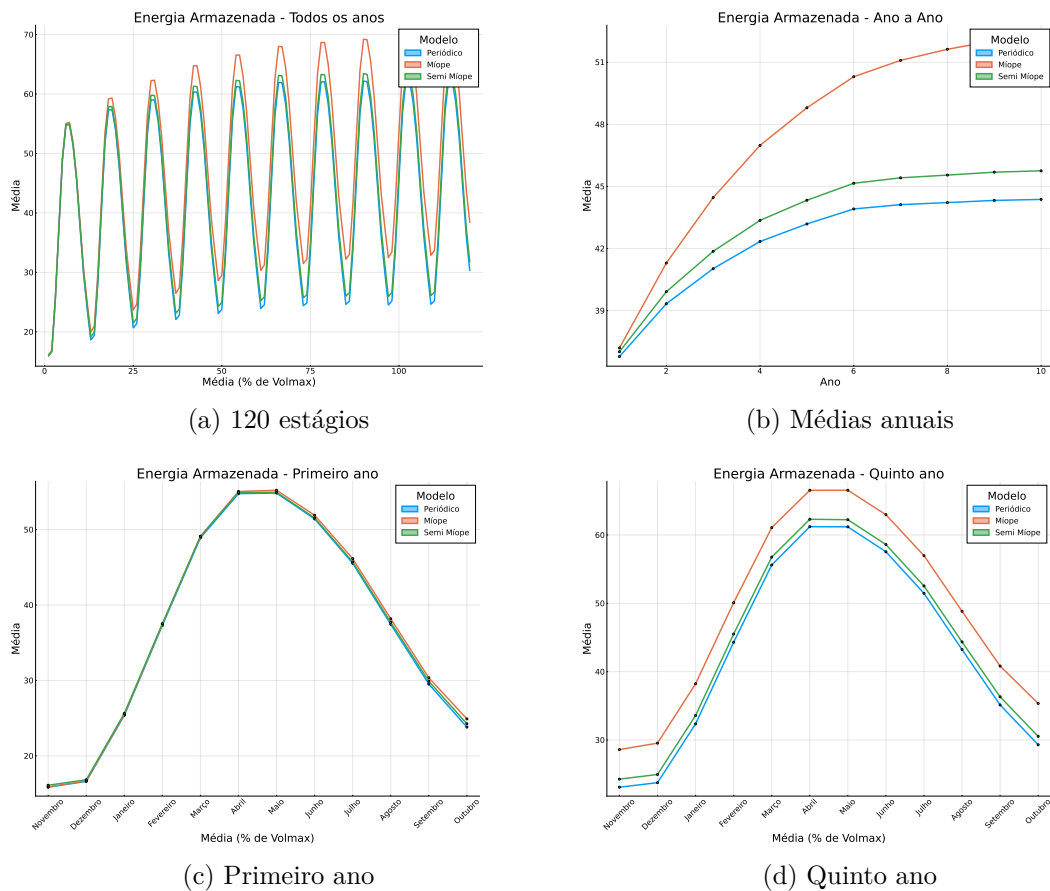


Figura 7.11: Comparativo Energia Armazenada

A grande diferença aparece entre o modelo Míope e os outros dois, o Periódico e o Semi-Míope. Vemos que, mesmo que ao longo do primeiro ano a diferença seja bem pequena, esta aumenta de forma dramática ao longo dos anos seguintes, convergindo para um patamar radicalmente diferente de energia armazenada, se comparado com o nível atingido pelos outros dois modelos. Uma consequência deste maior nível de armazenamento é que o modelo Míope também inicia mais distante do seu nível “alvo”, correspondente a um estado estacionário mais cheio, e portanto demora mais para estabilizar o armazenamento médio.

Apesar de esta ser de fato a diferença mais importante, também vale notar que o modelo Semi-Miope, mesmo não sendo tão extremo como o modelo Miope, também terá o estado estacionário mais alto, divergindo progressivamente do caso Periódico. De forma geral, a adição de um estágio a mais durante a etapa *forward* já reduz de forma significativa a diferença que vimos do caso Miope para o caso Periódico.

Este comportamento é compatível com o que veremos na geração térmica, apresentada na figura 7.12.

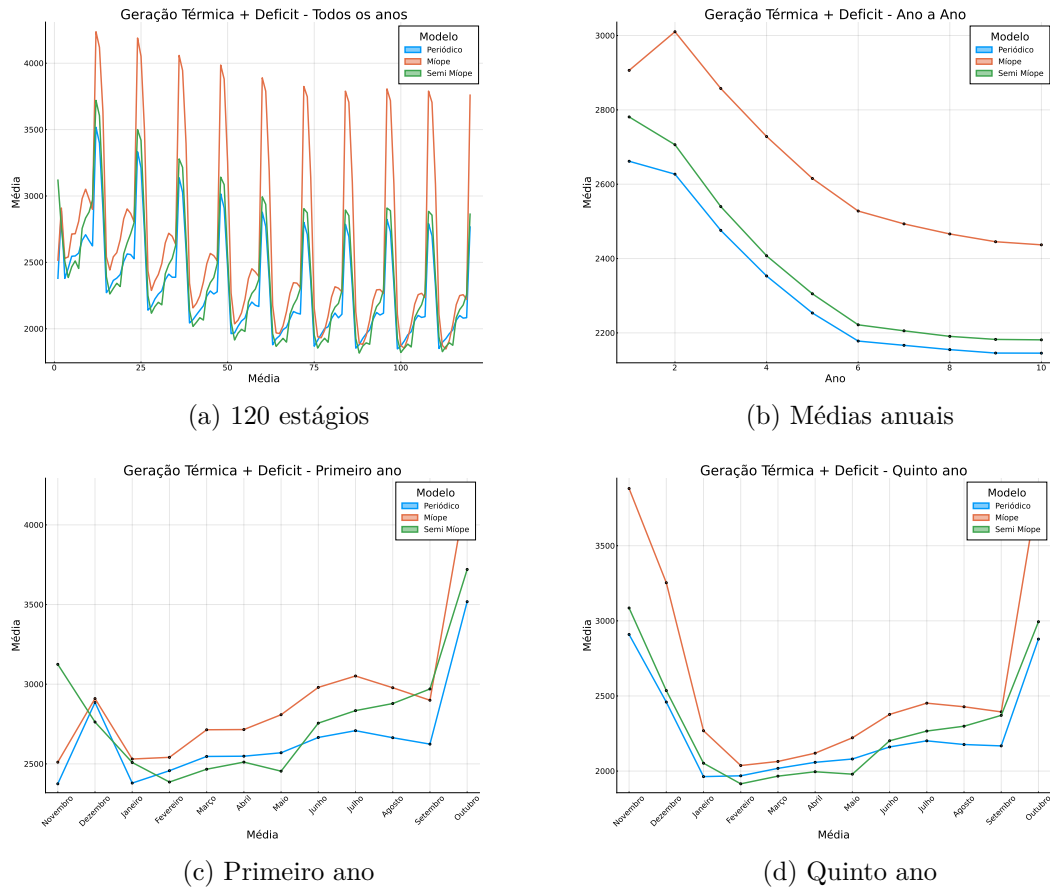


Figura 7.12: Comparativo Geração Térmica

Como esperado, a geração térmica será bem maior para o modelo Miope, o que irá permitir um maior nível de armazenamento médio, se comparado com os outros dois modelos. Novamente, o modelo Semi-Miope terá efeitos similares, mas reduzidos, se comparado com o modelo Miope.

Entretanto, aqui vemos algo além do observado nas energias armazenadas: ao longo do primeiro ano, a diferença de geração térmica entre o caso Miope e o Periódico é menor do que nos anos seguintes, e se dá de forma aproximadamente igualmente distribuída ao longo dos 12 meses. No quinto ano, já mais próximo do regime estacionário, a diferença de geração térmica está muito mais concentrada nos dois primeiros meses, e no último mês. Já o modelo Semi-Miope também usa mais energia térmica nos “extremos”, mas tem um uso levemente inferior nos meses úmidos.

Naturalmente, essa diferença de despacho também se reflete nos custos da operação, que podemos ver na figura 7.13. Aqui, tanto a diferença de custos ao longo de todo o primeiro ano, bem como a diferença nos meses de Fevereiro a Setembro, é bastante reduzida. Além disso, o modelo

Semi-Míope acaba sendo até menos caro do que o modelo Periódico ao longo dos anos finais, mesmo que esta diferença seja bastante pequena. Ainda assim, vale notar a radical diferença de custos do modelo Míope para os outros dois, e que portanto podemos atribuir, principalmente, aos maiores custos nos meses de Outubro a Janeiro, que também são os picos de geração térmica e déficit.

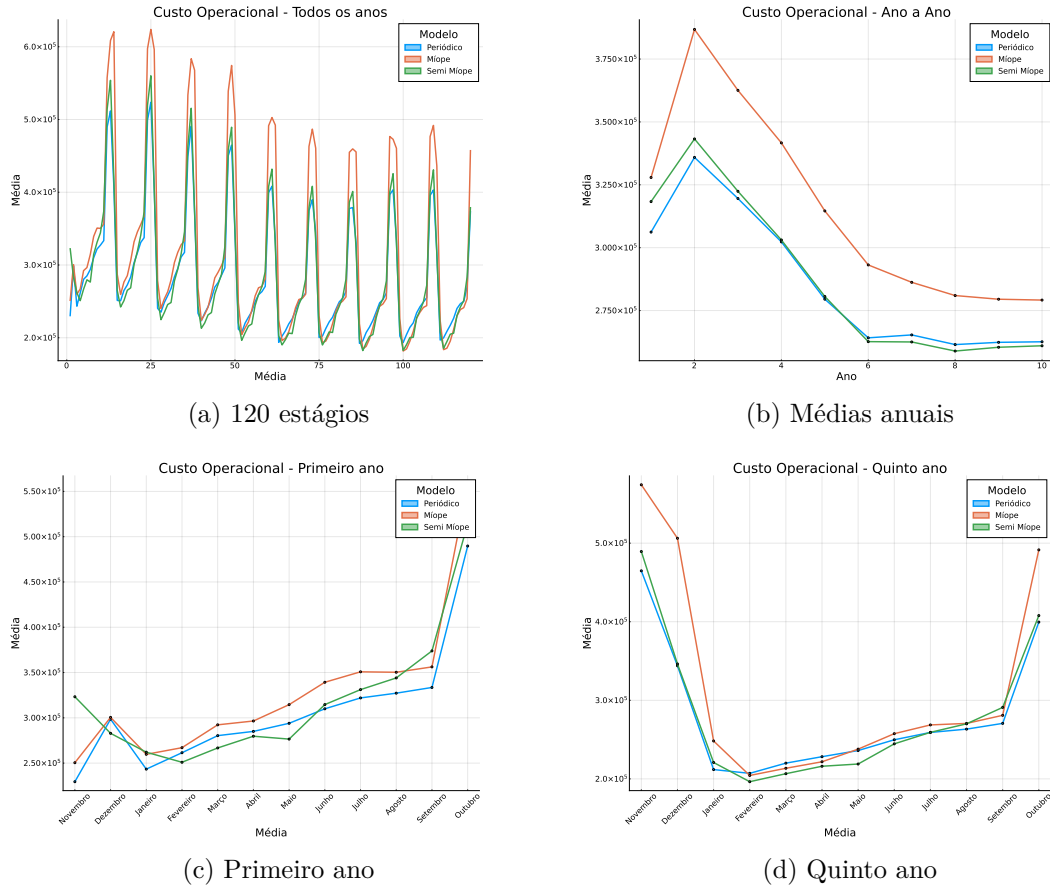


Figura 7.13: Comparativo Custo de Operação

Esta diferença de política, em que o modelo míope busca armazenamentos mais altos, pode ser compreendida ao analisar as funções de custo futuro construídas em cada um dos casos. Como estas funções dependem de quatro variáveis de armazenamento, é impossível realizar um gráfico que as representem diretamente; para ainda assim fazermos uma análise, utilizaremos os gráficos das *envoltórias superior e inferior* destas funções, como apresentadas na seção A. Estas envoltórias serão *projeções* do gráfico da função de custo futuro, permitindo uma análise visual das mesmas; vale observar que a envoltória inferior é convexa por construção, mas a superior, por sua natureza de “achatar variáveis”, não será nem convexa nem côncava.

No nosso caso, as envoltórias serão construídas de acordo com a energia armazenada *total* dos quatro REEs. Assim, a envoltória inferior corresponde ao *menor custo* para este armazenamento total (ou seja, à “*melhor*” distribuição entre os reservatórios), enquanto que a envoltória superior corresponde ao *maior custo* (e, analogamente, à “*pior*” distribuição da energia).

Vejamos, primeiramente, a diferença na função de custo futuro do primeiro mês de cálculo da política, ou seja, Novembro. Aqui, vale lembrar que o modelo Periódico terá realizado cortes em diferentes estados, pois este terá feito *forwards* de tamanho bastante grandes, e com isto terá a

possibilidade de explorar tanto armazenamentos mais altos como armazenamentos mais baixos ao final da estação seca. Por outro lado, o modelo Míope irá apenas ter os cortes correspondentes à condição inicial, e portanto sempre muito próximos. Já o modelo Semi-Míope, mesmo visitando mais cenários, terá apenas o “tempo” de um ano para atingir um outro estado de armazenamento, e portanto já terá cortes em regiões mais diversas, mas ainda assim provavelmente sem cobrir todo o espectro de possibilidades.

Para ter uma ideia desta diferença, apresentamos na figura 7.14 o gráfico de violinos para o armazenamento total ao final de Novembro, ao longo dos dez anos de simulação do modelo periódico.

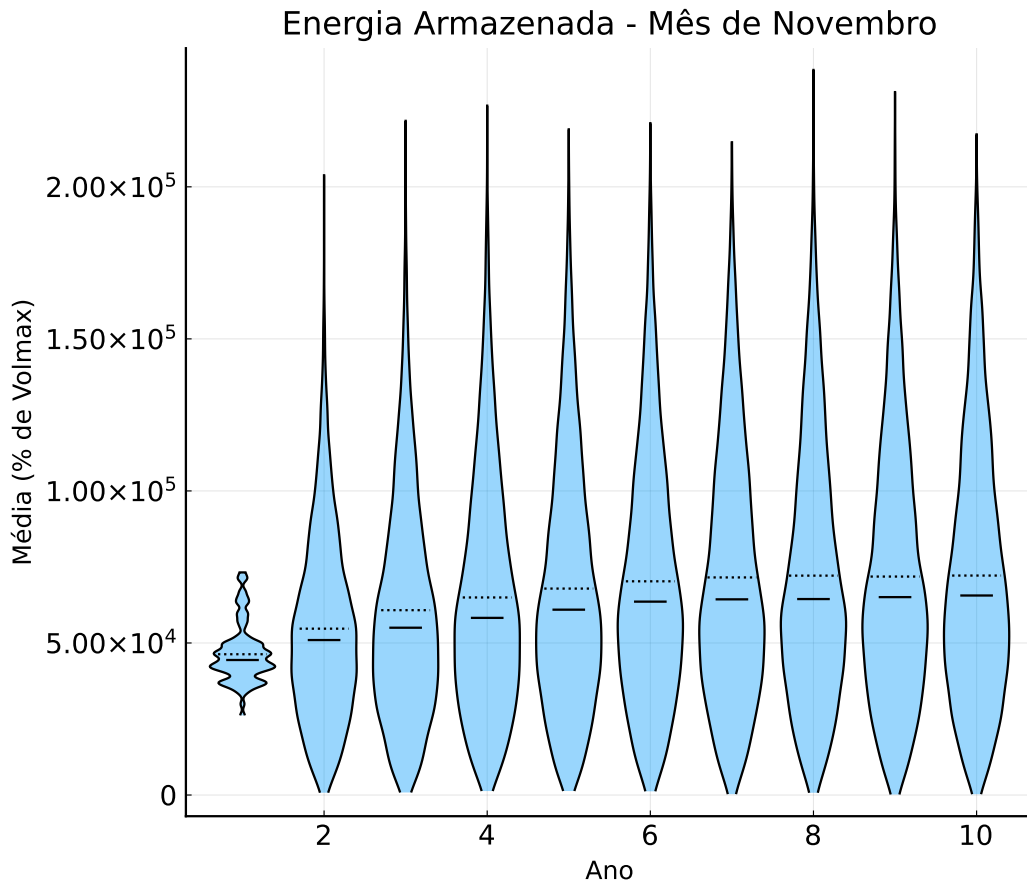


Figura 7.14: Distribuição da Energia Armazenada em Novembro, simulação num horizonte de 10 anos

Agora, vejamos o gráfico das envoltórias das funções de custo futuro na figura 7.15.

Aqui, novamente, vemos uma grande diferença qualitativa entre as funções de custo futuro do caso Periódico e do caso Semi-Míope para o caso Míope. De fato, a função no caso Míope tem uma forma muito mais simples, indicando que de fato ela tem cortes muito próximos, enquanto que as outras duas apresentam uma variação mais significativa perto dos volumes armazenados mais baixos, refletindo que, especialmente quando há pouco recurso, é especialmente importante que este esteja “bem distribuído” para que os custos não fiquem maiores ainda.

Mas esta não é a maior diferença, nem a mais relevante para explicar os comportamentos vistos nos gráficos anteriores. De fato, é importante notar que o caso Míope atribui custo zero para armazenamentos suficientemente altos, e, de forma complementar, a derivada sendo praticamente

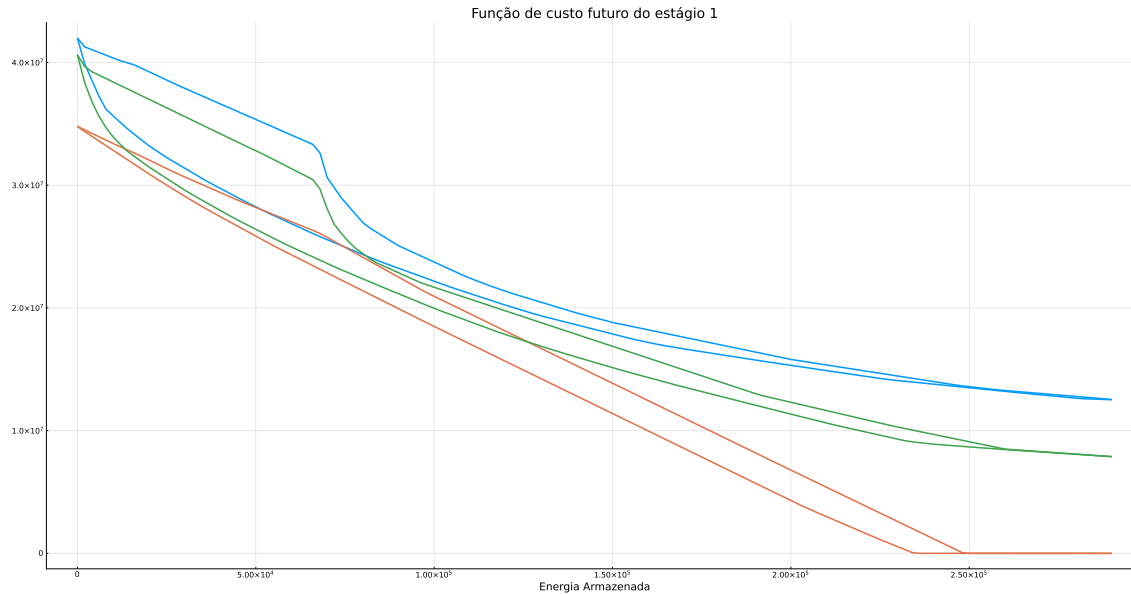


Figura 7.15: Comparação das envoltórias das Funções de Custo Futuro do primeiro mês, Novembro, para diferentes comprimentos da etapa *Forward*

constante, também dá um valor da água praticamente constante ao longo de todo o resto do armazenamento. Isto tem dois efeitos: primeiramente, o esforço neste mês para “encher o reservatório” será bem maior, pois corresponde a um custo marginal que só será igualado nos outros casos quando o armazenamento estiver por volta de 20%. Em segundo lugar, também não há incentivo para armazenar acima de 85% da energia armazenada máxima, já que para este caso o valor da água é nulo. Como este mês corresponde ao final da estação seca, em que o reservatório estará em geral mais vazio, o efeito dominante é o primeiro.

Isto explica, como vimos, o gasto significativamente maior dedicado à geração térmica neste mês, e a correspondente maior energia armazenada que se acumula ao longo dos anos da simulação.

Já o caso Semi-Míope não tem estes efeitos de forma tão pronunciada. Primeiramente, o valor da água irá, de fato, decair paulatinamente conforme a energia armazenada aumenta, sem chegar a zero em momento nenhum. E, mesmo que haja uma diferença de valores em termos absolutos, o que importa para a decisão de operação é justamente o valor da água informado pela função de custo futuro. Assim, estas são bastante similares, a menos de uma pequena defasagem, até praticamente a metade do valor de energia armazenada, quando o modelo Semi-Míope começa a ficar mais defasado. Novamente, uma queda mais forte da função de custo futuro representa um valor da água maior por mais tempo, o que, mesmo que não seja tão enfático como no caso do modelo Míope, também explica um armazenamento e uso de geração térmica maiores.

Podemos contrastar também as funções de custo futuro do estágio 12. Vale notar que, como todos os modelos são periódicos, todos eles irão utilizar a função de custo futuro do estágio 1 como a do estágio 13, convenientemente descontada, e portanto a função que acabamos de analisar terá um impacto na construção da função de custo futuro do estágio 12. As três funções, com a mesma convenção, podem ser vistas na figura 7.16.

Aqui, após 12 estágios, e ao final da estação seca do primeiro ano, vemos que todas as funções indicam custos com bastante variação para um baixo armazenamento, o que é coerente com o esperado. Também podemos ver que o modelo Míope ainda atribui custo zero para armazenamentos

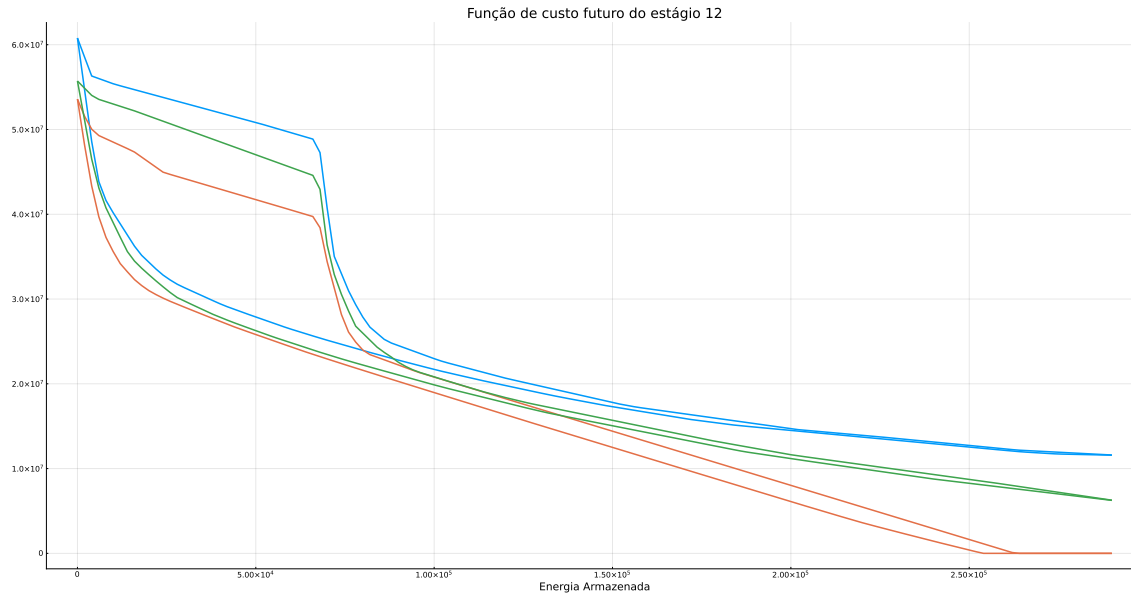


Figura 7.16: Comparação das envoltórias das Funções de Custo Futuro do décimo segundo mês, Outubro, para diferentes comprimentos da etapa *Forward*

muito grandes, e que de 40% a 80%, aproximadamente, o modelo ainda será bastante linear. Isso mais uma vez se explica porque a solução do problema inicia de um estado de armazenamento relativamente baixo, e portanto, ao final da estação seca, é esperado que haja apenas estados também baixos visitados.

Já o modelo Semi-Míope, que não visita mais de uma vez este último estágio, irá também ter os cortes essencialmente dados nesta mesma região de baixo armazenamento, mas, por já contar com uma estimativa melhor dos custos do que seria o estágio 13, não tem uma inclinação tão acentuada para os cortes que estarão ativos no caso de *alto* armazenamento. O modelo Periódico, por sua vez, tem a oportunidade de passar diversas vezes no último estágio, e com isto tem a possibilidade de gerar séries positivas de longa duração, o que permitirá chegar, algumas poucas vezes, em estados de maior armazenamento, mesmo que seja o final da estação seca. Isto faz que haja uma leve curvatura nesta região, correspondente a um valor da água paulatinamente decrescente com o armazenamento, o que é muito mais consistente.

8 Conclusão

Neste projeto, apresentamos tanto a formulação de problemas periódicos de otimização estocástica, como as modificações na PDDE para resolvê-los, e uma caracterização do ponto de vista teórico das soluções. Em seguida, exploramos diversos casos de estudo, ressaltando as vantagens da aplicação desta metodologia ao problema de planejamento da operação.

A mudança fundamental no caso periódico é utilizar cortes não apenas de um estágio para um anterior, mas também *para um estágio futuro*. Esta alteração permite, ao mesmo tempo, capturar a estrutura periódica do problema e traduzir a informação de um horizonte infinito em um número finito de funções de custo futuro. Com isto, também aumentamos a possibilidade de compartilhamento de cortes, o que acelera a convergência para a solução.

Mais ainda, dada esta estrutura periódica e fixadas as funções de custo futuro ao final do cálculo da política, podemos realizar simulações para estudar como isto se reflete no planejamento da operação. Em teoria, sabemos que as trajetórias simuladas correspondem a uma cadeia de Markov; em geral, estas podem possuir diversas distribuições estacionárias, mas é razoável esperar que, em muitos casos, esta distribuição seja única e, ademais, que independentemente das condições iniciais o sistema convirja para esta distribuição.

Assim, a partir do protótipo na biblioteca `DisjHTPlan`, modelamos alguns problemas de planejamento da operação, com períodos 1 (anual), 2 (semestral) e 12 (mensal). Pudemos observar que em todos os casos de estudo a energia armazenada convergiu a um estado estacionário, após um número variável de estágios segundo cada caso, conforme observado nos gráficos das seções 6 e 7. Nos casos de período 2 e 12, observamos que os armazenamentos apresentaram uma “oscilação natural”, revelando a convergência para o estado estacionário. Esta convergência também foi verificada pelos testes de Kolmogorov-Smirnov, que quantificaram, estatisticamente, a proximidade das distribuições de probabilidade de energia armazenada ao longo dos diferentes estágios.

Também pudemos comparar outras variáveis de interesse das simulações, como deficit, geração térmica e hidroelétrica, entre outras. Mesmo que estas nem sempre levem o mesmo número de períodos para convergir, este sempre foi, no máximo, igual ao tempo que levou a energia armazenada. Pudemos observar alguns casos, especialmente para a geração térmica, onde esta se estabilizou antes da convergência da energia armazenada.

Além de modelos neutros a risco, também apresentamos resultados para modelos com aversão ao risco, e analisamos as principais diferenças que esta mudança na função objetivo acarreta, tanto para o regime estacionário, como para a política de operação e os custos decorrentes. Com isto, o modelo se torna bastante mais conservador, aumentando significativamente o armazenamento, e incorrendo em aproximadamente um terço a mais de custo operativo. Isso se reflete em uma redução significativa da variabilidade da operação, que, já próximo do regime estacionário, fica em torno de 30% da variabilidade do modelo neutro a risco.

8.1 Propriedades do modelo periódico

Um dos casos estudados diz respeito à mudança das condições iniciais, tanto para a simulação como para o próprio cálculo da política, e das funções de custo futuro, feito na seção 6.2.2. Ainda que seja um exemplo em dimensão um, ele ilustra de forma bastante clara que as funções de custo futuro estimadas em cada um dos casos são, essencialmente, as mesmas. Isso é bastante coerente com a própria noção de função de custo futuro, pois ela deve depender apenas do *estado do sistema*, e não da condição inicial que foi usada para a simulação. Naturalmente, é relevante que as funções de custo futuro estejam bem-aproximadas na vizinhança dos pontos onde elas serão utilizadas, mas

esta metodologia traz, de forma intrínseca e sem custo adicional perceptível, a possibilidade de construir boas aproximações das funções de custo futuro.

Nesta mesma linha, os casos de estudo da seção 7.5 também enfatizam a necessidade de bem aproximar as funções de custo futuro. Ao alterar o tamanho das etapas *forward*, mantendo idêntica a quantidade de cortes, vimos que é importante que estes sejam feitos em todos os estados possíveis de armazenamento, pois caso contrário, mesmo usando um modelo periódico, estas funções irão carregar uma informação incompleta. Daí, mesmo que o modelo pareça convergir, a operação resultante será significativamente desfavorável.

Uma outra mudança interessante que podemos ver a partir dos casos estudados é o efeito das restrições de armazenamento e transmissão, comparando o modelo totalmente agregado da seção 6.4 com o modelo de 4REEs da seção 7, ambos no caso neutro a risco. Mesmo que os reservatórios ainda amortecem as variações de afluência, o modelo de 4REEs não consegue mais operar em um regime “essencialmente constante” de geração térmica ao longo do ano.

8.2 Comparação com modelos de horizonte finito

Também analisamos o comportamento das soluções obtidas pelos modelos periódicos, comparando com as soluções dos modelos de horizonte finito. Como este último é afetado pelo efeito de final do horizonte, muitas vezes ele não possui nem dois períodos para os quais as distribuições da energia armazenada estejam próximas. Naturalmente, isso também se reflete em outras variáveis de decisão, que também não apresentam um padrão tão regular de convergência como visto no caso periódico.

Por causa disto, o modelo periódico apresentou maior regularidade de operação, tanto nos casos simples da seção 6 como nos mais próximos do modelo do setor, da seção 7. Isto pode ser destacado com uma gestão mais adaptada dos recursos, que distribui de forma mais homogênea a geração térmica ao longo do ano, e com isto reduz a necessidade de uso de unidades com maior CVU em períodos de escassez, especialmente ao final da estação seca e início da estação úmida. A resultante é que, mesmo que o uso total de geração térmica seja maior no modelo periódico, como vimos na figura 7.2b, ao se considerar a não-linearidade dos custos de geração térmica, a operação que concentra mais geração no final da estação seca, obtida através do modelo de 120 estágios, se revela significativamente mais dispendiosa, o que observamos tanto no agregado anual como mensal da figura 7.3.

Outra observação interessante é que a operação já se revela diferente no primeiro ano e, de fato, já desde o primeiro mês, com uma tendência de maior geração térmica para o regime periódico, levando a um maior nível de armazenamento. Mesmo esta sendo uma pequena diferença, ao se acumular ao longo de alguns anos de operação, podemos notar que o sistema chegará a estados diferentes de armazenamento, e com diferentes capacidades de amortecer o impacto das incertezas.

Enfim, vale notar que estes estudos já indicam que o número de estágios necessários para “neutralizar” o efeito de bordo causado pelo final de horizonte é superior ao equivalente de 7 anos. De fato, se partindo de uma condição inicial fixa, o modelo de horizonte periódico leva mais do que este tempo para estar, estatisticamente, próximo do estado estacionário, por simetria, o efeito de final de horizonte também só poderá ser estatisticamente desconsiderado se houver mais do que este intervalo de tempo para o sistema “reverter à média”. Naturalmente, uma forma garantida de evitar estes problemas, e obter uma operação mais consistente, é utilizar um modelo periódico ao final do período de planejamento.

8.3 Estabilidade numérica

Os modelos mais completos da seção 7 também apresentam maior dificuldade numérica para serem resolvidos. Em parte por aumentar a dimensão da variável de estado, em parte também pela necessidade de um maior número de cortes para representar as funções de custo futuro, e enfim pela própria natureza dos problemas, ao longo do algoritmo de PDDE os PLs que são resolvidos se tornam cada vez mais complexos e, com isto, mais sujeitos a erro numérico.

Ao resolver uma quantidade também muito grande de problemas de programação linear, não é surpreendente que possa haver erros numéricos devido à precisão de cálculo com ponto flutuante. Na maior parte dos casos, é relativamente simples reinicializar o *solver* para que ele parta de um ponto que seja melhor condicionado, mas, naturalmente, este procedimento pode levar à perda de performance. Assim, um equilíbrio foi obtido ao estabelecer um tempo máximo para que o *solver* convirja, e apenas reinicializar em caso de erro ou de limite de tempo. Vale notar que, ao manter o estado *solver*, tiramos partido da possibilidade de *partida a quente*, ou seja, de uma inicialização próxima da solução ótima, o que acelera de forma significativa as iterações, em especial a construção dos cortes durante a etapa *backwards*.

8.4 Trabalhos futuros

A primeira direção, do ponto de vista de modelagem do problema é continuar o desenvolvimento do protótipo para que este possa lidar com modelos estocásticos mais complexos. De fato, para o caso agregado, notamos que o modelo de energias afluentes independentes não dá resultados coerentes ao mudarmos a escala de discretização do tempo, convergindo mais rapidamente ao estado estacionário, essencialmente em 4 anos no caso mensal em vez de mais de 8 no caso anual.

A razão que vemos para isto é que o modelo independente, para uma discretização mensal, gera anos de fluências muito menos extremas do que o modelo independente para uma discretização anual. Assim, ao implementar modelos autorregressivos para a série de fluências, podemos obter resultados mais consistentes para a incerteza utilizada e, com isto, também ter resultados mais compatíveis com discretizações distintas.

Enfim, e neste mesmo espírito, iremos estudar a possibilidade de incorporar modelos com diferentes escalas temporais ao longo do horizonte de planejamento, para podermos generalizar a aplicação de modelos periódicos.

Na mesma linha, podemos ver o modelo independente como uma política de operação que tem como entrada um ruído (ou perturbação) que sempre segue a mesma distribuição. Ora, no caso autorregressivo, ao partir de um estado conhecido, será necessário primeiro que o próprio processo estocástico das fluências convirja à sua distribuição estacionária para que o processo das energias armazenadas, dependente dele, possa convergir. Assim, é natural esperar que a convergência do estado de armazenamento para o regime estacionário seja ainda mais lenta do que a observada neste trabalho, e seria interessante quantificar este efeito.

Enfim, em uma direção complementar, é natural considerar o impacto do fator de desconto β na solução do problema. Um valor de β mais próximo de um necessitaria, em teoria, simulações *forward* mais longas, mas, como indicamos no final da seção 3, poderia não ser necessário para um problema específico. De toda forma, já há uma discussão em curso no setor elétrico sobre o papel deste parâmetro e, mesmo que do ponto de vista matemático ele seja importante para a convergência do problema aqui formulado, a sua interpretação econômica e seu impacto na decisão de planejamento também é importante de ser analisado.

A Envoltórias da função de custo futuro

Uma das maiores dificuldades ao se comparar a solução de problemas de otimização estocástica multi-estágio através de algoritmos de decomposição é poder comparar as diferentes funções de custo futuro que são construídas ao longo do processo. A principal razão é que elas são objetos em alta dimensão, e portanto qualquer representação de ambas será parcial. Nesta seção, apresentamos duas formas de construir uma visualização gráfica das funções de custo futuro.

A.1 Envoltória inferior

Considere uma função (convexa) descrita por cortes como habitual na PDDE:

$$\Omega(x) = \max_i \{a_i \cdot x + b_i\}.$$

Também podemos representar Ω como um problema de otimização:

$$\begin{aligned} \Omega(x) = \min_{\theta} \quad & \theta \\ \text{s.a} \quad & \theta \geq a_i \cdot x + b_i \text{ para todo } i. \end{aligned} \tag{A.1}$$

Daqui, podemos introduzir um parâmetro y , de dimensão baixa (preferencialmente 1 ou 2), que representará uma *agregação* das variáveis originais x . Um exemplo desta construção é, por exemplo, considerar $y = \sum_j x_j$, se as variáveis x_j correspondem a diferentes energias armazenadas, e portanto y representa o armazenamento total do sistema, fazendo abstração dos diferentes reservatórios (equivalentes ou individuais).

Seguindo este exemplo, podemos definir

$$\begin{aligned} \phi(y) = \min_x \quad & \Omega(x) \\ \text{s.a} \quad & \sum_j x_j = y \\ & x \in X \end{aligned}$$

onde X é o conjunto viável das variáveis de estado x . Como a função Ω é *convexa*, e a restrição que impusemos também, o problema resultante é convexo. Mais ainda, como a dependência deste problema na variável y se dá apenas pela variação do lado direito, a própria função $\phi(y)$ assim definida também será convexa.

Graças à representação (A.1), que também é um problema de minimização, obtemos uma formulação direta e explícita para o PL que permite calcular o mínimo da função Ω ao longo de $\sum_j x_j = y$:

$$\begin{aligned} \phi(y) = \min_{x, \theta} \quad & \theta \\ \text{s.a} \quad & \theta \geq a_i \cdot x + b_i \text{ para todo } i \\ & \sum_j x_j = y \\ & x \in X. \end{aligned} \tag{A.2}$$

Retornando ao contexto do exemplo, a função $\phi(y)$ representa o mínimo do custo futuro dado que o armazenamento total seja igual a y . Em particular, o valor das variáveis x_j encontrados ao resolver o PL (A.2) dá a “melhor distribuição possível” da energia armazenada nos reservatórios, restrita ao armazenamento total ser y .

Observação A.1. Naturalmente, este exemplo pode ser generalizado. Em primeiro lugar, qualquer relação *linear* entre as variáveis x e y resulta em um problema convexo, e portanto computacionalmente viável — ainda mais que assumimos o próprio cálculo das funções de custo futuro dentro dos

problemas de cada estágio da PDDE. Da mesma forma, podemos empregar qualquer representação *convexa* que tenhamos da função $\mathfrak{Q}(x)$ como um problema de *minimização*: assim, poderíamos ter desigualdades convexas mais gerais, e não apenas cortes lineares, para descrever $\mathfrak{Q}$.

Em segundo lugar, e de forma um pouco mais delicada, também poderíamos usar relações não-lineares entre x e y . Por exemplo, uma relação da forma $y = f(x)$, com f convexa, poderia ser substituída pela restrição convexa $y \geq f(x)$, *sob a condição que a função $\mathfrak{Q}(x)$ seja decrescente em x e que f seja crescente em x .*

A.2 Envoltória superior

Naturalmente, dada uma agregação $f(x) = y$, também podemos considerar o *pior caso*, ou seja,

$$\begin{aligned} U(y) = \max_x \quad & \mathfrak{Q}(x) \\ \text{s.a} \quad & f(x) = y \\ & x \in X. \end{aligned}$$

Retomando a interpretação de x como energia armazenada, esta situação corresponde à pior distribuição desta energia nos reservatórios disponíveis. Assim, ao calcular uma envoltória inferior e superior, temos como avaliar o impacto de ter um armazenamento melhor ou pior distribuído.

Para a envoltória superior, no entanto, a situação é mais complexa. O problema de maximizar uma função convexa em um domínio convexo é teoricamente simples: o máximo é atingido em um ponto extremo. Entretanto, o conjunto viável

$$X_y := X \cap \{x \mid f(x) = y\}$$

pode possuir muitos pontos extremos, o que dificulta uma solução computacional eficiente. Se, ao contrário, não desejarmos enumerar os pontos extremos, deveremos resolver um problema não-convexo.

Portanto, há duas possibilidades. Se a dimensão de x for baixa, e as restrições em X simples, é possível que haja poucos pontos extremos. Neste caso, enumeramos estes pontos extremos x_k de X_y , e calculamos $\mathfrak{Q}(x_k)$ para cada um deles, e tomamos o máximo.

Se a dimensão de X for alta, podemos utilizar a formulação de maximização de $\mathfrak{Q}$:

$$\begin{aligned} U(y) = \max_{x, \theta_i, z_i} \quad & \sum \theta_i \\ \text{s.a} \quad & f(x) = y \\ & \theta_i \leq a_i \cdot x + b_i \\ & \theta_i \leq Q_{\min}(1 - z_i) \\ & \sum z_i = 1 \\ & x \in X, z_i \in \{0, 1\}, \end{aligned}$$

onde $Q_{\min}$ é uma cota inferior para Q , por exemplo a empregada para iniciar o algoritmo de PDDE. Nesta formulação, representamos cada corte por uma variável θ_i separada, e incluímos as variáveis z_i , binárias, para representar a disjunção de que $\mathfrak{Q}(x)$ será igual a apenas um dos cortes. (No caso em que mais de um corte esteja ativo, qualquer um deles dá o mesmo valor para θ_i) Este é um problema Inteiro-Misto, com tantas variáveis binárias como cortes.

B Cadeias de Markov

Nesta seção, buscamos apresentar de forma sucinta os principais aspectos da teoria de cadeias de Markov em espaço de estados gerais. Nossa apresentação é guiada pelas seguintes perguntas, que

são fundamentais no estudo de tal teoria, e que se relacionam intimamente com a noção de estados estacionários e regimes periódicos, apresentadas na seção 4.

1. Dada uma cadeia de Markov, quais são as condições necessárias e suficientes para que exista uma distribuição de probabilidade estacionária π ?
2. Se a distribuição existe, em quais condições ela é única?
3. Se a distribuição for única, em quais condições a cadeia de Markov “converge” para ela?

A noção de convergência para uma cadeia de Markov será detalhada ao longo desta seção, pois precisamos introduzir alguns conceitos antes de poder dar uma definição precisa da mesma.

Como uma boa quantidade de resultados obtidos para cadeias de Markov em espaços de estados enumeráveis podem ser traduzidos para as cadeias de Markov em espaços de estado contínuos, vamos buscar, sempre que possível, enunciar e provar os resultados para o caso enumerável, e trabalhar diretamente com o caso contínuo quando for estritamente necessário.

B.1 Cadeias de Markov com espaço de estados discretos

O conceito fundamental na teoria das cadeia de Markov é o de probabilidades de transição. Antes de chegarmos a ele, precisamos de algumas definições associadas à teoria das matrizes.

Definição B.1 (Matrizes estocásticas). Seja S um conjunto enumerável. A matriz $\mathcal{P} = (\mathcal{P}_{ij})$ de dimensões $\#S \times \#S$ é dita *estocástica* se, para cada $i \in S$, temos

$$\sum_{j \in S} \mathcal{P}_{ij} = 1 \quad \text{e} \quad \mathcal{P}_{ij} \geq 0.$$

Um exemplo de matriz estocástica para o espaço de estados $S := \{A, B, C\}$ é o seguinte

$$\mathcal{P} := \begin{bmatrix} 1/3 & 1/3 & 1/3 \\ 2/5 & 2/5 & 1/5 \\ 0 & 1 & 0 \end{bmatrix}. \quad (\text{B.1})$$

O conceito de matriz estocástica será ligado futuramente com o de matriz de transição de uma cadeia de Markov. Um segundo conceito primitivo nesta teoria é o seguinte.

Definição B.2 (Distribuição de probabilidades em um conjunto). Seja S um conjunto enumerável. O vetor $\pi = (\pi(i))_{i \in S}$ é dito uma *distribuição de probabilidade sobre S* se satisfaz as seguintes condições:

$$\sum_{i \in S} \pi(i) = 1 \quad \text{e} \quad \pi(i) \geq 0.$$

Um exemplo de distribuição de probabilidades no espaço de estados $S = \{A, B, C\}$ é

$$\pi = (1/3, 1/3, 1/3).$$

Com estes objetos básicos definidos, podemos passar para a definição do principal conceito da teoria das cadeias de Markov. Intuitivamente, cadeias de Markov são processos estocásticos cujo estado futuro só depende do estado atual, e não do passado do processo. Esta condição é formalizada a partir da seguinte definição:

Definição B.3 (Cadeia de Markov). Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidades e $X = (X_t)_{t \in \mathbb{N}}$ um processo estocástico definido neste espaço e tomando valores no conjunto enumerável S . Dizemos que o processo X é uma cadeia de Markov se satisfaz a seguinte propriedade:

$$P(X_{t+1} = j_{t+1} | X_t = j_t, \dots, X_0 = j_0) = P(X_{t+1} = j_{t+1} | X_t = j_t)$$

para todo $t \in \mathbb{N}$ e $j_t \in S$.

Para não sobrecarregar a notação, quando nos referirmos a uma cadeia de Markov, assumimos também o espaço de probabilidades $(\Omega, \mathcal{O}, P)$, e usaremos a expressão “seja (X_t) uma cadeia de Markov com valores em S ” para tal.

Definição B.4 (Probabilidade de transição). Seja (X_t) uma cadeia de Markov com valores em S . Para todo $t \in \mathbb{N}$, a função

$$P_{t,t+1}(i, j) = P(X_{t+1} = j | X_t = i)$$

é chamada de probabilidade de transição.

Em palavras, $P_{t,t+1}(i, j)$ é a probabilidade de a cadeia, estando no estado i no tempo t , faça a transição para o estado j no tempo $t + 1$. Uma classe importante de cadeias de Markov são as cadeias com probabilidade de transição estacionária.

Definição B.5 (Transições estacionárias). Seja (X_t) uma cadeia de Markov com valores em S . Dizemos que a cadeia (X_n) tem *probabilidades de transição estacionárias* se, para todo $t \in \mathbb{N}$ e todos os estados i e j de S , temos

$$P(X_{t+1} = j | X_t = i) = P(X_1 = j | X_0 = i).$$

A definição acima nos diz que as probabilidades de transição de tal cadeia de Markov não variam com o tempo, o que também é conhecido como *probabilidades homogêneas no tempo*. A próxima definição vai nos permitir associar as probabilidades de transição estacionárias de uma cadeia de Markov com o conceito de matrizes estocásticas apresentada no início desta seção.

Definição B.6 (Matriz de transição). Seja (X_t) uma cadeia de Markov com probabilidades de transição estacionárias. A matriz $P = (P_{ij})$ definida por

$$P_{ij} = P(X_1 = j | X_0 = i)$$

é chamada de *matriz de transição* associada à cadeia de Markov (X_t) .

Das definições B.1 e B.6, vemos que as matrizes de transição de uma cadeia de Markov são matrizes estocásticas. De fato, temos $P_{ij} \geq 0$ para todos i e j , e também

$$\begin{aligned} \sum_{j \in S} P_{ij} &= \sum_{j \in S} P(X_1 = j | X_0 = i) = \frac{\sum_{j \in S} P(X_1 = j, X_0 = i)}{P(X_0 = i)} \\ &= \frac{P(\bigcup_{j \in S} \{X_1 = j, X_0 = i\})}{P(X_0 = i)} = \frac{P(X_0 = i)}{P(X_0 = i)} = 1, \end{aligned}$$

o que mostra que $P = (P_{ij})$ é uma matriz estocástica.

É comum representar cadeias de Markov com transições estacionárias através de um grafo, onde os nós são os estados de S , e as arestas indicam a probabilidade de transição de um nó para o seguinte. Assim, o exemplo do espaço de estados $S = \{A, B, C\}$ e a matriz dada na equação (B.1) pode ser ilustrado pela figura B.1.

A próxima definição conecta o conceito de vetor de distribuição de probabilidades, apresentado na definição B.2, com as cadeias de Markov.

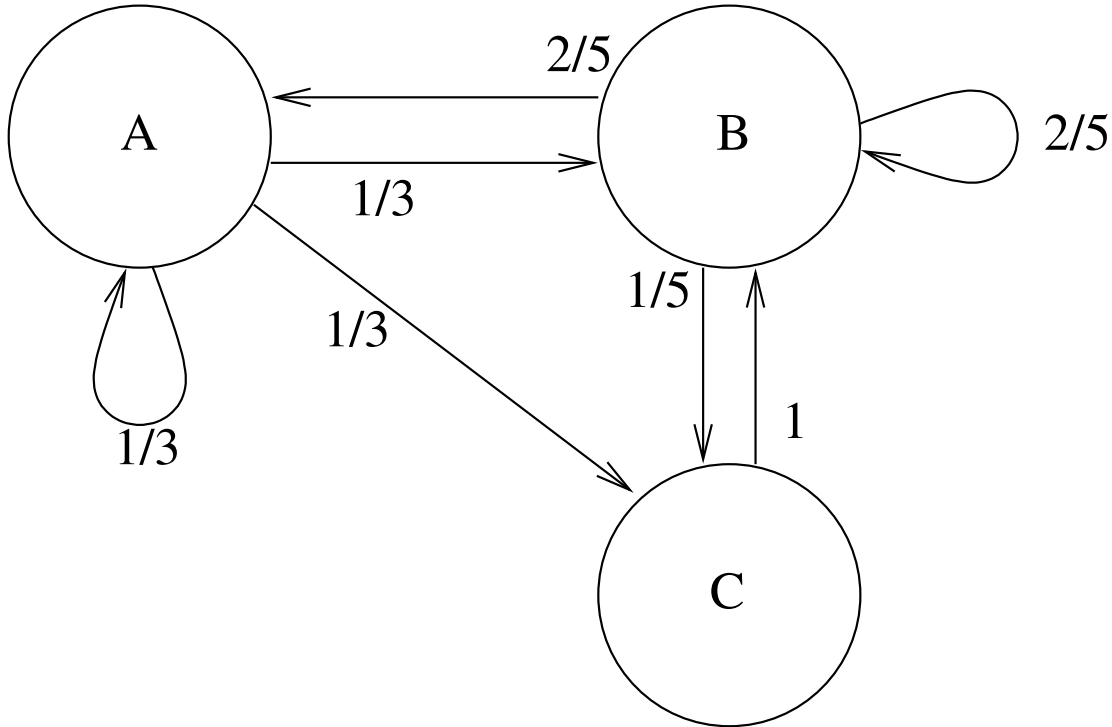


Figura B.1: Uma cadeia de Markov com 3 estados.

Definição B.7 (Probabilidades associadas a uma cadeia de Markov). Seja (X_t) uma cadeia de Markov com valores em S . Para todo $t \in \mathbb{N}$, denotamos por

$$\pi^t = (\pi^t(i))_{i \in S} = (P(X_t = i))_{i \in S}$$

o vetor de probabilidades sobre S que dá a probabilidade de a cadeia assumir o valor i no tempo t .

Para uma cadeia de Markov (X_t) com probabilidades de transição estacionárias, podemos escrever o vetor π^t da distribuição de probabilidades da cadeia de Markov X no tempo t em termos da t -ésima potência da matriz de transição P , e da distribuição da cadeia no tempo 0:

$$\pi^t = \pi^0 P^t. \quad (\text{B.2})$$

Com estas definições, podemos tornar mais precisa a última questão, relativa à *convergência* da cadeia de Markov. De fato, conforme $t \rightarrow \infty$, podemos investigar se o vetor de probabilidades π^t converge.

Antes de abordar esta questão, vamos tratar da existência e unicidade de uma distribuição de probabilidade invariante para a cadeia de Markov X .

B.2 Distribuição de probabilidade invariante para uma Cadeia de Markov

Começamos esta subseção com a definição de distribuição de probabilidades estacionária para uma cadeia de Markov.

Definição B.8 (Probabilidade estacionária). Seja (X_t) uma cadeia de Markov com valores em S . Um vetor de probabilidades $\pi = (\pi(i))_{i \in S}$ sobre o espaço de estados S é dito *estacionário* se

$$\pi \mathcal{P} = \pi,$$

em que $\mathcal{P}$ é a matriz de transição associada à cadeia de Markov (X_t) .

Note que se um vetor de probabilidades π é invariante pela matriz de transição $\mathcal{P}$, então ele também é invariante para qualquer potência desta matriz. De fato, para $t = 2$ temos

$$\pi \mathcal{P}^2 = (\pi \mathcal{P}) \mathcal{P} = \pi \mathcal{P} = \pi.$$

Por indução, temos que para qualquer $t > 0$ vale $\pi \mathcal{P}^t = \pi$.

A questão que se coloca agora é a seguinte: fixada uma matriz de transição $\mathcal{P}$, sempre existe um vetor de probabilidades estacionárias π ? A resposta como veremos depende de S ser finito ou não. Se S for finito, então sempre existe pelo menos um tal vetor invariante; já no caso de S ser infinito, nem sempre isto ocorre. Como veremos mais à frente, o passeio aleatório em $\mathbb{Z}$ com probabilidades de transição $1/2$ não possui vetor de probabilidade invariantes.

Os próximos resultados dependem apenas do espaço de estados e do fato de a matriz de transição $\mathcal{P}$ ser uma matriz estocástica, e por isso não faremos menção a uma cadeia de Markov (X_t) particular.

Definição B.9 (Simplexo de probabilidades). Seja S um conjunto enumerável. O *simplexo de probabilidades* sobre S é o conjunto

$$\Sigma_S := \left\{ (p_i)_{i \in S} \in [0, 1]^S \mid \sum_{i \in S} p_i = 1 \right\}.$$

Quando for claro pelo contexto, denotamos os elementos de Σ_S por p , e também por π .

Da definição B.7, vemos que as probabilidades π^t associadas a uma cadeia de Markov de fato pertencem ao simplexo Σ_S . Mais ainda, a evolução destas probabilidades pode ser vista como a aplicação sucessiva de matrizes de transição sobre estas probabilidades, o que, no caso de probabilidades estacionárias, é dado pela equação B.2. Vamos formalizar estas noções através do seguinte lema:

Lema B.10. *Seja $\mathcal{P}$ uma matriz de transição. A função $T : \Sigma_S \rightarrow \Sigma_S$ dada por*

$$T(p) = p\mathcal{P}$$

esta bem definida.

Demonstração. Considere e_i um vetor que tem todas as coordenadas iguais a zero, menos a i -ésima coordenada, igual a um. Temos que $e_i \in \Sigma_S$, e daí o vetor

$$T(e_i) = e_i \mathcal{P} = (\mathcal{P}_{ij})_{j \in S}.$$

Como $\mathcal{P}$ é estocástica, vemos que todas as coordenadas de $T(e_i)$ são positivas, e somam um, logo $T(e_i) \in \Sigma_S$.

Para o caso geral, note que $p \in \Sigma_S$ pode ser escrito como $p = \sum_{i \in S} p_i e_i$, com $\sum_{i \in S} p_i = 1$. Daí

$$T(p) = T\left(\sum_{i \in S} p_i e_i\right) = \sum_{i \in S} p_i T(e_i).$$

Como todos os p_i são positivos, vemos que todas as coordenadas de $T(p)$ serão positivas. Mais ainda, a soma das coordenadas será

$$\sum_{i \in S} p_i \cdot 1 = 1,$$

pois a soma das coordenadas de cada $T(e_i)$ é um. Isto implica que $T(p) \in \Sigma_S$, e segue-se que a aplicação T está bem definida. $\square$

Com a aplicação T em mãos, podemos estabelecer um primeiro resultado de existência do vetor de probabilidades estacionário.

Teorema B.11. *Se S é finito, então para toda matriz estocástica $\mathcal{P}$ existe ao menos um vetor π que é invariante por $\mathcal{P}$.*

Demonstração. Quando S tem cardinalidade d finita, o conjunto Σ_S é um subconjunto convexo, fechado e limitado de $\mathbb{R}^d$. Como a aplicação T definida no lema B.10 é contínua (pois é linear) segue-se pelo teorema do ponto fixo de Brouwer que existe um ponto fixo π para T , isto é, um vetor $\pi \in \Sigma_S$ tal que $T(\pi) = \pi$.

Da definição de T , concluímos que

$$\pi\mathcal{P} = \pi. \quad \square$$

O teorema acima afirma que, quando o espaço de estados é finito, então existe um vetor de probabilidades invariante por $\mathcal{P}$, que pode não ser único. Mais adiante mostraremos que, quando S é *infinito*, nem sempre existe um tal vetor de probabilidades invariante. O contraexemplo clássico é o passeio aleatório simétrico em $\mathbb{Z}$.

O próximo passo ainda neste contexto em que S finito é estabelecer as condições de unicidade para o vetor de probabilidades invariante. A partir daí, veremos como estender estes resultados para espaços de estados mais gerais. O primeiro passo nesta definição é analisar as matrizes de transição.

Definição B.12 (Matrizes regulares). Uma matriz estocástica $\mathcal{P}$ é dita *regular* se existe $k \in \mathbb{N}$ tal que a k -ésima potência de $\mathcal{P}$, ou seja, $\mathcal{P}^k$, tem todas as entradas positivas.

Esta definição desempenha uma papel importante na teoria das cadeias de Markov, e parte do nosso esforço posterior será em descrever condições mais simples que impliquem a regularidade de uma matriz estocástica. A importância da definição B.12 se deve ao seguinte resultado.

Teorema B.13 (Unicidade da distribuição estacionária para matrizes regulares, [LL, página 75]). *Se o espaço de estados S é finito e a matriz estocástica $\mathcal{P}$ é regular, então existe um único vetor $\pi \in \Sigma_S$ tal que $\pi\mathcal{P} = \pi$.*

Nestas condições, o vetor π tem todas as entradas estritamente positivas.

A propriedade de uma matriz estocástica ser *regular* garante que uma matriz de dimensão finita tenha um *único* vetor de probabilidades invariantes π . Combinando este teorema a algumas propriedades de álgebra linear podemos concluir a análise das cadeias de Markov respondendo à terceira pergunta inicial, a saber, sob quais condições a cadeia de Markov converge. Em equações, queremos saber quando, para todo $p \in \Sigma_S$, temos

$$\lim_{t \rightarrow +\infty} p\mathcal{P}^t = \pi.$$

Teorema B.14 (Convergência de cadeias de Markov em espaços de estado finitos, [LL, página 80]). *Se o espaço de estados S tem cardinalidade finita d e a matriz estocástica $\mathcal{P}$ é regular, então:*

1. $\mathcal{P}$ tem um único vetor de probabilidades invariante π , correspondente a um autovalor igual a 1, e as componentes de π são todas positivas;
2. os outros autovalores de $\mathcal{P}$ tem todos norma menor do que 1;
3. dado $p \in \Sigma_S$ a sequência $(p\mathcal{P}^t)_{t \in \mathbb{N}}$ converge para o ponto fixo π , isto é,

$$\lim_{t \rightarrow +\infty} p\mathcal{P}^t = \pi.$$

Como afirmamos antes da definição B.12, no caso do espaço de estados não ser finito o vetor de probabilidades estacionário pode não existir, como mostra o exemplo a seguir.

Exemplo B.15. O passeio aleatório simétrico em $\mathbb{Z}$ não possui distribuição estacionária.

Resolução. Neste caso, temos $S = \mathbb{Z}$, e as entradas da matriz de transição associadas ao passeio são dadas por

$$\mathcal{P}_{i,i-1} = \frac{1}{2}, \quad \mathcal{P}_{i,i+1} = \frac{1}{2} \quad \text{e} \quad \mathcal{P}_{ij} = 0 \quad \forall j \neq i+1, i-1.$$

Seja π tal que $\pi\mathcal{P} = \pi$. Daí

$$\pi(0) = \frac{1}{2}(\pi(-1) + \pi(1)),$$

e por sua vez

$$\pi(1) = \frac{1}{2}(\pi(0) + \pi(2)) \quad \text{e} \quad \pi(-1) = \frac{1}{2}(\pi(-2) + \pi(0)).$$

Substituindo,

$$\pi(0) = \frac{1}{2}(\pi(-2) + \pi(2)).$$

Por indução, também teremos que

$$\pi(0) = \frac{1}{2}(\pi(-k) + \pi(k))$$

para todo $k \in \mathbb{Z}$.

Ora, como $\sum_{n \in \mathbb{Z}} \pi(n) = 1$, os valores de $\pi(n)$ tendem a zero no infinito, mais precisamente:

$$\lim_{k \rightarrow +\infty} \pi(k) = 0 = \lim_{k \rightarrow -\infty} \pi(k).$$

Mas daí

$$\pi(0) = \lim_{k \rightarrow +\infty} \frac{\pi(-k) + \pi(k)}{2} = 0.$$

Como o mesmo argumento pode ser repetido para qualquer $k \in \mathbb{Z}$, teríamos que $\pi(k) = 0$, e não existe vetor de probabilidades estacionário. $\square$

Para determinar a existência de vetores de probabilidade estacionários para matrizes estocásticas mais gerais do que as regulares, vamos precisar dos conceitos de *irredutibilidade*, *recorrência* e *aperiodicidade*, que apresentaremos a seguir.

B.2.1 Cadeias de Markov irredutíveis

Intuitivamente, uma cadeia de Markov é *irredutível* se ela não pode ser decomposta em cadeias menores (com espaço de estados $S' \subset S$). Para formalizar esta decomposição, vamos definir duas relações entre os estados de uma cadeia de Markov com transições estacionárias.

Definição B.16 (Estados que se comunicam). Considere o espaço de estados S e a matriz estocástica $\mathcal{P}$.

1. Dados $i, j \in S$, se existe $t \in \mathbb{N}$ tal que $\mathcal{P}_{ij}^t > 0$, então dizemos que o estado i *conduz* ao estado j , e denotamos esta relação por $i \rightarrow j$.
2. Dados $i, j \in S$, se i conduz a j e j conduz a i , então dizemos que os estados i e j se *comunicam* e denotamos isso por $i \leftrightarrow j$.

A relação de comunicação é uma relação de equivalência sobre S , como mostra o lema abaixo.

Lema B.17. *A relação comunicação apresentada na definição B.16 é uma relação de equivalência.*

Demonstração. De fato temos que

1. $\mathcal{P}_{ii}^0 = 1 > 0$, logo $i \leftrightarrow i$.
2. A definição é simétrica: $i \leftrightarrow j$ se e somente se $j \leftrightarrow i$.
3. Se $i \leftrightarrow j$ e $j \leftrightarrow k$, então existem inteiros n e m tais que $\mathcal{P}_{ij}^n > 0$ e $\mathcal{P}_{jk}^m > 0$. Daí, como $\mathcal{P}$ tem todas as entradas positivas, vemos que

$$\mathcal{P}_{ik}^{n+m} = \sum_{r \in S} \mathcal{P}_{ir}^n \mathcal{P}_{rk}^m > \mathcal{P}_{ij}^n \mathcal{P}_{jk}^m > 0,$$

o que implica $i \rightarrow k$.

Para $k \rightarrow i$ raciocinamos de maneira análoga. □

O resultado acima implica que o espaço de estados S é particionado em classes de equivalência segundo a relação de comunicação. Com isto, podemos passar para a definição do conceito de irredutibilidade.

Definição B.18 (Matrizes de transição irredutíveis). Seja $\mathcal{P}$ uma matriz estocástica. Se a relação de comunicação $\leftrightarrow$ gera apenas uma classe de equivalência, dizemos que $\mathcal{P}$ é *irredutível*.

Quando a matriz $\mathcal{P}$ é regular, temos apenas uma classe de equivalência, e portanto $\mathcal{P}$ é irredutível. Além disso, se existe alguma potência k tal que $\mathcal{P}^k$ tenha todas as entradas positivas, então também vemos que $\mathcal{P}$ será regular. Usando novamente a matriz 3×3 que demos como exemplo, vemos que

$$\mathcal{P}^2 = \begin{bmatrix} 11/45 & 26/45 & 8/45 \\ 22/75 & 37/75 & 16/75 \\ 2/5 & 2/5 & 1/5 \end{bmatrix}$$

e portanto a cadeia de Markov definida por esta matriz é irredutível.

A figura B.2 dá um exemplo de cadeia de Markov redutível: nela, os estados B , C e D se comunicam todos, mas o estado A não pode ser alcançado a partir de nenhum deles, e portanto forma uma classe de equivalência separada.

B.2.2 Cadeias de Markov recorrentes e transientes

Para podermos apresentar o conceito de recorrência, e o complementar, de transiência, vamos precisar introduzir alguns novos objetos.

Definição B.19 (Tempo de primeira passagem). Seja (X_t) uma cadeia de Markov com valores em S . Dado $A \subseteq S$, o *tempo de primeira passagem por A* é definido por

$$T_A(\omega) := \min\{t \geq 1 : X_t(\omega) \in A\},$$

com a convenção que o mínimo de um conjunto vazio é $+\infty$.

Note que T_A é uma variável aleatória estendida, ou seja, pode assumir o valor $+\infty$.

Estamos particularmente interessados no caso $A = \{i\}$, pois a partir deste caso temos a seguinte definição.

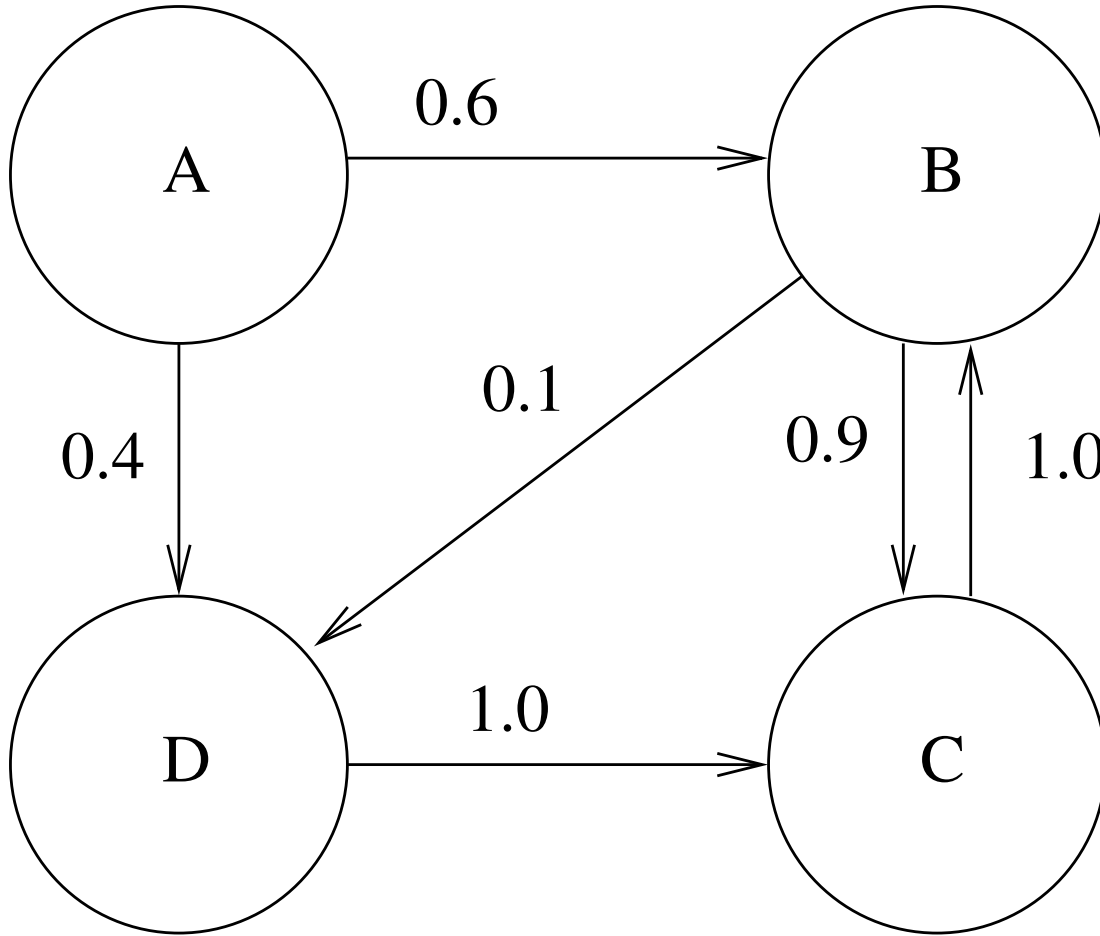


Figura B.2: Uma cadeia de Markov redutível com 4 estados.

Definição B.20 (Probabilidades de passagem). Seja (X_t) uma cadeia de Markov com espaço de estados S . Para cada $t \in \mathbb{N}$, a função $f^t : S \times S \rightarrow [0, 1]$ é a probabilidade de atingir o estado j pela primeira vez no tempo t dado que a cadeia estava no estado i no tempo zero. Mais precisamente,

$$f^t(i, j) = P \left(\{X_t = j\} \cap \left(\bigcap_{r=1}^{t-1} \{X_r \neq j\} \right) \middle| X_0 = i \right).$$

É comum denotarmos $f^t(i, j)$ por f_{ij}^t , por analogia com as matrizes de transição que também são indexadas por ij .

A partir das funções f^t , construímos as seguintes funções para $i, j \in S$.

$$f_{ii}^* := \sum_{t=1}^{+\infty} f_{ii}^t \quad \text{e} \quad f_{ij}^* := \sum_{t=1}^{+\infty} f_{ij}^t. \quad (\text{B.3})$$

Tanto f_{ii}^* quanto f_{ij}^* estão bem definidas, pois

$$f_{ii}^* = P(T_i < +\infty | X_0 = i) \quad \text{e, analogamente,} \quad f_{ij}^* = P(T_j < +\infty | X_0 = i).$$

Em palavras, f_{ii}^* é a probabilidade de *retornarmos* ao estado i em um tempo finito, dado que começamos em i , e analogamente f_{ij}^* é a probabilidade de atingirmos o estado j em um tempo finito dado que começamos em i .

Definição B.21 (Estados recorrentes e transientes). Seja (X_t) uma cadeia de Markov com espaço de estados S . Dizemos que o estado $i \in S$ é *recorrente* se $f_{ii}^* = 1$, e *transiente* caso contrário.

Ou seja, um estado é recorrente se, com probabilidade 1, retornamos a ele em um tempo finito, dado que começamos neste estado.

Apesar de a definição B.21 ter um grande apelo intuitivo, ela não é simples de trabalhar do ponto de vista computacional. Devemos buscar formas mais simples de caracterizar as propriedades de recorrência ou transiência de um estado $i \in S$.

Definição B.22 (Número de visitas). Seja (X_t) uma cadeia de Markov com espaço de estados S . A variável aleatória $N(i) : \Omega \rightarrow \mathbb{N} \cup \{+\infty\}$ definida por

$$N(i) := \sum_{t=1}^{+\infty} \mathbb{1}_{\{i\}}(X_t) \quad (\text{B.4})$$

em que

$$\mathbb{1}_{\{i\}}(X_t) := \begin{cases} 1 & \text{se } X_t = i, \\ 0 & \text{caso contrário.} \end{cases}$$

é o *número de visitas* que a cadeia de Markov X fez ao estado i .

Lema B.23. *Seja (X_t) uma cadeia de Markov com espaço de estados S e matriz de transição $\mathcal{P}$. Então para todo estado $i \in S$,*

$$\mathbb{E}[N(j)|X_0 = i] = \sum_{t=1}^{+\infty} \mathcal{P}_{ij}^t,$$

sempre que $P(N(j) = +\infty|X_0 = i) = 0$.

Demonstração. Inicialmente, note que a sequência $(s(X)_m)_{m \geq 0}$ definida por

$$s(X)_m := \sum_{t=1}^m \mathbb{1}_{\{j\}}(X_t)$$

é não decrescente e converge para $\sum_{t=1}^{+\infty} \mathbb{1}_{\{j\}}(X_t)$.

Usando $\mathbb{E}_i$ para indicar a esperança condicional a $X_0 = i$, temos pelo teorema da convergência monótona que

$$\begin{aligned} \mathbb{E}_i[N(j)] &= \mathbb{E}_i \left[\sum_{t=1}^{+\infty} \mathbb{1}_{\{j\}}(X_t) \right] = \mathbb{E}_i \left[\lim_{m \rightarrow +\infty} \sum_{t=1}^m \mathbb{1}_{\{j\}}(X_t) \right] \\ &= \lim_{m \rightarrow +\infty} \sum_{t=1}^m \mathbb{E}_i[\mathbb{1}_{\{j\}}(X_t)] = \lim_{m \rightarrow +\infty} \sum_{t=1}^m P(X_t = j|X_0 = i) = \sum_{t=1}^{+\infty} \mathcal{P}_{ij}^t. \quad \square \end{aligned}$$

De forma intuitiva, podemos esperar que, se o número médio de visitas a um determinado estado é arbitrariamente grande, então ele deveria ser recorrente. Os próximos resultados vão justamente confirmar esta intuição.

Teorema B.24 (Estados recorrentes, [LL, página 110]). *Seja (X_t) uma cadeia de Markov com espaço de estados S e matriz de transição $\mathcal{P}$. O estado $i \in S$ é recorrente se, e somente se,*

$$\sum_{t=1}^{+\infty} \mathcal{P}_{ij}^t = +\infty.$$

Uma consequência direta do resultado acima é a seguinte.

Corolário B.25. *Seja (X_t) uma cadeia de Markov com espaço de estados S . O estado $i \in S$ é recorrente se, e somente se, $\mathbb{E}_i[N(i)] = +\infty$.*

A propriedade de um estado ser recorrente não teria tanta importância na análise de cadeias de Markov se esta propriedade fosse apenas do estado em si, e não da *classe de comunicação* à qual este estado pertence. O próximo resultado nos diz que esta característica é de fato uma propriedade de classe.

Teorema B.26 (Recorrência das classes de comunicação, [LL, página 112]). *Seja (X_t) uma cadeia de Markov com espaço de estados S . Dados dois estados $i, j \in S$ que se comunicam, então*

1. *i é recorrente se, e somente se j é recorrente.*
2. *i é transiente se, e somente se j é transiente.*

Até o momento, não vimos a conexão entre a existência de uma distribuição de probabilidades invariante e a propriedade de recorrência. Já sabemos que esta é uma propriedade de classe segundo a relação de comunicação, e a propriedade de uma cadeia ser irredutível nos dá justamente a existência de uma única classe de comunicação. O próximo resultado nos diz que ser recorrente é uma consequência da cadeia de Markov ser irredutível e ter uma distribuição invariante.

Teorema B.27 ([LL, página 117]). *Seja (X_t) uma cadeia de Markov irredutível com espaço de estados S . Se a cadeia (X_t) possui uma distribuição invariante $\pi \in \Sigma_S$, então a cadeia é recorrente.*

Teorema B.28. *Seja (X_t) uma cadeia de Markov com espaço de estados S e matriz de transição $\mathcal{P}$. Se o estado $j \in S$ for transiente, então*

$$\lim_{t \rightarrow +\infty} \mathcal{P}_{i,j}^t = 0$$

para todo $i \in S$ fixo.

Demonstração. Se o estado $j \in S$ é transiente então

$$\mathbb{E}_i[N(j)] = \sum_{t=1}^{+\infty} \mathcal{P}_{ij}^t < +\infty$$

Logo a série acima é absolutamente convergente e segue-se que

$$\lim_{t \rightarrow +\infty} \mathcal{P}_{i,j}^t = 0. \quad \square$$

O conceito de recorrência se subdivide em dois subconceitos, a saber: recorrente nulo e recorrente positivo. Para estudá-los, precisamos da noção de *aperiodicidade* que passamos a expor na próximo subseção.

B.2.3 Cadeias de Markov periódicas e aperiódicas

Começamos esta subseção definindo o que entendemos por período de um estado.

Definição B.29 (Período de um estado). *Seja (X_t) uma cadeia de Markov com espaço de estados S e matriz de transição $\mathcal{P}$. O período de um estado $i \in S$ é definido por*

$$d(i) = \begin{cases} \text{mdc}\{t \geq 1 \mid \mathcal{P}_{ii}^t > 0\} & \text{se existir um } t \geq 1 \text{ tal que } \mathcal{P}_{ii}^t > 0, \\ 0 & \text{caso contrário} \end{cases}$$

Ou seja, o período de um estado $i \in S$ é o máximo divisor comum dos tempos $t \geq 1$ tais que a cadeia começando em i retorna a i em t passos.

Definição B.30. Seja (X_t) uma cadeia de Markov com espaço de estados S . Dizemos que um estado $i \in S$ é *aperiódico* se $d(i) = 1$ e que é *periódico* se $d(i) > 1$.

Assim como no caso do conceito de recorrência, o conceito de aperiodicidade também é uma propriedade de classe segundo a relação de comunicação.

Lema B.31 (O período é uma propriedade de classe, [LL, página 134]). *Seja (X_t) uma cadeia de Markov com espaço de estados S . Se os estados $i, j \in S$ são tais que $i \leftrightarrow j$, então $d(i) = d(j)$.*

O teorema a seguir é um dos principais resultados para cadeias de Markov a espaço de estados discretos. Ele nos permite, em conjunto com o conceito de irreducibilidade, garantir que a matriz de transição de uma cadeia de Markov com espaço de estados finito é regular, o que nos dará, novamente, a unicidade da distribuição de probabilidade invariante.

Teorema B.32 ([LL, página 135]). *Seja (X_t) uma cadeia de Markov com espaço de estados finito S e matriz de transição $\mathcal{P}$ irreducível. Considere o estado $i \in S$ e seja $d = d(i)$ o seu período. Então existe $t_i > 0$ tal que para todo $t \geq t_i$ temos que $\mathcal{P}_{ii}^{td} > 0$. Além disto, existe $N > 0$ tal que, para todo $m \geq N$, $\mathcal{P}_{jj}^{md} > 0$ para qualquer estado $j \in S$.*

Deste teorema, podemos deduzir a conexão que afirmarmos no parágrafo anterior entre os conceitos de irreducibilidade e aperiodicidade com o de regularidade.

Corolário B.33. *Seja (X_t) uma cadeia de Markov com espaço de estados finito S e matriz de transição $\mathcal{P}$ irreducível e aperiódica. Então existe $t_0 \geq 1$ tal que se $t \geq t_0$, todas as entradas da matriz $\mathcal{P}^t$ são positivas.*

Em particular, a matriz de transição $\mathcal{P}$ é regular.

Os conceitos de irreducibilidade e aperiodicidade nos permitem estudar distribuições de probabilidade invariantes para espaços de estados mais gerais que finitos, e veremos um primeiro exemplo para espaços de estados enumeráveis.

B.2.4 Recorrência nula e positiva

Os últimos conceitos necessários para responder as três perguntas da introdução são os de recorrência nula e positiva. Eles nos permitirão estudar estas questões fundamentais em espaços de estados mais gerais, como por exemplo infinito enumeráveis.

Definição B.34. Seja (X_t) uma cadeia de Markov com espaço de estados S e matriz de transição $\mathcal{P}$. Considere $T_i : \Omega \rightarrow \mathbb{N}$ o tempo de primeiro retorno ao estado $i \in S$. Definimos o tempo médio de primeiro retorno ao estado $i \in S$ por

$$\mu_i = \mathbb{E}_i[T_i] = \sum_{t=1}^{+\infty} t f_{ii}^t$$

em que f_{ii}^t é a probabilidade de primeiro retorno ao estado i dado que a cadeia começou em i .

Com a definição do tempo médio de primeiro retorno, podemos definir os conceitos de recorrência nula e positiva.

Definição B.35 (Recorrência nula e positiva). Seja (X_t) uma cadeia de Markov com espaço de estados S . Se $i \in S$ é um estado recorrente, então

1. i é recorrente *positivo* quando $\mu_i < +\infty$.
2. i é recorrente *nulo* quando $\mu_i = +\infty$.

Da mesma forma que anteriormente, esta é uma propriedade de classe:

Lema B.36 ([LL, página 148]). *Seja (X_t) uma cadeia de Markov. Um estado i é recorrente nulo (resp. positivo) se e somente se j é recorrente nulo (resp. positivo) para todo estado j que se comunica com i .*

Podemos então passar para o principal resultado para cadeias de Markov com espaço de estados enumeráveis:

Teorema B.37 (Recorrência e distribuição invariante em espaços de estados enumeráveis [LL, página 168]). *Seja (X_t) uma cadeia de Markov com espaço de estados S e matriz de transição $\mathcal{P}$. Suponha que a cadeia seja irredutível, aperiódica e recorrente. Então*

1. $\lim_{t \rightarrow +\infty} \mathcal{P}_{ii}^t = 0$ se $\mathbb{E}_i[T_i] = +\infty$ (ou seja é recorrente nula).
2. $\lim_{t \rightarrow +\infty} \mathcal{P}_{ii}^t = \frac{1}{\mathbb{E}_i[T_i]} = \pi_i$ se $\mathbb{E}_i[T_i] < +\infty$ (ou seja é recorrente positiva).

O teorema acima nos diz que, para um espaço de estados enumerável, uma cadeia de Markov a tempo discreto possui uma distribuição estacionária quando for irredutível, aperiódica e recorrente positiva.

B.3 Cadeias de Markov com espaço de estados contínuo

Nesta seção vamos tratar do caso de espaço de estados contínuos e tempo discreto. Assim como na seção B.1, onde começamos definindo o conceito de matriz estocástica para depois introduzirmos o conceito de probabilidade de transição, vamos começar esta seção introduzindo o conceito análogo ao de matriz de transição que será fundamental para o caso contínuo.

No que segue S é um espaço de estados contínuo e $\mathcal{B}(S)$ a σ -álgebra definida sobre este espaço.

Definição B.38. Seja $(S, \mathcal{B}(S))$ um espaço mensurável. Chamamos de núcleo de transição a função $K : S \times \mathcal{B}(S) \rightarrow [0, 1]$ se esta possui as seguintes propriedades.

- 1 Para todo $x \in S$ temos que $K(x, \cdot)$ é uma medida de probabilidade.
- 2 Para todo $A \in \mathcal{B}(S)$ temos que $K(\cdot, A)$ é uma função mensurável.

Olhando para o item 1, podemos notar a semelhança desta definição com a definição B.1 de matriz estocástica. Para um espaço de estados discreto S , tínhamos $\mathcal{P}_{ij} \geq 0$ para todos i e j de S , e também que, fixado $i \in S$, ocorria que $\sum_{j \in S} \mathcal{P}_{ij} = 1$. Agora, considerando $S \times \mathcal{B}(S)$ como novo “espaço de estados”, temos pelo item 1 que, para todo par $(x, A) \in S \times \mathcal{B}(S)$, o núcleo de transição $K(x, A)$ é positivo, e também fixado $x \in S$ temos $K(x, S) = 1$, que é a soma sobre todos os pontos do espaço de estados contínuo.

Observação B.39. Quando ocorrer que fixado $x \in S$ a medida de probabilidade $K(x, \cdot)$ for absolutamente contínua em relação a medida de Lebesgue vamos escrever $K(x, y)$ representando a densidade de probabilidade associada ao núcleo de transição.

Observação B.40. Quando o espaço de estados é contínuo chamamos as cadeias de Markov de processos de Markov.

Definição B.41. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo estocástico tomando valores no espaço de estados contínuo S . Dizemos que X é um processo de Markov se satisfaz a seguinte propriedade.

$$P(X_{t+1} \in A | x_t, \dots, x_0) = P(X_{t+1} \in A | x_t) = \int_A K(x_t, dx_{t+1}) \quad (\text{B.5})$$

Quando o núcleo de transição for absolutamente contínuo em relação à medida de Lebesgue escrevemos a última igualdade na equação B.5 do seguinte modo

$$P(X_{t+1} \in A | x_t) = \int_A K(x_t, x_{t+1}) dx_{t+1}$$

em que agora K representa uma densidade de probabilidade.

Definição B.42. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Dizemos que o processo de Markov X é homogêneo se para toda a sequência $(t_i)_{0 \leq i \leq k}$ de pontos em $\mathbb{N}$ não decrescentes temos que

$$(X_{t_1}, \dots, X_{t_k}) \sim (X_{t_1-t_0}, \dots, X_{t_k-t_0})$$

Assim como no caso discreto vamos considerar em nossa análise apenas processos de Markov homogêneo.

Exemplo B.43. Um caso simples de processos de Markov em espaço de estados contínuo é o modelo AR(1) comumente utilizado na análise de séries temporais. Dado $\theta \in \mathbb{R}$ e uma sequência de variáveis aleatórias $(\epsilon_t)_{t \in \mathbb{N}}$ com distribuição $\epsilon_t \sim N(0, \sigma^2)$ definimos o processo de Markov $X^\theta = (X_t^\theta)_{t \in \mathbb{N}}$ tomando valores em $\mathbb{R}$ por

$$X_t^\theta = \theta X_{t-1}^\theta + \epsilon_t.$$

Note que o núcleo de transição fica dado por $X_t^\theta | x_{t-1} \sim N(\theta x_{t-1}, \sigma^2)$ que neste caso tem densidade normal.

No caso discreto para introduzirmos o primeiro conceito necessário ao estudo da existência e unicidade da distribuição invariante de uma cadeia de Markov, a saber, o conceito de Regularidade, tivemos que definir o conceito de potência de uma matriz. Pois a regularidade da cadeia consistia na existência de $k \in \mathbb{N}$ tal que todas as entradas da matriz de transição $\mathcal{P}$ eram positivas.

Já sabemos que no caso de espaço de estados contínuo o conceito de matriz de estocástica (transição no caso de cadeias de Markov) é substituído pelo conceito de núcleos de transição. Neste último caso o conceito de potência de uma matriz é substituído pelo de convolução de medidas.

Definição B.44. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Considere $K : S \times \mathcal{B}(S) \rightarrow [0, 1]$ o núcleo de transição associado ao processo X . Dado o par $(x, A) \in S \times \mathcal{B}(S)$ e denotando $K^1(x, A) = K(x, A)$ definimos o núcleo de transição para a t -ésima transição por

$$K^t(x, A) = \int_S K^{t-1}(y, A) K(x, dy)$$

B.3.1 Cadeias de Markov irredutíveis

A definição do conceito de irredutibilidade para cadeias de Markov com espaço de estados discreto estava assentado sobre a ideia de que o espaço de estados podia ser particionado em classes de equivalência segundo a relação de comunicação " $\leftrightarrow$ ". Embora podemos olhar para os núcleos de transição $K(x, A) > 0$ para decidir se um determinado conjunto A é atingido pelo processo de Markov dado que este começou no ponto x temos que esta seria apenas a noção de condução " $\rightarrow$ " e teríamos dificuldades para definir o que poderia ser a condução de " A " para " x ". Por isso vamos precisar de uma noção de irredutibilidade que não dependa do conceito de classe de comunicação.

Definição B.45. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Considere $K : S \times \mathcal{B}(S) \mapsto [0, 1]$ o núcleo de transição associado ao processo de Markov X . Dado uma medida $\varphi : \mathcal{B}(S) \mapsto [0, 1]$ dizemos que o processo de Markov X é φ -irredutível se ocorre que para todo $A \in \mathcal{B}(S)$ com $\varphi(A) > 0$ existe $t \in \mathbb{N}$ tal que $K^t(x, A) > 0$ para todo $x \in S$.

A definição acima garante que conjuntos suficientemente grandes segundo a medida φ são sempre visitados pelo processo com probabilidade positiva não importando o ponto em que o processo começou.

Relembre que tempos de parada e número de visitas a um determinado estado desempenhavam um papel importante no contexto discreto. No caso contínuo isto não será diferente.

Definição B.46. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Dado $A \in \mathcal{B}(S)$, dizemos que a variável aleatória $\tau_A : \Omega \mapsto \mathbb{N}$ definida por

$$\tau_A := \min\{t \geq 1 : X_t(\omega) \in A\}$$

é um tempo de parada.

Em palavras é o primeiro tempo no qual o processo entra no conjunto A . Apesar de não mencionado na definição acima estamos implicitamente considerando que $\tau_A = +\infty$ se $X_t \notin A$ para todo $t \in \mathbb{N}$.

Definição B.47. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Dado $A \in \mathcal{B}(S)$, dizemos que a variável aleatória $n_A : \Omega \mapsto \mathbb{N}$ definida por

$$n_A := \sum_{t=1}^{+\infty} \mathbf{1}_A(X_t)$$

é o número de visitas ao estado n_A .

Associado as definições B.46 e B.47 temos as quantidades

$$\mathbb{E}_x[n_A] \quad \text{e} \quad P_x(\tau_A < +\infty)$$

que são respectivamente o número médio de visitas ao conjunto A dado que o processo começou no ponto x e a probabilidade de se visitar o conjunto A pela primeira vez em um tempo finito dado que o processo começou no ponto x .

Analogamente ao caso discreto sempre que definimos um conceito que representa uma nova propriedade do processo devemos mostrar que esta propriedade é intrínseca ao processo. Olhando

para a definição B.45 temos a impressão que a característica do processo ser irredutível depende da medida auxiliar φ . O próximo teorema nos dirá que este não é o caso. Mas antes de enunciá-lo vamos precisar de um conceito auxiliar que nos permitirá obter subprocessos de processos de Markov que desfrutem de propriedades mais "amigáveis" do que o processo original.

Definição B.48. Sejam $(S, \mathcal{B}(S))$ um espaço mensurável e $K : S \times \mathcal{B}(S) \mapsto [0, 1]$ um núcleo de transição. Fixado $\epsilon > 0$ dizemos que o núcleo de transição $K_\epsilon : S \times \mathcal{B}(S) \mapsto [0, 1]$ é um resolvente se possui a seguinte forma.

$$K_\epsilon(x, A) := (1 - \epsilon) \sum_{i=0}^{+\infty} \epsilon^i K^i(x, A), \quad 0 < \epsilon < 1$$

e o processo com este núcleo é chamado de K_ϵ -processo.

Teorema B.49. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Se o processo de Markov X é φ -irredutível então existe uma medida de probabilidade ψ tal que

- 1 O processo de Markov X é ψ -irredutível.
- 2 Se existe uma medida ρ tal que X é ρ -irredutível, então ρ é dominada por ψ , ou seja, $\rho \ll \psi$.
- 3 Se $\psi(A) = 0$, então $\psi\left(\{y : P_y(\tau_A < +\infty) > 0\}\right) = 0$.
- 4 A medida ψ é equivalente a

$$\psi_0(A) := \int_S K_{\frac{1}{2}}(x, A) \varphi(dx) \quad \forall A \in \mathcal{B}(S).$$

Ou seja, $\psi_0 \ll \psi$ e $\psi \ll \psi_0$.

Demonstração. Ver referência [MT09] página 83. □

Ou seja a propriedade de ser φ -irredutível não depende de uma particular medida φ bastando que exista um tal medida para que o processo de Markov seja irredutível.

Exemplo B.50. (Continuação do exemplo B.43). Como para o modelo AR(1) os núcleos de transição são dados por normais com média θx_{t-1} e variância σ^2 temos que considerando φ a medida de Lebesgue podemos escrever o núcleo K como

$$K(x_{t-1}, A) = \int_A \frac{1}{\sqrt{2\pi}\sigma} e^{-\frac{(\theta x_{t-1} - u)^2}{2\sigma^2}} \varphi(du)$$

Dado que a densidade da normal é sempre positiva para todo $x_{t-1} \in S$ temos que $K(x_{t-1}, A) > 0$ somente se $\varphi(A) > 0$. Portanto o processo de Markov $X^\theta = (X_t^\theta)_{t \in \mathbb{N}}$ do modelo AR(1) é φ -irredutível.

B.3.2 Átomos e conjunto pequenos

A próxima definição apresenta um conceito que apesar de pouco intuitivo é fundamental para podermos transportar resultados obtidos das cadeias de Markov (espaço de estados discretos) para os processos de Markov (espaço de estados contínuo).

Definição B.51. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Considere $K : S \times \mathcal{B}(S) \rightarrow [0, 1]$ o núcleo de transição associado ao processo de Markov X . Dizemos que um evento $\alpha \in \mathcal{B}(S)$ é um átomo se existe uma medida de probabilidade $\nu > 0$ tal que

$$K(x, A) = \nu(A) \quad \forall x \in \alpha \quad \text{e} \quad \forall A \in \mathcal{B}(S).$$

Além disso se X é φ -irredutível então dizemos que o átomo α é acessível se $\varphi(\alpha) > 0$.

Em palavras a definição acima diz que um átomo é um subconjunto mensurável de S com a seguinte propriedade: quando o processo começa em um dos seus pontos a probabilidade do processo visitar a partir daí qualquer subconjunto mensurável do espaço de estados passa a não depender do ponto onde o processo começou.

Esse conceito é simples de visualizar no caso discreto. De fato, considere S um espaço de estados discreto. Olhando para a matriz de transição $\mathcal{P}$ considere a linha i tal que suas entradas são positivas. Neste caso temos que $\alpha = \{i\} \subseteq S$ é um átomo, pois é mensurável em relação a σ -álgebra 2^S (conjunto das partes de S) e sabemos que

$$\sum_{j \in S} \mathcal{P}_{ij} = \sum_{j \in S} \nu(j) = 1$$

bastando tomar $\nu(j) = \mathcal{P}_{ij}$.

Percebemos também que este conceito é muito forte para se utilizado na maioria dos casos de interesse. Por adotamos uma relaxação da definição B.51 conhecida como condição *minorizing*.

Definição B.52. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Considere $K : S \times \mathcal{B}(S) \rightarrow [0, 1]$ o núcleo de transição associado ao processo de Markov X . Dizemos que a condição *minorizing* é atendida quando existem um evento $C \in \mathcal{B}(S)$, uma constante $\epsilon > 0$ e uma medida de probabilidade $\nu > 0$ tal que

$$K(x, A) \geq \epsilon \nu(A) \quad \forall x \in C \quad \text{e} \quad \forall A \in \mathcal{B}(S).$$

A condição *minorizing* por sua vez leva a uma última simplificação conhecida como conjuntos pequenos.

Definição B.53. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Considere $K : S \times \mathcal{B}(S) \rightarrow [0, 1]$ o núcleo de transição associado ao processo de Markov X . Um conjunto $C \in \mathcal{B}(S)$ é dito pequeno se existe $m \in \mathbb{N}$ e uma medida de probabilidade $\nu_m > 0$ tal que

$$K^m(x, A) \geq \nu_m(A) \quad \forall x \in C \quad \text{e} \quad \forall A \in \mathcal{B}(S).$$

Observação B.54. É possível calcular os conjuntos pequenos para o exemplo B.43 (modelo AR(1)). Não faremos aqui em função de demandar um certo trabalho algébrico que não tornaria a compreensão sobre este tipo de conjunto mais clara.

Até agora não dissemos como os conceitos de conjuntos pequenos e de φ -irredutibilidade se conectam. Isso é feito no próximo resultado.

Teorema B.55. *Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov φ -irredutível tomando valores no espaço de estados contínuo S . Para todo $A \in \mathcal{B}(S)$ tal que $\varphi(A) > 0$ existem $m \in \mathbb{N}$, uma medida de probabilidade $\nu_m > 0$ e um conjunto pequeno $C \subset A$ tal que $\nu_m(C) > 0$.*

Demonstração. Ver referência [MT09] página 105. □

B.3.3 Ciclos e períodos

Na subseção B.2.3 definimos no contexto de espaço de estados discreto o conceito de periodicidade. Como sabemos uma cadeia de Markov pode ter um comportamento periódico ou aperiódico e como esta é uma propriedade de classe (com respeito a relação de comunicação $\leftrightarrow$) tínhamos que este conceito se ligava ao de irredutibilidade quando a cadeia de Markov em questão apresentava uma única classe de comunicação. O equivalente a este conceito no caso de espaço de estados contínuos é o de ciclos que discutiremos brevemente nesta subseção.

Definição B.56. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov φ -irredutível tomando valores no espaço de estados contínuo S . Dizemos que o processo X tem um ciclo de tamanho $d \geq 1$ se existem um conjunto pequeno $C \in \mathcal{B}(S)$, um inteiro $M \in \mathbb{N}$ e uma medida ν_M tal que

$$d := \text{mdc}\{m \geq 1 : \exists \delta_m > 0 \text{ tal que } C \text{ é pequeno para } \nu_m \geq \delta_m \nu_M\}$$

Pode ser mostrado que o período d independe do conjunto pequeno C e portanto é uma propriedade intrínseca do processo X . Apesar de usarmos o nome ciclo também diremos que o processo é aperiódico quando $d = 1$.

Observação B.57. Se existe um conjunto pequeno $C \in \mathcal{B}(S)$, então pela definição B.53 existem um inteiro $m \in \mathbb{N}$ e uma medida de probabilidade $\nu_m > 0$ tal que

$$K^m(x, A) \geq \nu_m(A) \quad \forall x \in C \quad \text{e} \quad \forall A \in \mathcal{B}(S).$$

se ocorre que $m = 1$ então $d = 1$ e o processo de Markov X é dito ser fortemente aperiódico.

Observação B.58. Como estamos considerando um processo φ -irredutível (para podermos falar em período) temos pelo teorema B.55 que $C \subset A$. Então a observação anterior nos diz que quando um processo é fortemente aperiódico podemos ir A para A em um "passo". Esse fato será fundamental na obtenção da distribuição de probabilidade do processo de Markov.

B.3.4 Transiência a recorrência

O conceito de transiência para processo de Markov é o que apresentar maior semelhança com sua versão discreta. Relembre que na definição B.47 apresentamos a seguinte variável aleatória

$$n_A := \sum_{t=1}^{+\infty} \mathbb{1}_A(X_t)$$

que representa o número de visitas a um evento A pelo processo X dado que este processo começou no ponto x . No comentário após a definição B.47 também apresentamos a quantidade $\mathbb{E}_x[n_A]$ que representa o número médio de visitas ao conjunto A pelo processo X dado que o processo começou no ponto x . Estas ideias são importante para a próxima definição.

Definição B.59. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Um evento $A \in \mathcal{B}(S)$ é dito ser recorrente se $\mathbb{E}_x[n_A] = +\infty$ para todo $x \in A$. O mesmo evento é dito ser uniformemente transiente se existe uma constante $M > 0$ tal que $\mathbb{E}_x[n_A] < M$ para todo $x \in A$. Por fim A é dito ser transiente se existe uma cobertura para este conjunto por conjuntos uniformemente transientes, ou seja,

$$A \subset \bigcup_{i=0}^{+\infty} B_i$$

onde os B_i são conjunto uniformemente transientes.

Notemos a semelhança desta definição com a afirmação apresentada no corolário B.25 que nos diz que um estado i é recorrente se, e somente se $\mathbb{E}_i[n_{\{i\}}] = +\infty$. Além disso o teorema B.26 apresentado logo após o corolário B.25 nos diz que no caso discreto a propriedade de ser recorrente é uma propriedade de classe e não dependia de um estado em particular. Este resultado é transportado para o caso contínuo através do conceito de conjunto atômico apresentado na introdução da subseção B.3.2.

Teorema B.60. *Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov φ -irredutível com átomo acessível α e tomando valores no espaço de estados contínuo S . Então as seguintes afirmações valem.*

- 1 Se α for recorrente, então todo conjunto $A \in \mathcal{B}(S)$ tal que $\varphi(A) > 0$ é recorrente.
- 2 Se α for transiente, então S é transiente.

Demonstração. Ver referência [MT09] página 179. □

Observe que α desempenha um papel semelhante ao da relação de equivalência $\leftrightarrow$. Na próxima definição estendemos a definição de recorrência de eventos para o processo como um todo.

Definição B.61. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Um processo de Markov é recorrente se ocorre que

- 1 Se existe uma medida φ tal que X é φ -irredutível.
- 2 Para todo $A \in \mathcal{B}(S)$ tal que $\psi(A) > 0$ temos $\mathbb{E}_x[n_A] = +\infty$.

E dizemos que este processo é transiente se ele é φ -irredutível e S um conjunto transiente.

No caso de espaço de estados discreto provamos que as propriedades recorrente e transiente formam uma dicotomia. No caso de espaço de estados contínuo essa característica se mantém como indica o seguinte resultado

Teorema B.62. *Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov φ -irredutível tomando valores no espaço de estados contínuo S . Então o processo X é recorrente ou transiente.*

Demonstração. Ver referência [MT09] página 181. □

B.3.5 Recorrência de Harris

As classificações do conceito de recorrência vão se sofisticando conforme a complexidade do espaço de estados. Vimos isso na passagem do espaço de estados discreto finito, quando precisamos apenas do critério de recorrência dado na definição B.21, para o espaço de estados discreto enumerável onde foi necessário refinarmos a definição de recorrência em B.59. Agora no tratamento de espaço de estados contínuos será necessário mais uma vez refinarmos este conceito.

Definição B.63. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Um conjunto $A \in \mathcal{B}(S)$ é dito ser Harris recorrente se

$$P_x(n_A = +\infty) = 1 \quad \forall x \in A.$$

O processo de Markov é dito ser Harris recorrente se ele é φ -irredutível e todo o evento em $\mathcal{B}(S)$ é Harris recorrente.

Em palavras um evento A será Harris recorrente se com probabilidade 1 o processo de Markov X que começou em algum ponto deste conjunto torna a visitar este conjunto uma infinidade de vezes. Também temos que se um processo de Markov é Harris recorrente então ele é recorrente segundo a definição B.59.

Teorema B.64. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Se para todo conjunto $A \in \mathcal{B}(S)$ temos $P_x(\tau_A < +\infty) = 1$ para todo $x \in A$. Então $P_x(n_A = +\infty) = 1$ para todo $x \in S$ e o processo X é Harris recorrente.

Demonstração. Ver referência [MT09] página 201. □

Observe que a propriedade de ser Harris recorrente só é necessária quando temos um espaço de estados contínuo, pois no caso discreto temos que $\mathbb{E}_x[n_x] = +\infty$ se, e somente se $P_x(\tau_x < +\infty) = 1$.

Teorema B.65. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov φ -irredutível tomando valores no espaço de estados contínuo S . Considere que exista um conjunto pequeno $C \in \mathcal{B}(S)$ tal que $P_x(\tau_C < +\infty) = 1$ para todo $x \in S$. Então o processo X é Harris recorrente.

Demonstração. Ver referência [MT09] página 205. □

B.3.6 Distribuição invariante

Todo o esforço conceitual que fizemos até agora teve como único propósito, assim como no caso do espaço de estados discreto, estabelecer as condições para que as perguntas feitas no início desta seção, a saber, as condições necessárias e suficientes para a existência e unicidade da distribuição invariante de um processo de Markov. Começamos então pela definição de distribuição invariante para um processo de Markov em um espaço de estados contínuo.

Definição B.66. Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Considere $K : S \times \mathcal{B}(S) \rightarrow [0, 1]$ o núcleo de transição associado ao processo de Markov X . A medida de probabilidade σ -finita $\pi : \mathcal{B}(S) \rightarrow [0, 1]$ satisfazendo a equação

$$\pi(B) = \int_S K(x, B) \pi(dx) \quad \forall B \in \mathcal{B}(S)$$

é denominada de distribuição invariante do processo X .

Quando existe uma distribuição invariante para um processo de Markov φ -irredutível o processo é dito ser positivo. Processos recorrentes que não admitem uma distribuição invariante finita são chamados de processos nulos. Relembre que pelo teorema B.62 um processo φ -irredutível é recorrente ou transiente. Então podemos de fato falar em processos recorrentes positivos e nulos assim como nas cadeias de Markov.

Teorema B.67. *Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Se o processo X é positivo então é recorrente.*

Demonstração. Ver referência [MT09] página 231. $\square$

Note que o resultado acima é justamente o análogo contínuo do teorema B.27 apresentado para cadeias com espaço de estados discretos. Podemos finalmente apresentar o primeiro grande resultado com respeito a existência de uma medida invariante para um processo de Markov.

Teorema B.68. *Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov φ -irredutível com átomo α e tomando valores no espaço de estados contínuo S . Então o processo é recorrente positivo se, e somente se $\mathbb{E}_\alpha[\tau_\alpha] < +\infty$. Além disso se π é uma distribuição invariante para X então*

$$\pi(\alpha) = \frac{1}{\mathbb{E}_\alpha[\tau_\alpha]}$$

Demonstração. Ver referência [MT09] página 236. $\square$

O resultado acima nos diz que quando o evento considerado é atômico temos uma representação simples para a distribuição invariante. Como no caso discreto todos os eventos são atômicos temos que

$$\pi(i) = \frac{1}{\mathbb{E}_i[\tau_i]}$$

e recuperamos o resultado do teorema B.37. O segundo grande resultado desta seção é o seguinte.

Teorema B.69. *Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov φ -irredutível tomando valores no espaço de estados contínuo S . Considere também que X é recorrente. Então existe uma única distribuição invariante π e esta medida tem representação para todo $A \in \mathcal{B}(S)$ com $\varphi(A) > 0$ dada por*

$$\pi(B) = \int_A \mathbb{E}_w \left[\sum_{t=1}^{\tau_A} \mathbb{1}_B(X_t) \right] \pi(dw).$$

Demonstração. Ver referência [MT09] página 246. $\square$

B.3.7 Convergência para a distribuição invariante

Tão importante quanto a existência e unicidade da distribuição de probabilidade invariante para um processo de Markov e saber em quais condições temos a convergência para esta distribuição dado que o processo começou com uma distribuição μ , isto é, $X_0 \sim \mu$. Para isso precisamos dizer em qual métrica essa convergência se realiza. No caso dos processos de Markov a convergência se dá segundo a norma da variação total que lembramos abaixo.

Definição B.70. Seja $(S, \mathcal{B}(S))$ um espaço mensurável. Considere μ_1 e μ_2 medidas definidas sobre este espaço. Então a norma de variação total é definida por

$$\|\mu_1 - \mu_2\|_{TV} := \sup_{A \in \mathcal{B}(S)} \{|\mu_1(A) - \mu_2(A)|\}.$$

Assim como no caso discreto (ver teorema B.37) em que para assegurarmos a convergência da sequência $(p\mathcal{P}^t)_{t \in \mathbb{N}}$ para a distribuição invariante π precisávamos que a cadeia de Markov relacionada possuísse as propriedades de ser recorrente, irredutível e aperiódica temos que no caso de espaço de estados contínuos aproximadamente as mesmas propriedades serão demandadas para que a sequência dada pelo "produto" dos núcleos de transição pela distribuição inicial possa convergir para a distribuição invariante. Mais especificamente as propriedades mais importantes são a recorrência de Harris e a aperiodicidade.

Teorema B.71. *Sejam $(\Omega, \mathcal{O}, P)$ um espaço de probabilidade, $(S, \mathcal{B}(S))$ um espaço mensurável e $X = (X_t)_{t \in \mathbb{N}}$ um processo de Markov tomando valores no espaço de estados contínuo S . Considere $K : S \times \mathcal{B}(S) \rightarrow [0, 1]$ o núcleo de transição associado ao processo de Markov X . Considere que o processo X é Harris recorrente, positivo e aperiódico. Então temos que*

$$\lim_{t \rightarrow \infty} \left\| \int_S K^t(x, \cdot) \mu(dx) - \pi \right\|_{TV} = 0.$$

para toda distribuição inicial μ .

Demonstração. Ver referência [MT09] página 328. □

Exemplo B.72. (Continuação do exemplo B.43) Quando o parâmetro $\theta \in \mathbb{R}$ no modelo AR(1) tem norma menor que 1, ou seja, $|\theta| < 1$ temos que dado qualquer distribuição inicial os núcleos de transição $N(\theta x_{t-1}, \sigma^2)$ converge para a distribuição invariante do processo $X^\theta = (X_t^\theta)_{t \in \mathbb{N}}$ que é dada por $N(0, \sigma^2/(1 - \theta^2))$.

C Testes Estatísticos

Para avaliar a convergência ao estado estacionário, descrito na seção anterior, iremos empregar testes estatísticos que permitam quantificar se uma distribuição está próxima ou não da distribuição estacionária. Mais ainda, como em geral não possuímos uma fórmula para uma eventual distribuição estacionária, o que poderemos fazer é comparar os resultados de simulação, para um grande número de cenários, e avaliar a convergência.

Nesta seção, apresentaremos os aspectos fundamentais da *teoria da bondade do ajuste*, e estudaremos formas de aplicá-la à inferência da distribuição invariante de processos de Markov.

C.1 Testes da Bondade do Ajuste

Os testes de bondade do ajuste têm sua origem no problema de determinar se uma amostra é proveniente ou não de uma distribuição de probabilidade de um determinado tipo. Note que esta é uma situação diferente da que encontramos na inferência paramétrica. De fato, na inferência paramétrica já temos informações de qual família de distribuições de probabilidade a amostra pertence: o que falta é terminarmos o membro exato da família que originou os dados. Por exemplo, se já sabemos que os dados da amostra pertencem a uma distribuição normal, o que nos resta é estimar a média e a variância para fixar o membro da família normal que originou a nossa amostra.

Uma situação mais difícil ocorre quando não possuímos nem mesmo a informação de qual família pode ter originado os dados da amostra. Neste caso, com base na amostra, precisamos construir estimadores para a *distribuição de probabilidade desconhecida*, e não para parâmetros desconhecidos. Portanto, os testes de bondade do ajuste pertencem ao campo da Estatística Não-Paramétrica.

O estimador fundamental na classe dos testes de bondade do ajuste, e que serve de base para todos os demais métodos da classe, é a *distribuição empírica* dos dados da amostra. Antes de apresentarmos sua definição, vamos estabelecer algumas convenções e notações que serão utilizadas ao longo do texto.

1. A tripla $(\Omega, \mathcal{O}, P)$ é chamada de *espaço de probabilidade*; cada $\omega \in \Omega$ representa um evento de interesse; $\mathcal{O}$ é a família dos subconjuntos *mensuráveis* de Ω , ou seja, aqueles aos quais atribuiremos probabilidades; e $P : \mathcal{O} \rightarrow [0, 1]$ é a probabilidade de ocorrência dos eventos que estamos estudando.
2. Considere $X : \Omega \rightarrow \mathbb{R}$ uma variável aleatória. A probabilidade

$$P_X(A) := P(X^{-1}(A))$$

está definida para subconjuntos A de $\mathbb{R}$, a partir da probabilidade original P e da variável aleatória X , e é dita *distribuição* da variável aleatória X .

3. A *função de distribuição acumulada* da variável aleatória X é denotada F_X e definida pela “probabilidade acumulada”:

$$F_X(a) = P_X(-\infty, a] = P(X \leq a).$$

Note que F_X é crescente, indo de zero a um conforme a vai de $-\infty$ a $+\infty$.

4. Seja $n \in \mathbb{N}$. Uma sequência de variáveis aleatórias $(X_i)_{1 \leq i \leq n}$, independentes e identicamente distribuídas, com distribuição P_X , é uma *amostra aleatória* da variável X . Ou seja, temos uma sequência de n “cópias” independentes da variável X .

Com isto, podemos passar para a definição do principal estimador na teoria da bondade do ajuste:

Definição C.1 (Função de distribuição empírica). Seja $(X_i)_{1 \leq i \leq n}$ uma amostra aleatória da variável X . A variável aleatória $\hat{F}_n : \mathbb{R} \times \Omega^n \rightarrow [0, 1]$, definida por

$$\hat{F}_n(x, \tilde{\omega}) := \frac{1}{n} \sum_{i=1}^n \mathbb{1}_{\{X_i(\tilde{\omega}) \leq x\}} \quad (\text{C.1})$$

é chamada de *distribuição empírica* da amostra. O termo $\mathbb{1}_{\{X_i(\tilde{\omega}) \leq x\}}$ é chamado de *indicadora* do evento $X_i \leq x$, e vale 1 caso $X_i(\tilde{\omega}) \leq x$ e zero caso contrário.

Vamos dissertar um pouco sobre os elementos que aparecem na fórmula (C.1) e faremos isso por meio de um exemplo.

Comentário C.2. O exemplo C.3 que vamos apresentar a seguir deveria ser apresentado no contexto das *séries temporais* pois, a rigor, envolve uma variável aleatória evoluindo no tempo. Mas para dar um exemplo simples e ao mesmo tempo significativo vamos desconsiderar os aspectos de séries temporais envolvidos.

Exemplo C.3. Digamos que estamos interessados em estudar a demanda por energia elétrica no mês de Janeiro ao longo dos anos. Fixando um horizonte de estudo que vai de Janeiro de 1900 até Janeiro de 2100 podemos definir Ω como

$$\Omega := \{1900, 1901, \dots, 2022, \dots, 2100\}.$$

Definimos agora a variável aleatória $X_{\text{Jan}} : \Omega \mapsto \mathbb{N}$ como a demanda por energia elétrica em kWh (Quilowatts/hora) no mês de Janeiro. Por exemplo $X_{\text{Jan}}(1953)$ é a demanda por energia elétrica em kWh no mês de Janeiro no ano de 1953 e este valor é um número natural. Desejamos responder perguntas do tipo: qual é a probabilidade de termos uma demanda por energia elétrica no mês de Janeiro menor ou igual a 1500 kWh? Inicialmente devemos munir o espaço Ω com uma medida de probabilidade, por exemplo

$$P(\{1900\}) = p_{1900}, \dots, P(\{2100\}) = p_{2100}$$

com $p_{1900} + \dots + p_{2100} = 1$ e todos estes valores não negativos. Note que $\{1900\}$ e $\{2100\}$ são subconjuntos de Ω que sabemos medir (pertence a família $\mathcal{O}$). Agora dado $n \in \mathbb{N}$ a probabilidade de termos uma demanda de energia elétrica igual a n é dada por

$$P_{X_{\text{Jan}}}(\{n\}) = P(X_{\text{Jan}}^{-1}(\{n\})) = P(\{\omega \in \Omega : X_{\text{Jan}}(\omega) = n\}).$$

Por exemplo, poderíamos ter para $n = 1000$ o seguinte resultado

$$\begin{aligned} P_{X_{\text{Jan}}}(\{1000\}) &= P(\{\omega \in \Omega : X_{\text{Jan}}(\omega) = 1000\}) = P(\{1950, 1980, 2010\}) = \\ &= P(\{1950\}) + P(\{1980\}) + P(\{2010\}) = p_{1950} + p_{1980} + p_{2010} \end{aligned}$$

Ou seja,

$$P_{X_{\text{Jan}}}(\{1000\}) = p_{1950} + p_{1980} + p_{2010}.$$

Quanto a função de distribuição $F_{X_{\text{Jan}}} : \mathbb{R} \mapsto [0, 1]$ podemos escrever esta como

$$F_{X_{\text{Jan}}}(x) = \sum_{n \leq [x]} P_{X_{\text{Jan}}}(\{n\}).$$

Observe que apesar de termos uma fórmula para $F_{X_{\text{Jan}}}$ esta não nos ajuda muito pois não sabemos os valores $P_{X_{\text{Jan}}}(\{n\})$ e portanto não temos como calcular $F_{X_{\text{Jan}}}(1500)$. Além disso note que na verdade a fórmula acima deveria ter sido escrita como

$$F_{X_{\text{Jan}}}(x|\theta) = \sum_{n \leq [x]} P_{X_{\text{Jan}}}(\{n\}|\theta).$$

em que $\theta = (p_{1900}, \dots, p_{2100})$, ou seja, θ é o parâmetro desconhecido que devemos estimar. Com as informações que temos disponíveis a estimação dessa quantidade desconhecida pode ser uma tarefa não trivial. É precisamente aqui que entram os estimadores não paramétricos como a função de distribuição empírica. Vejamos como isso é feito. Considere que temos uma amostra aleatória de tamanho 10 da variável X_{Jan} . Relembrando do item 4 descrito antes da observação C.2 temos que cada termo da sequência $(X_i)_{1 \leq i \leq 10}$ é definido no espaço Ω^{10} e toma valores em $\mathbb{N}$. Digamos que temos a seguinte observação

$$\tilde{\omega} = (1940, 1956, 1987, 1964, 1997, 2001, 2010, 1922, 2018, 1988)$$

e associada a ela temos os seguintes valores de demanda de energia elétrica no mês de Janeiro em kWh

$$(1000, 2123, 9125, 3056, 15256, 20256, 22536, 860, 25000, 17030).$$

Ou seja, por exemplo, para $i = 4$ temos $X_4(\tilde{\omega}) = 3056$. Então para esta realização de $\tilde{\omega}$ temos a seguinte distribuição empírica

$$\hat{F}_{10}(x, \tilde{\omega}) = \frac{1}{10} (\mathbb{1}_{\{1000 \leq x\}}(x) + \mathbb{1}_{\{2123 \leq x\}}(x) + \dots + \mathbb{1}_{\{17030 \leq x\}}(x)).$$

Assim poderíamos usar a função de distribuição empírica $\hat{F}_{10}$ ao invés da distribuição $F_{X_{\text{Jan}}}$, mas para isso é preciso provar que para uma amostra de tamanho 10 a função de distribuição empírica $\hat{F}_{10}$ está suficientemente próxima de $F_{X_{\text{Jan}}}$ em algum sentido apropriado.

Como mencionado no final do exemplo C.3, para podermos utilizar a função de distribuição empírica como uma aproximação da função de distribuição real dos dados precisamos provar resultados que garantam que, quando aumentamos o tamanho da amostra $(X_i)_{1 \leq i \leq n}$, a sequência de funções de distribuição empíricas $(\hat{F}_n)_{n \geq 1}$ “converge” para F_X . Esta propriedade é conhecida como *consistência forte* do estimador. Na prática, se tivermos uma amostra aleatória (X_i) suficientemente grande podemos aproximar F_X tão bem quanto desejarmos. Este é justamente o conteúdo do próximo teorema:

Teorema C.4 (Teorema de Glivenko-Cantelli [LR05, p. 441]). *Seja $(X_i)_{1 \leq i \leq n}$ uma amostra aleatória de X . Considere $\hat{F}_n$ a função de distribuição empírica definida pela equação (C.1). Então*

$$\lim_{n \rightarrow +\infty} \sup_{x \in \mathbb{R}} \{|\hat{F}_n(x) - F_X(x)|\} = 0 \quad q.c. \quad (\text{C.2})$$

O teorema de Glivenko-Cantelli é um bom exemplo dos resultados que vamos explorar, pois na base dos testes de bondade do ajuste está a noção de aproximar funções.

Antes de apresentarmos mais ferramentas para estudar aproximações da função de distribuição, vamos analisar um pouco mais as propriedades pontuais (ou seja, para um $x \in \mathbb{R}$ fixo) da função de distribuição empírica. Esses resultados são a primeira etapa para obter os principais *resultados funcionais* sobre a distribuição empírica, ou seja, que permitem aproximar F_X como função.

Começamos com um resultado simples que nos diz que para cada $x \in \mathbb{R}$, $\hat{F}_n(x)$ é um estimador *não-enviesado* para o valor $F_X(x) \in [0, 1]$. Ou seja, temos $\mathbb{E}[\hat{F}_n(x)] = F_X(x)$, e portanto o estimador $\hat{F}_n(x)$, na média, nem subestima nem superestima o valor (correto, mas desconhecido) de $F_X(x)$. A razão disto é simples: como cada amostra X_i segue a mesma lei F_X , o valor esperado de $\mathbb{1}_{\{X_i \leq x\}}$ é exatamente a probabilidade de que $X_i \leq x$. Daí, aplicando a definição de $\hat{F}_n$, temos que

$$\mathbb{E}[\hat{F}_n(x)] = \mathbb{E}\left[\frac{1}{n} \sum_{i=1}^n \mathbb{1}_{\{X_i \leq x\}}\right] = \frac{1}{n} \sum_{i=1}^n \mathbb{E}[\mathbb{1}_{\{X_i \leq x\}}] = \frac{1}{n} \sum_{i=1}^n P(X_i \leq x) = \frac{1}{n} \sum_{i=1}^n F_X(x) = F_X(x).$$

Um outro resultado importante é que, para cada $x \in \mathbb{R}$, a variável aleatória $n\hat{F}_n(x)$ tem distribuição de probabilidade Binomial. De fato, inicialmente observe que

$$n\hat{F}_n(x) = \sum_{i=1}^n \mathbb{1}_{\{X_i \leq x\}}$$

pode ser entendido como uma soma de Bernoullies, pois as variáveis aleatórias $\mathbb{1}_{\{X_i \leq x\}}$ assumem apenas os valores 0 ou 1. Então a variância de $n\hat{F}_n(x)$ é dada por

$$\text{Var}\left[\sum_{i=1}^n \mathbb{1}_{\{X_i \leq x\}}\right] = \sum_{i=1}^n \text{Var}[\mathbb{1}_{\{X_i \leq x\}}] = \sum_{i=1}^n (F_X(x)(1 - F_X(x))) = n(F_X(x)(1 - F_X(x))).$$

Dividindo por n a variável, dividimos por n^2 sua variância, e portanto

$$\text{Var}[\hat{F}_n(x)] = \frac{1}{n}(F_X(x)(1 - F_X(x))).$$

Esse resultado implica que a variância do estimador $\hat{F}_n(x)$ tende a zero quando o tamanho da amostra tende para infinito.

Outro resultado assintótico importante é a normalidade assintótica de $\sqrt{n}\hat{F}_n(x)$. Para ver isto, basta notar que

$$\begin{aligned} \sqrt{n}(\hat{F}_n(x) - F_X(x)) &= \sqrt{n}\left(\frac{1}{n} \sum_{i=1}^n \mathbb{1}_{\{X_i \leq x\}}(x) - F_X(x)\right) = \\ &= \sqrt{n}\left(\frac{1}{n} \sum_{i=1}^n \mathbb{1}_{\{X_i \leq x\}}(x) - \frac{1}{n} \sum_{i=1}^n F_X(x)\right) = \sqrt{n}\left(\frac{1}{n} \sum_{i=1}^n (\mathbb{1}_{\{X_i \leq x\}}(x) - F_X(x))\right) = \\ &= \frac{1}{\sqrt{n}} \sum_{i=1}^n \left((\mathbb{1}_{\{X_i \leq x\}}(x) - F_X(x))\right). \end{aligned}$$

Pelo *teorema central do limite* segue-se que

$$\sqrt{n}(\hat{F}_n(x) - F_X(x)) = \frac{1}{\sqrt{n}} \sum_{i=1}^n \left((\mathbb{1}_{\{X_i \leq x\}}(x) - F_X(x))\right) \xrightarrow{d} N(0, (F_X(x)(1 - F_X(x))))$$

em que a notação $\xrightarrow{d}$ significa convergência em distribuição. De posse desses resultados pontuais podemos passar para os resultados funcionais que estão inseridos na teoria dos *processos empíricos*. Vamos estudar os aspectos básicos dessa teoria na próxima subseção.

C.2 Processos Empíricos

A necessidade do estudo dos processos empíricos reside no fato que as distribuições assintóticas de probabilidade dos testes estatísticos que vamos estudar são determinadas através de resultados estabelecidos para os processos empíricos. No contexto em que estamos trabalhando estes processos surgem de modo natural da seguinte forma. Fixado $n \in \mathbb{N}$ considere o processo estocástico $(\mathbb{B}_n(x))_{x \in \mathbb{R}}$ definido por

$$\mathbb{B}_n(x) := \sqrt{n}(\hat{F}_n(x) - F_X(x)) \quad (\text{C.3})$$

em que as funções $\hat{F}_n$ e F_X são definidas como na seção anterior. O processo $(\mathbb{B}_n(x))_{x \in \mathbb{R}}$ é chamado de *processo empírico* pois é um processo estocástico que depende de uma amostra aleatória da variável X . Assim como as distribuições empíricas os processos empíricos também são funções aleatórias. Desejamos obter resultados assintóticos para funções aleatórias, pois vamos construir futuramente testes estatísticos que buscam determinar quando duas funções de distribuição de probabilidade estão suficientemente próximas (são estatisticamente a mesma função). Inicialmente note que dado $x \in \mathbb{R}$ temos pelo teorema central do limite (visto na seção anterior) que

$$\mathbb{B}_n(x) \xrightarrow{d} \mathbb{B}(x)$$

em que $\mathbb{B}(x) \sim N(0, (F_X(x)(1 - F_X(x))))$. Agora observe que dado $(x_1, x_2, \dots, x_m) \in \mathbb{R}^m$ com $x_1 < \dots < x_m$ segue-se pelo teorema central do limite multivariado que

$$(\mathbb{B}_n(x_1), \dots, \mathbb{B}_n(x_m)) \xrightarrow{d} (\mathbb{B}(x_1), \dots, \mathbb{B}(x_m)).$$

Intuitivamente quanto maior o número de pontos $x_1 < \dots < x_m$ escolhidos melhor seria a nossa "aproximação" da função aleatória limite $\mathbb{B}$. As aspas foram inseridas pois a convergência é em distribuição então não há uma real aproximação entre os vetores aleatórios. Apesar de intuitivo esse procedimento não é suficiente para estudarmos a convergência da função aleatória $\mathbb{B}_n$ para a função $\mathbb{B}$. Para este estudo precisamos do *teorema central do limite funcional* ou *teorema de Donsker*. Informalmente este teorema pode ser enunciado como.

Teorema C.5 (Teorema de Donsker- Informal). *Considere a sequência $(\mathbb{B}_n)_{n \geq 1}$ de processos empíricos definida como nos parágrafos anteriores. Então temos que esta sequência converge fracamente para o processo Gaussiano $\mathbb{B}$ com função de média identicamente nula e função de covariância dada por*

$$c(x, y) = \text{Cov}[\mathbb{B}(x), \mathbb{B}(y)] = F_X(x \wedge y) - F_X(x)F_X(y) \quad \forall x, y \in \mathbb{R}.$$

Em geral, por comodidade, denotamos um processo ser Gaussiano com função de média zero e função de covariância $c(x, y)$ por $\mathbb{B}(x) \sim \mathcal{GP}(0, c(x, x))$, em que $c(x, x) = F_X(x) - (F_X(x))^2$ é a função de variância. O leitor poderia se perguntar a esta altura como o teorema C.5 pode nos auxiliar na determinação das distribuições assintóticas de probabilidades dos estimadores em que temos interesse. Antes de responder a esta pergunta considere o seguinte lema auxiliar que em função do nível das técnicas envolvidas na sua construção matemática rigorosa vamos apresentar no mesmo "espírito" informal do teorema de Donsker.

Lema C.6 (Continuous mapping theorem-Informal). *Para toda função contínua $g : E \rightarrow \mathbb{R}$ definida no espaço métrico E temos que se $\mathbb{B}_n \xrightarrow{w} \mathbb{B}$ então $g(\mathbb{B}_n) \xrightarrow{d} g(\mathbb{B})$.*

No resultado acima a notação $\xrightarrow{w}$ significa convergência fraca enquanto que a notação $\xrightarrow{d}$ significa convergência em distribuição. O motivo desta separação se deve ao fato que a convergência em distribuição tem por base as funções de distribuição que estão por sua vez definidas na reta. Já a convergência fraca pode ser definida em espaços métricos mais gerais (espaços em que podemos definir uma noção de distância).

Para responder a pergunta colocada anteriormente vamos usar como exemplo um estatística que será estudada em seções futuras. A estatística de *Kolmogorov-Smirnov* é definida por

$$T_n^{\text{KS}}(X) := \sqrt{n} \sup_{x \in \mathbb{R}} \{ |\hat{F}_n(x) - F_X(x)| \}.$$

Essa estatística é uma das mais clássicas na teoria da bondade do ajuste e serve de base de compreensão para estatísticas mais recentes dentro dessa teoria. Considerando $\mathfrak{F}$ o espaço das funções de distribuição podemos definir sobre esse espaço funções do tipo $g : \mathfrak{F} \rightarrow \mathbb{R}_+$ dadas por

$$g(F) := \sup_{x \in \mathbb{R}} \{ |F(x)| \}$$

Relembrando a definição do processo empírico $\mathbb{B}_n$ temos da fórmula acima e da expressão para a estatística de Kolmogorov-Smirnov que

$$T_n^{\text{KS}}(X) = \sup_{x \in \mathbb{R}} \{ |\sqrt{n}(\hat{F}_n(x) - F_X(x))| \} = \sup_{x \in \mathbb{R}} \{ |\mathbb{B}_n(x)| \} = g(\mathbb{B}_n)$$

Ou seja,

$$T_n^{\text{KS}}(X) = g(\mathbb{B}_n)$$

Pelo teorema C.5 (teorema de Donsker) segue-se que $\mathbb{B}_n \xrightarrow{w} \mathbb{B}$ e pelo lema C.6 obtemos que

$$T_n^{\text{KS}}(X) \xrightarrow{d} g(\mathbb{B}).$$

Então assintoticamente temos que a função de distribuição de $T_n^{\text{KS}}(X)$ coincide com a função de distribuição de $g(\mathbb{B})$. Como podemos intuir obter a função de distribuição de uma função de um processo Gaussiano é, em geral, uma tarefa não trivial. Para algumas estatísticas como a de Kolmogorov-Smirnov isso pode ser feito apesar da expressão do função de distribuição obtida ser complexa e um pouco difícil de trabalhar.

De posse das informações essenciais sobre processos empíricos passamos na próxima subseção para o estudo de uma ferramenta de inferência central na teoria dos testes de bondade do ajuste e da teoria estatística em geral que são os testes de hipótese.

C.3 Testes de hipótese na teoria da bondade do ajuste.

Como observado nas seções anteriores a teoria dos testes de bondade do ajuste busca estimar a função de distribuição de uma variável aleatória (o caso vetorial será tratado nas seções finais) através da função de distribuição empírica. Vimos que esse estimador $\hat{F}_n$ é fortemente consistente (teorema de Glivenko-Cantelli) o que significa que $\hat{F}_n$ converge quase certamente para F_X uniformemente em $x \in \mathbb{R}$. Agora considere $d_{\text{sup}} : \mathfrak{F} \times \mathfrak{F} \rightarrow [0, +\infty)$ uma métrica definida por

$$d_{\text{sup}}(F, G) := \sup_{x \in \mathbb{R}} \{|F(x) - G(x)|\}. \quad (\text{C.4})$$

Então temos pela discussão feita no parágrafo anterior que

$$\lim_{n \rightarrow +\infty} d_{\text{sup}}(\hat{F}_n, F_X) = 0 \quad \text{q.c.} \quad (\text{C.5})$$

Esse resultado nos motiva a construir o seguinte estudo: dado uma amostra $(X_i)_{0 \leq i \leq n}$ desejamos saber se essa amostra tem função de distribuição F . Ou seja, desejamos saber se $F_X = F$. Por exemplo, poderíamos conjecturar que a distribuição dos dados observados vem de uma distribuição normal e portanto nesse caso estamos considerando que F pertence a família das distribuições normais. O leitor atento já deve ter percebido que temos o seguinte problema: sabemos que as densidades de probabilidade da distribuição normal apresentam a característica curva em forma de “sino” mas a posição e forma desse “sino” dependem da média e da variância. Logo não é possível montar um procedimento de inferência que nos diga algo como: “os dados obtidos tem origem em uma distribuição de probabilidade cuja a densidade tem uma curva em forma de sino”. O que é comumente feito para resolver esse problema é estimar a média e a variância da distribuição F utilizando os próprios dados da amostra. Por exemplo, podemos utilizar a média $\bar{x}$ e variância $\hat{\sigma}^2$ amostrais para fixar $F(\cdot | \bar{x}, \hat{\sigma}^2)$ na classe das distribuições normais e depois efetuamos os procedimentos de inferência que iremos descrever nos próximos parágrafos.

Formalmente identificamos essa *hipótese* com a notação $H_0 : F_X = F$ de onde temos que

$$d_{\text{sup}}(F_X, F) = 0 \Leftrightarrow H_0 : F_X = F.$$

Ou seja, a distância entre a distribuição F_X dos dados e da distribuição conjecturada F será zero na métrica do supremo se, e somente se a hipótese nula for verdadeira, isto é, $F_X = F$. Como consequência do limite (C.5) e da igualdade acima podemos escrever

$$\lim_{n \rightarrow +\infty} d_{\text{sup}}(\hat{F}_n, F) = 0 \quad \text{q.c} \Leftrightarrow H_0 : F_X = F. \quad (\text{C.6})$$

Então para um tamanho de amostra suficientemente grande e uma constante $c > 0$ podemos considerar (de um modo a ser tornado preciso mais adiante) que a hipótese $H_0 : F_X = F$ é verdadeira quando $d_{\text{sup}}(\hat{F}_n, F) < c$. Essa é a intuição dos testes de hipótese no contexto dos testes da bondade do ajuste. Como a teoria dos testes de hipótese é muito extensa vamos nos preocupar apenas com os aspectos mais básicos dessa teoria. Um tratamento completo sobre esses assuntos pode ser encontrado em [LR05]. Um teste de hipótese é composto dos seguintes objetos:

- 1 Um par de hipóteses mutuamente exclusivas chamadas respectivamente de hipótese *nula* e *alternativa*. Utilizamos a notação H_0 para hipótese nula e H_1 para alternativa. Nestas notas estamos interessados na seguinte estrutura

$$\begin{cases} H_0 : F_X = F \\ H_1 : F_X \neq F \end{cases} \quad (\text{C.7})$$

que é chamada de hipótese *nula simples vs alternativa composta*.

- 2 O segundo objeto a ser definido é uma *estatística* e como discutimos (quando introduzimos a função de distribuição empírica) essa é uma função da amostra aleatória. Nos testes da bondade do ajuste para uma amostra usamos a seguinte fórmula geral

$$T_n(X) := c(n)d_{\text{geral}}(\hat{F}_n, F). \quad (\text{C.8})$$

em que $d_{\text{geral}} : \mathfrak{F} \times \mathfrak{F} \mapsto [0, +\infty)$ é uma métrica geral e $c(\cdot)$ é uma função do tamanho da amostra e auxilia nos resultados assintóticos. Quando estivermos trabalhando com duas amostra independentes de tamanhos n_1 e n_2 vamos utilizar a seguinte formulação geral

$$T_{n_1, n_2}(X^{(1)}, X^{(2)}) := c(n_1, n_2)d_{\text{geral}}(\hat{F}_{n_1}, \hat{F}_{n_2}). \quad (\text{C.9})$$

Com os principais objetos já definidos vamos estudar o processo inferencial dos testes de hipótese no contexto da *teoria da decisão estatística*. Dentro do escopo dessa teoria definimos um critério para decidir aceitar ou rejeitar a hipótese nula H_0 . Ao tomarmos essa decisão podemos estar cometendo os seguintes tipos de erros:

- I O *erro do tipo I* ocorre quando a amostra que temos disponível nos induz (segundo o critério que nós estabelecemos) a rejeitar a hipótese nula quando essa é verdadeira.
- II O *erro do tipo II* ocorre quando a amostra que temos disponível nos induz (segundo o critério que nós estabelecemos) a aceitar a hipótese nula quando essa é falsa.

Em geral na maioria das aplicações o erro que mais buscamos evitar é o erro do tipo I. Pois na modelagem dos testes de hipótese costumamos fazer a afirmação mais forte na hipótese nula. Portanto faz sentido fixarmos uma probabilidade (uma tolerância) para a chance de cometermos o erro do tipo I. Chamamos essa probabilidade de *nível de significância* do teste e a denotamos por $\alpha \in (0, 1)$. Como α é a probabilidade de cometermos um erro vamos escolher sempre valores pequenos sendo os mais tradicionais 0,01 e 0,05.

Passamos agora para a construção do critério de decisão. Como destacado nas seções anteriores a determinação da distribuição assintótica de probabilidade da estatística que escolhemos utilizar é fundamental. Vamos denotar por $P_{T_n(X)}$ a probabilidade assintótica da estatística $T_n(X)$. Como o caso para mais de uma amostra é praticamente análogo ao caso de uma amostra vamos nos concentrar nesta seção no segundo caso, ou seja, uma amostra. Pela conclusão feita após a equação

(C.6) temos que se a distância $d_{\text{geral}}(\hat{F}_n, F)$ for maior que uma constante positiva então deveríamos rejeitar a hipótese nula. Logo podemos probabilisticamente escrever o erro do tipo I para n suficientemente grande como

$$P_{T_n(X)}(T_n(X) > c | H_0 : F_X = F) = \alpha. \quad (\text{C.10})$$

O modo adequado de lermos a equação acima é o seguinte: ao nível de significância $\alpha \in (0, 1)$ cometemos o erro do tipo I (rejeitar a hipótese nula quando essa é verdadeira) quando o valor da estatística $T_n(X)$ utilizada no teste for maior que a contante $c > 0$ (a ser calculada) e isso ocorre justamente com probabilidade α . Nesse contexto o evento $\mathcal{RC}(c) := \{T_n(X) > c\}$ é chamado de *região crítica* e o seu complementar denotado por $\mathcal{RA}(c)$ é chamado de *região de aceitação*. Com base nesses conjuntos podemos montar a seguinte *função de decisão*.

$$d_\alpha(\tilde{x}) := \begin{cases} \text{rejeitar } H_0 & ; \quad t(\tilde{x}) > c \\ \text{aceitar } H_0 & ; \quad t(\tilde{x}) \leq c \end{cases} \quad (\text{C.11})$$

em que $t(\tilde{x}) \in \mathbb{R}$ é o valor efetivamente calculado da estatística T_n na realização $\tilde{x}$ (em que $\tilde{x} \in \mathbb{R}^n$ é o valor observado do vetor aleatório $(X_1, \dots, X_n)$). A função de decisão d_α opera do seguinte modo: ao nível de significância $\alpha \in (0, 1)$ e dispondo de uma amostra $\tilde{x}$ decidimos rejeitar a hipótese nula se o valor da estatística T_n efetivamente calculado, ou seja, $t(\tilde{x})$ é maior que uma constante $c > 0$ e ao fazer isso podemos estar cometendo um erro do tipo I com probabilidade α . Nos resta apenas calcularmos o valor da constante $c > 0$ pois o nível de significância α sempre deve ser dado a priori, ou seja, antes que qualquer procedimento inferencial seja executado. Olhando para equação (C.10) percebemos que a única quantidade desconhecida é justamente a constante $c > 0$. No caso da distribuição assintótica de probabilidade $P_{T_n(X)}$ ter densidade podemos determinar o valor da constante c com certa facilidade. Inicialmente note que

$$F_{T_n(X)}(c) = P_{T_n(X)}(T_n(X) \leq c | H_0 : F_X = F) = 1 - \alpha$$

Ou seja,

$$F_{T_n(X)}(c) = 1 - \alpha.$$

Como nesse caso a função de distribuição possui inversa obtemos da igualdade acima o seguinte resultado

$$c = F_{T_n(X)}^{-1}(1 - \alpha).$$

Assim nesse contexto a função de decisão d_α está completamente especificada.

Até agora não discutimos de modo mais aprofundado sobre o erro do tipo II e apesar da nossa estratégia de decisão priorizar evitar o erro do tipo I temos que o erro do tipo II é fundamental pois é através dele que construímos o importante conceito de *poder de um teste*. Como esse conceito é usado para a comparação de testes estatísticos dentro da teoria da bondade do ajuste vamos dedicar a próxima subseção para tratarmos sobre ele.

C.3.1 Poder de um teste estatístico

O poder de um teste estatístico é um dos conceitos mais finos dentro da teoria dos testes de hipóteses e também ocupa um lugar destaque na teoria da bondade do ajuste. Logo em função da sutileza desse conceito vamos proceder da forma mais cautelosa e clara possível. Começamos pela definição da *função poder*.

Definição C.7 (Função poder). Considere $\mathfrak{F}$ o espaço das funções de distribuição e $\mathcal{RC}(c)$ a região crítica delineada pela estatística $T_n(X)$ através da fixação de nível de significância $\alpha \in (0, 1)$ e da determinação da constante $c > 0$ como mostrado na seção anterior. A função $\pi_{\alpha,n} : \mathfrak{F} \rightarrow [0, 1]$ definida por

$$\pi_{\alpha,n}(F) := P_{T_n(X)}(\mathcal{RC}(c)|F) \quad (\text{C.12})$$

é chamada de *função poder*.

A função (C.12) tem uma série de propriedades importantes, mas no contexto em que estamos trabalhando a propriedade mais interessante e que nos ajudar a entender como essa função é usada na prática é a seguinte.

Lema C.8 (Monotonicidade da função poder - Informal). *Sejam F_1 e F_2 funções de distribuição pertencentes ao conjunto $\mathfrak{F}$. Considere que F_X é a função de distribuição da amostra que possuímos. Então segue-se que*

$$d_{\text{geral}}(F_1, F_X) \geq d_{\text{geral}}(F_2, F_X) \Rightarrow \pi_{\alpha,n}(F_1) \geq \pi_{\alpha,n}(F_2) \quad (\text{C.13})$$

Pensando um pouco sobre a definição da função poder percebemos o quando a lema (C.8) é intuitivo. Um modo simples de ver isso é o seguinte: as funções de distribuição F_1 e F_2 são as distribuições que conjecturamos para os dados da amostra que segue de fato a distribuição F_X . Digamos que tanto $F_1 \neq F_X$ quanto $F_2 \neq F_X$ mas que $d_{\text{geral}}(F_1, F_X) \geq d_{\text{geral}}(F_2, F_X)$, ou seja, a distribuição F_1 está mais “errada” do que a distribuição F_2 quando comparada com F_X . Então a desigualdade $\pi_{\alpha,n}(F_1) \geq \pi_{\alpha,n}(F_2)$ nos diz que temos uma maior probabilidade de tomar uma decisão correta quando conjecturamos F_1 para distribuição verdadeira dos dados do que quando conjecturamos F_2 . Para ver isso escrevendo $\pi_{\alpha,n}(F_1)$ explicitamente temos que

$$\pi_{\alpha,n}(F_1) = P_{T_n(X)}(\mathcal{RC}(c)|F_1) = P_{T_n(X)}(\mathcal{RC}(c)|H_1 : F_1 \neq F_X).$$

Então a equação acima nos diz que: tomamos a decisão correta de rejeitar a hipótese nula dado que a hipótese alternativa é verdadeira com probabilidade $\pi_{\alpha,n}(F_1)$. Agora se F_1 apesar de diferente F_X for suficientemente próxima então as chances de tomar a decisão correta vão ser menores. É importante notarmos que a probabilidade $\pi_{\alpha,n}$ depende do tamanho da amostra e do nível de significância adotado a priori. Além é claro da estatística que resolvemos usar para testar as hipóteses. Isso significa que para um dado nível de significância $\alpha \in (0, 1)$ se tivermos que F_1 está muito próxima de F_X então para podermos perceber essa diferença vamos precisar aumentar consideravelmente o tamanho da amostra.

Com as considerações feitas no parágrafo anterior podemos passar para a principal forma como a função poder é utilizada na prática, a saber, na comparação entre estatísticas para testes de hipóteses. Fixados o tamanho da amostra e o nível de significância podemos verificar, com a ajuda da função poder, como duas estatísticas $T_n^1(X)$ e $T_n^2(X)$ se comportam na hora de discernir a função de distribuição conjecturada e a função de distribuição verdadeira dos dados. Note que nesse caso estamos analisando o pares $(T_n^1(X), P_{T_n^1(X)})$ e $(T_n^2(X), P_{T_n^2(X)})$. Antes de vermos com essa comparação deve proceder vamos modificar um pouco a função de decisão (C.11) e passarmos a escreve-lá como

$$d_\alpha(\tilde{x}) := \begin{cases} \text{rejeitar } H_0 & ; \tilde{x} \in \mathcal{RC}(c), \\ \text{aceitar } H_0 & ; \tilde{x} \in \mathcal{RA}(c). \end{cases} \quad (\text{C.14})$$

Sabemos que dado um nível de significância e um tamanho de amostra podemos, para um determinado par $(T_n^1(X), P_{T_n^1(X)})$, encontrar uma constante $c > 0$ e assim definir as regiões críticas e de

aceitação. Com isso escrevemos a função de decisão (C.14) que apesar de depender da estatística escolhida e da sua distribuição de probabilidade nos fornece um critério geral de decisão.

Fixado uma função de distribuição F conjecturada e uma métrica, como por exemplo, a métrica do supremo podemos então olhar para a função poder associada a cada estatística e fazer o seguinte estudo: quando mais distante a função F estiver ser correta, ou seja, ser igual a F_X como as probabilidades $\pi_{\alpha,n}^1(F)$ e $\pi_{\alpha,n}^2(F)$ se comportam? Em quais caso temos $\pi_{\alpha,n}^1(F) > \pi_{\alpha,n}^2(F)$? Como podemos perceber as respostas a essas questões são, em geral, difíceis de responder. Por tanto não tentaremos conduzir esses estudos em relação as estatísticas que serão apresentadas nas próximas seções, mas vamos informar o leitor o que os estudos empíricos mais recentes tem revelado. Por fim, gostaríamos de salientar que uma outra questão muito importante sobre esse tipo de análise (talvez mais importante que as duas questões anteriores) ocorre quando F é diferente de F_X mais está suficientemente próxima. Nesse caso desejamos saber qual das duas estatísticas $T_n^1(X)$ e $T_n^2(X)$ é sensível o suficiente para capturar essa diferença para um mesmo tamanho de amostra. Ou seja, qual é mais poderosa para nos auxiliar a tomar a decisão correta.

C.3.2 Critério de decisão via p-Valor

Além do critério de decisão via região crítica que foi abordado nas subseções anteriores temos que um dos critérios mais simples de ser formulado é através do *p-Valor*. Começamos então por definir esse conceito. Quando o teste que desejamos realizar tem uma estrutura do tipo “rejeitar H_0 se $T_n(X) \geq t(\tilde{x})$ ” o p-Valor pode ser matematicamente definido pela seguinte fórmula.

Definição C.9 (p-Valor). Considere T_n a estatística definida na equação (C.8). Chamamos de *p-Valor* a seguinte probabilidade

$$p(\tilde{x}) := P_{T_n(X)}(T_n(X) \geq t(\tilde{x}) | H_0 : F_X = F). \quad (C.15)$$

em que $P_{T_n(X)}$ é a distribuição de probabilidade da estatística $T_n(X)$.

Então o p-Valor é: a probabilidade de termos obtido um valor igual ou maior a $t(\tilde{x})$ da estatística $T_n(X)$ dado que a hipótese nula H_0 é verdadeira. Vamos agora usar o p-Valor para construir um critério de decisão simples. Por exemplo, considere que estamos estudando uma dada população e queremos saber se uma determinada característica A está presente nessa população ou não. Então construímos uma hipótese nula que afirma: “ H_0 : a população tem a característica A ”. Considere também que essa característica pode ser detectada por uma estatística $T_n(X)$. Pode-se pensar então da seguinte forma: se a probabilidade $p(\tilde{x})$ for pequena podemos considerar que as chances da estatística $T_n(X)$ assumir valores maiores ou iguais a $t(\tilde{x})$ dado que a população apresenta a característica afirmada pela hipótese H_0 são muito baixas e portanto a hipótese nula deveria ser descartada.

Para termos um critério de decisão precisamos tornar rigorosa a expressão “se a probabilidade $p(\tilde{x})$ for pequena” para isso utilizamos o conceito já introduzido anteriormente de *nível de significância* $\alpha \in (0, 1)$ que, como sabemos, é um valor escolhido a priori (antes de começarmos a análise estatística) e dizemos que a probabilidade $p(\tilde{x})$ é pequena se $p(\tilde{x}) < \alpha$ e é significativa caso contrário. Então dada uma realização da amostra $\tilde{x} = (x_i)_{0 \leq i \leq n}$ formulamos o seguinte critério de decisão.

$$d_\alpha(\tilde{x}) := \begin{cases} \text{aceitar } H_0 & ; \quad p(\tilde{x}) \geq \alpha \\ \text{rejeitar } H_0 & ; \quad p(\tilde{x}) < \alpha. \end{cases} \quad (C.16)$$

Uma forma simples de entender a função de decisão dada na equação (C.16) é olhar para a relação entre p-Valor e região crítica. Por exemplo, se $p(\tilde{x}) < \alpha$ então temos que

$$P_{T_n(X)}(T_n(X) \geq t(\tilde{x}) | H_0 : F_X = F) < P_{T_n(X)}(T_n(X) > c | H_0 : F_X = F)$$

Essa desigualdade implica que

$$\{T_n(X) \geq t(\tilde{x})\} \subset \{T_n(X) > c\} = \mathcal{RC}(c)$$

que por sua vez implica que $t(\tilde{x}) > c$. Assim em conformidade com a função de decisão (C.11) rejeitamos a hipótese nula. Por outro lado, se $p(\tilde{x}) \geq \alpha$ então temos que

$$P_{T_n(X)}(T_n(X) \geq t(\tilde{x}) | H_0 : F_X = F) \geq P_{T_n(X)}(T_n(X) > c | H_0 : F_X = F).$$

Multiplicando ambos os lados da desigualdade acima por -1 obtemos

$$-P_{T_n(X)}(T_n(X) \geq t(\tilde{x}) | H_0 : F_X = F) \leq -P_{T_n(X)}(T_n(X) > c | H_0 : F_X = F).$$

Agora somando 1 a ambos os lados da desigualdade acima

$$1 - P_{T_n(X)}(T_n(X) \geq t(\tilde{x}) | H_0 : F_X = F) \leq 1 - P_{T_n(X)}(T_n(X) > c | H_0 : F_X = F).$$

De onde segue-se que

$$P_{T_n(X)}(T_n(X) < t(\tilde{x}) | H_0 : F_X = F) \leq P_{T_n(X)}(T_n(X) \leq c | H_0 : F_X = F).$$

Essa desigualdade implica que

$$\{T_n(X) < t(\tilde{x})\} \subseteq \{T_n(X) \leq c\} = \mathcal{RA}(c)$$

que por sua vez implica que $t(\tilde{x}) \leq c$. Assim mais um vez em conformidade com a função de decisão (C.11) aceitamos a hipótese nula. Uma outra observação importante a ser feita sobre o p-Valor ocorre quando é possível obter a distribuição exata da estatística $T_n(X)$ que nesse caso vai depender do tamanho da amostra n . Nesse contexto para que um determinado fenômeno possa ser detectado por um teste estatístico é necessário termos em mãos uma amostra suficientemente grande. Um exemplo importante em que isso ocorre é quando conjecturamos que a distribuição dos dados que temos em mãos é F mas na realidade a distribuição é G com as duas estando muito próximas na métrica do supremo. Assim para que uma estatística como a de Kolmogorov-Smirnov possa detectar essa sutil diferença será necessário que a amostra não seja pequena.

C.4 Testes de hipótese quando desejamos comparar duas distribuições.

O nosso principal interesse nestas notas no que tange aos testes de hipótese são os testes que nos permitem saber se duas amostras aleatórias independentes são provenientes de uma mesma distribuição ou não. A formulação desses testes é um pouco distinta da apresentada em (C.7) (uma amostra) pois aqui conjecturamos que a primeira amostra $X^{(1)}$ é proveniente de uma função de distribuição F_1 e que a segunda amostra $X^{(2)}$ (que é independente da primeira) é proveniente de uma função de distribuição F_2 . Desejamos então saber se estamos ou não falando da mesma distribuição, isto é, $F_1 = F_2$. Nesse contexto a formulação natural do teste de hipóteses que necessitamos é essa

$$\begin{cases} H_0 : F_1 = F_2 \\ H_1 : F_1 \neq F_2 \end{cases} \quad (\text{C.17})$$

A partir daqui podemos repetir boa parte do desenvolvimento teórico feito para o caso (C.7) fazendo obviamente as devidas mudanças para que as fórmulas continuem valendo como por exemplo na fórmula (C.10) em passamos a escrevê-la como

$$P_{T_n(X)}(T_n(X) > c | H_0 : F_1 = F_2) = \alpha. \quad (\text{C.18})$$

C.5 Testes para uma amostra.

Nesta subseção vamos dissertar sobre os testes da bondade do ajuste para o caso de uma amostra. Isso significa que apesar de não conseguirmos especificar através da nossa experiência e dos dados disponíveis uma distribuição de probabilidade dentro de uma classe de distribuições nós somos capazes de pelo menos conjecturar a qual classe a distribuição de probabilidade dos dados pertence. Por isso o nome *testes para uma amostra* pois temos uma conjectura sobre qual é a distribuição dos dados e através de uma amostra desejamos refutar ou não essa conjectura.

Nesse contexto vamos apresentar um teste que tem por base duas estatísticas distintas. Como podemos perceber na seção sobre testes de hipótese na teoria da bondade do ajuste o que fazemos em geral é comparar a distância entre a função empírica e a função de distribuição conjecturada para os dados rejeitando a hipótese nula caso essa distância seja maior que um determinado valor crítico calculado a priori. A vantagem desse tipo de teste é que ele nos permite entender claramente os motivos que nos levaram a rejeitar a hipótese nula e além disso nos permite apresentar uma fórmula simples para o p-Valor que em geral tem uma fórmula muito complexa quando se utiliza outros testes. Quanto as estatísticas vamos trabalhar com a estatística de Kolmogorov-Smirnov e a estatística de Anderson-Darling. Deixaremos os comentários mais detalhados sobre cada uma delas nas subseções dedicadas a elas.

C.5.1 Teste de Kolmogorov-Smirnov.

Nesta seção vamos apresentar uma das primeiras estatísticas propostas na teoria da bondade do ajuste que é a estatística de *Kolmogorov-Smirnov*. Já apresentamos a fórmula matemática dessa estatística na seção sobre processos empíricos mas para comodidade do leitor vamos reapresentar ela aqui em forma de definição.

Definição C.10 (Estatística de Kolmogorov-Smirnov). A estatística $T_n^{KS}(X)$ definida por

$$T_n^{KS}(X) := \sqrt{n} \sup_{x \in \mathbb{R}} \{|\hat{F}_n(x) - F_X(x)|\} \quad (\text{C.19})$$

é chamada de Kolmogorov-Smirnov.

Como podemos perceber pela fórmula dada na equação (C.19) essa estatística tem por base a métrica do supremo que é uma das formas mais simples de se comparar duas funções. Além disso como observado na seção sobre processos empíricos essa fórmula nos permite obter a distribuição de probabilidade tanto exata quanto assintótica dessa estatística. No teorema abaixo apresentamos a distribuição de probabilidade exata e por comodidade de notação vamos denotar essa distribuição apenas por P ao invés de $P_{T_n^{KS}(X)}(\cdot | H_0 : F_X = F)$, mas o leitor deve ter sempre em mente que essa distribuição é obtida sob a hipótese nula ser verdadeira.

Teorema C.11 (Distribuição exata da estatística de Kolmogorov-Smirnov: Ver referência [Sha03] página 447 primeiro item do teorema 6.10.). *Seja $(X_i)_{0 \leq i \leq n}$ uma amostra da variável aleatória X com função de distribuição $F_X : \mathbb{R} \mapsto [0, 1]$ e considere $\hat{F}_n : \mathbb{R} \mapsto [0, 1]$ a função de distribuição empírica definida na equação (C.1). Por fim, seja $T_n^{KS}(X)$ a estatística de Kolmogorov-Smirnov. A distribuição de probabilidade exata dessa estatística é dada por*

$$P(T_n^{KS}(X) \leq t) = \left(n! \prod_{i=1}^n \int_{a_n^i(u_n, t)}^{b_n^i(u_n, t)} du_1 \dots du_n \right) \mathbf{1}_{(1/2n, 1)}(t) + \mathbf{1}_{[1, +\infty)}(t) \quad (\text{C.20})$$

em que

$$a_n^i(u_n, t) := \max\{0, (n - i + 1)/n - t\}, \quad b_n^i(u_n, t) := \min\{u_{n-i+2}, (n - i)/n + t\} \text{ e } u_{n+1} = 1$$

Como podemos notar, trabalhar diretamente com a distribuição de probabilidade exata da estatística de Kolmogorov-Smirnov para construir testes de hipótese e intervalos de confiança pode não ser viável na maioria dos casos. Por isso é comum considerar a versão assintótica dessa distribuição que possui uma estrutura mais simples.

Lema C.12 (Distribuição assintótica da estatística de Kolmogorov-Smirnov: Ver referência [Sha03] página 447 segundo item do teorema 6.10.). *Nas mesmas condições do teorema C.11 temos que o seguinte limite existe.*

$$\lim_{n \rightarrow +\infty} P(T_n^{KS}(X) \leq t) = 1 - 2 \sum_{i=1}^{+\infty} (-1)^{i-1} e^{-2i^2 t^2} \quad (\text{C.21})$$

Na prática são as aproximações finitas para a distribuição de probabilidade dada na equação (C.21) que são utilizadas para a construção dos testes de hipótese independentemente de terem por base regiões críticas ou p-Valor e a expressão para as funções de decisão são as mesmas que as apresentadas na seção dos testes de hipótese. Além disso essa expressão, a princípio, permite que outros estudos importantes sejam conduzidos como por exemplo sobre o Poder do teste de Kolmogorov-Sminorv.

Nessa subseção apresentamos os rudimentos do teste de Kolmogorov-Smirnov para uma amostra. Na próxima subseção vamos continuar estudando os testes para uma amostra mas dessa vez tendo por base métricas dadas por integrais.

C.5.2 Teste de Anderson-Darling para uma amostra.

Nesta subseção vamos apresentar uma estatística $T_n(X)$ para testes da bondade do ajuste que está inserida na classe dos assim chamados *testes integrais de processos empíricos*. Inicialmente relembre que a definição do processo empírico $\mathbb{B}$ definido na equação (C.3) era dada por

$$\mathbb{B}_n(x) := \sqrt{n}(\hat{F}_n(x) - F_X(x)).$$

Com base na função de distribuição F_X (conjecturada) e na função $w : (0, 1) \mapsto (0, +\infty)$, chamada de *função peso*, definimos o seguinte produto interno.

$$\langle F, G \rangle_{\text{AD}} := \int_{\mathbb{R}} F(x)G(x)w(F_X(x))dF_X(x) \quad (\text{C.22})$$

em que as funções F e G são contínuas a direita. Em geral é comum definirmos normas associadas a produtos internos. Considere $\mathfrak{B}$ o espaço dos processos empíricos $\mathbb{B}_n$, então definimos a norma $\|\cdot\|_{\text{AD}} : \mathfrak{B} \mapsto [0, +\infty)$ sobre esse espaço como

$$\|\mathbb{B}_n\|_{\text{AD}} := \sqrt{\langle \mathbb{B}_n, \mathbb{B}_n \rangle_{\text{AD}}} \quad (\text{C.23})$$

Com esse objeto podemos definir a estatística que iremos trabalhar nessa subseção.

Definição C.13 (Estatística de Anderson-Darling). A estatística $T_n^{\text{AD}}(X)$ definida por

$$T_n^{\text{AD}}(X) := \|\mathbb{B}_n\|_{\text{AD}}^2 \quad (\text{C.24})$$

é chamada de Anderson-Darling.

Pela fórmula do processo empírico $\mathbb{B}_n$ também podemos escrever a estatística (C.24) do seguinte modo

$$T_n^{AD}(X) = n \|\hat{F}_n - F_X\|_{AD}^2$$

o que torna explícito o fato de estarmos comparando a distância entre a função de distribuição empírica e da função de distribuição conjecturada para os dados. Para compreendermos um pouco melhor a estrutura dessa estatística vamos investigar o seu comportamento para uma função de distribuição simples. Considere por exemplo que $X \sim \text{Unif}[0, 1]$ então sua função de distribuição é dada por

$$F_X(x) = \begin{cases} 0 & x < 0 \\ x & 0 \leq x < 1 \\ 1 & x \geq 1 \end{cases}$$

Considerando a função peso constante $w(u) := 1$ (nesse caso (C.24) recebe o nome de *Cramer-von Mises*) temos que

$$T_n^{AD}(X) = \int_0^1 \mathbb{B}_n^2(x) dx = n \int_0^1 (\hat{F}_n(x) - F_X(x))^2 dx.$$

Dado que a integral pode ser entendida como uma “soma” podemos olhar a expressão acima da seguinte forma heurística: a estatística de Anderson-Darling em um caso simples é uma “soma” dos quadrados das diferenças entre a função de distribuição empírica e a função de distribuição conjecturada. Note que com essa interpretação percebemos que a estatística de Anderson-Darling busca punir mais as discrepâncias entre as funções $\hat{F}_n$ e F_X . Vamos ver agora o que acontece quando usamos funções peso mais elaboradas, como exemplo, a função

$$w(u) := \frac{1}{u(1-u)}. \quad (\text{C.25})$$

Essa função peso foi originalmente proposta na primeira definição da estatística de Anderson-Darling. No nosso exemplo simples da distribuição uniforme temos que

$$\begin{aligned} T_n^{AD}(X) &= n \int_0^1 (\hat{F}_n(x) - F_X(x))^2 w(F_X(x)) dx = n \int_0^1 (\hat{F}_n(x) - F_X(x))^2 w(x) dx = \\ &= n \int_0^1 \frac{(\hat{F}_n(x) - F_X(x))^2}{x(1-x)} dx \end{aligned}$$

Ou seja,

$$T_n^{AD}(X) = n \int_0^1 \frac{(\hat{F}_n(x) - F_X(x))^2}{x(1-x)} dx.$$

Observe que para a integral acima fazer sentido estamos considerando w definido em $(0, 1)$. Percebemos então que a interpretação heurística se mantém mais agora com um fator adicional pois para os pontos de $(0, 1)$ próximos de 0 ou 1 a função peso w assume valores muito grandes e portanto esses pontos serão considerados mais relevantes. Isso implica que a estatística de Anderson-Darling olhará a discrepância entre as funções $\hat{F}_n$ e F_X de modo mais “exigente ou severo” nesses pontos.

Apesar do exemplo dado no parágrafo anterior muito simples ele é bastante instrutivo pois mostra algumas características importantes da estatística de Anderson-Darling. Fazendo uma comparação simples com a estatística de Kolmogorov-Smirnov é possível perceber que o modo como a Anderson-Darling avalia as discrepâncias entre as funções $\hat{F}_n$ e F_X é mais sofisticado do que Kolmogorov-Smirnov pois nessa última estatística olhamos apenas para o “máximo” (pois estamos usando o supremo) da distância entre as funções $\hat{F}_n$ e F_X segundo o valor absoluto sem nos

importar em caracterizar o ponto x em que esse “máximo” foi observado. Por outro lado esse olhar mais sofisticado tem um preço pois a princípio poderíamos nos perguntar como efetivamente podemos utilizar a estatística de Anderson-Darling no processo de inferência. No caso da função peso (C.25) é possível obter uma expressão mais simples. Considere $(U_i)_{0 \leq i \leq n}$ uma sequência de variáveis aleatórias definidas por

$$U_i := F_X(X_i)$$

Sabemos da teoria das probabilidades que essas variáveis tem distribuição uniforme no intervalo $[0, 1]$. Agora seja $(U_{(i)})_{0 \leq i \leq n}$ a *estatística de ordem* associada a sequência $(U_i)_{0 \leq i \leq n}$ então é possível provar que a seguinte fórmula vale.

$$T_n^{AD}(X) = -n - \frac{1}{n} \sum_{i=1}^n \left((2i-1) \log(U_{(i)}) + (2n+1-2i) \log(1-U_{(i)}) \right). \quad (\text{C.26})$$

Mesmo com um expressão mais amigável para a estatística de Anderson-Darling obtenção a sua distribuição de probabilidade enfrenta muita dificuldades mesmo no caso assintótico. De fato, isso é algo que poderíamos esperar em função da complexidade da distribuição de probabilidade da estatística de Kolmogorov-Smirnov mesmo essa sendo bem mais simples. Mesmo assim aproximações da distribuição assintótica podem ser obtidas usando de forma mais elaborada as ideias apresentadas na seção sobre processos empíricos o que possibilita a implementação computacional do teste e de fato é possível encontrar bons pacotes estatísticos que oferecem uma implementação dessa estatística. Por fim um aspecto que gostaríamos de destacar é o poder do teste de Anderson-Darling frente a outros testes como o Kolmogorov-Smirnov. Na literatura sobre a teoria da bondade do ajuste é destacado que empiricamente se observa em geral que o teste de Anderson-Darling é mais poderoso do que o teste de Kolmogorov-Smirnov. Isso significa, conforme estudado na subseção sobre poder de um teste, que o teste de Anderson-Darling é capaz de perceber melhor as diferenças sutis entre a função de distribuição conjecturada na hipótese nula e qualquer outra função de distribuição que atenda a hipótese alternativa dado um determinado nível de significância.

C.6 Testes para duas amostras.

Em geral quando um estatístico decide usar as ferramentas disponíveis na teoria da bondade do ajuste esse encontra-se em uma situação de total desconhecimento sobre de qual família de distribuições a amostra que tem em mãos é proveniente. Nesse sentido os testes apresentados para uma amostra apesar de serem importantes são de difícil utilização prática. Ao longo desta seção vamos reformular os testes apresentados na seção anterior para o caso de duas amostras. Seguindo a mesma ordem adotada para os testes de uma amostra começaremos pelo teste de Kolmogorov-Smirnov.

C.6.1 Teste de Kolmogorov-Smirnov para duas amostras.

O teste de Kolmogorov-Smirnov para duas amostras segue uma lógica parecida ao da versão para uma amostra mas agora com a diferença que na impossibilidade de se conjecturar uma distribuição de probabilidade para os dados e assim compará-la com a distribuição empírica vamos buscar comparar duas funções de distribuição empíricas em que cada uma é definida sobre uma amostra aleatória independente da outra. Para ver como isso é feito observe que se as amostras independentes $X^{(1)}$ e $X^{(2)}$ são provenientes da mesma função de distribuição F_X então temos pela consistência forte (teorema C.4) do estimador dado pela função de distribuição empírica que esse converge uniformemente para F_X . Ou seja,

$$\lim_{n_1 \rightarrow +\infty} \sup_{x \in \mathbb{R}} \{ |\hat{F}_{n_1}^{(1)}(x) - F_X(x)| \} = 0 \text{ e } \lim_{n_2 \rightarrow +\infty} \sup_{x \in \mathbb{R}} \{ |\hat{F}_{n_2}^{(2)}(x) - F_X(x)| \} = 0 \quad \text{q.c.}$$

Consequentemente segue-se dos limites acima que

$$\lim_{n_1 \rightarrow +\infty} \lim_{n_2 \rightarrow +\infty} \sup_{x \in \mathbb{R}} \{|\hat{F}_{n_1}^{(1)}(x) - \hat{F}_{n_2}^{(2)}(x)|\} = 0 \quad \text{q.c.}$$

O limite acima de certa forma se parasse com o critério de Cauchy para a convergência uniforme de seqüências de funções que é utilizado justamente quando não sabemos quem é a função limite. Além disso assim como no caso de uma amostra temos que o fator envolvendo a raiz quadrada que aparece antes do supremo serve justamente para garantir a estabilidade assintótica dessa estatística e assim seja possível obter a sua distribuição assintótica. Passamos então a definição da estatística do teste de Kolmogorov-Smirnov para duas amostras.

Definição C.14 (Estatística de Kolmogorov-Smirnov para duas amostras). Considere $X^{(1)} := (X_i^{(1)})_{0 \leq i \leq n_1}$ e $X^{(2)} := (X_i^{(2)})_{0 \leq i \leq n_2}$ duas amostras aleatórias independentes com funções de distribuição $F_1 : \mathbb{R} \mapsto [0, 1]$ e $F_2 : \mathbb{R} \mapsto [0, 1]$ respectivamente. Considere também $\hat{F}_{n_1}^{(1)} : \mathbb{R} \mapsto [0, 1]$ e $\hat{F}_{n_2}^{(2)} : \mathbb{R} \mapsto [0, 1]$ as funções de distribuição empíricas associadas as amostras $X^{(1)}$ e $X^{(2)}$ respectivamente. Definimos a estatística de Kolmogorov-Smirnov para duas amostras por

$$T_{n_1, n_2}^{KS}(X^{(1)}, X^{(2)}) := \sqrt{\frac{n_1 n_2}{n_1 + n_2}} \sup_{x \in \mathbb{R}} \{|\hat{F}_{n_1}^{(1)}(x) - \hat{F}_{n_2}^{(2)}(x)|\} \quad (\text{C.27})$$

Note que a estatística dada na equação (C.27) está bem em linha com a intuição apresentada na introdução dessa seção. Como sabemos a hipótese nula a ser considerada é $H_0 : F_1 = F_2$ e portanto o teste de kolmogorov-Smirnov buscará coerentemente descartar a hipótese nula se a estatística (C.27) for maior que um determinado valor crítico que depende do nível de significância.

Apesar da estatística de Kolmogorov-Smirnov ter ficado um pouco mais complexa para o caso de duas amostras temos felizmente que a distribuição de probabilidade assintótica ainda pode ser calculada.

Lema C.15 (Distribuição assintótica da estatística de Kolmogorov-Smirnov para duas amostras: Ver referência [Sha03] página 449 comentários antes da seção 6.5.3.). *Nas mesmas condições da definição C.14 temos que o seguinte limite existe.*

$$\lim_{n_1, n_2 \rightarrow +\infty} P\left(T_{n_1, n_2}^{KS}(X^{(1)}, X^{(2)}) \leq t\right) = \sum_{i=-\infty}^{+\infty} (-1)^{i-1} e^{-2i^2 t^2} \quad (\text{C.28})$$

com $t > 0$.

Como observamos a utilização prática da distribuição assintótica (C.28) passa por aproximá-la por uma soma finita com um número adequado de termos.

C.6.2 Teste de Anderson-Darling para 2-amostras.

A extensão do teste de Anderson-Darling para o caso com duas amostra não tão direto com no caso do teste de Kolmogorov-Smirnov. Isso ocorre em função da fórmula do produto interno que define essa estatística depender da função de distribuição conjecturada para os dados. Como para duas amostras estamos em uma situação de total desconhecimento sobre a classe das distribuições de probabilidade a qual a distribuição dos dados pertence a fórmula para esse produto interno precisa ser alterada e dado que tudo que temos são duas amostras independentes a margem para essa alteração é um pouco apertada.

Inicialmente poderíamos considerar a possibilidade de substituir a função de distribuição conjecturada por uma das funções de distribuição empíricas em razão da hipótese nula afirmar que os dados que temos são provenientes da mesma amostra e como sempre calculamos a distribuição de probabilidade assintótica da estatística de Anderson-Darling sob a hipótese nula ser verdadeira (na verdade calculamos as distribuições assintóticas sempre desse modo) isso a princípio poderia resolver o problema. A questão é que a estatística deve se definida também levando em conta o que é afirmado na hipótese alternativa que no nosso caso diz que temos duas distribuições distintas para as amostras e não uma. Uma forma de contorna essas dificuldades é através de uma função de distribuição empírica que é a combinação convexa das distribuições empíricas de cada uma das amostras.

Definição C.16 (Função de distribuição empírica de agrupamento). Considere $X^{(1)} := (X_i^{(1)})_{0 \leq i \leq n_1}$ e $X^{(2)} := (X_i^{(2)})_{0 \leq i \leq n_2}$ duas amostras aleatórias independentes com funções de distribuição $F_1 : \mathbb{R} \mapsto [0, 1]$ e $F_2 : \mathbb{R} \mapsto [0, 1]$ respectivamente. Considere também $\hat{F}_{n_1}^{(1)} : \mathbb{R} \mapsto [0, 1]$ e $\hat{F}_{n_2}^{(2)} : \mathbb{R} \mapsto [0, 1]$ as funções de distribuição empíricas associadas as amostras $X^{(1)}$ e $X^{(2)}$ respectivamente. Definimos a função de distribuição empírica de agrupamento para duas amostra por

$$\hat{H}_{n_1, n_2}(x) = \frac{n_1}{n_1 + n_2} \hat{F}_{n_1}^{(1)}(x) + \frac{n_2}{n_1 + n_2} \hat{F}_{n_2}^{(2)}(x) \quad (\text{C.29})$$

Para ver como a função definida na equação (C.29) atende as condições que identificamos nos parágrafos iniciais dessa subseção vamos definir a versão não empírica dessa função

$$H_{n_1, n_2}(x) = \frac{n_1}{n_1 + n_2} F_1(x) + \frac{n_2}{n_1 + n_2} F_2(x) \quad (\text{C.30})$$

Agora note que sob a hipótese nula a função de distribuição H_{n_1, n_2} (pois é a combinação convexa de funções de distribuição) é simplesmente igual a F_1 e F_2 . No contexto da hipótese alternativa olhamos para uma combinação convexa das distribuições que temos com ponderação dada pela proporção das amostras que possuímos de cada uma das distribuições. Por outro lado como as funções de distribuição empíricas convergem uniformemente e quase certamente para as funções de distribuição populacionais (conjecturadas) temos que

$$\begin{aligned} & \left| \hat{H}_{n_1, n_2}(x) - H_{n_1, n_2}(x) \right| = \\ & = \left| \frac{n_1}{n_1 + n_2} \hat{F}_{n_1}^{(1)}(x) + \frac{n_2}{n_1 + n_2} \hat{F}_{n_2}^{(2)}(x) - \left(\frac{n_1}{n_1 + n_2} F_1(x) + \frac{n_2}{n_1 + n_2} F_2(x) \right) \right| = \\ & = \left| \frac{n_1}{n_1 + n_2} (\hat{F}_{n_1}^{(1)}(x) - F_1(x)) + \frac{n_2}{n_1 + n_2} (\hat{F}_{n_2}^{(2)}(x) - F_2(x)) \right| \leq \\ & \leq \frac{n_1}{n_1 + n_2} |\hat{F}_{n_1}^{(1)}(x) - F_1(x)| + \frac{n_2}{n_1 + n_2} |\hat{F}_{n_2}^{(2)}(x) - F_2(x)| \leq \\ & \leq |\hat{F}_{n_1}^{(1)}(x) - F_1(x)| + |\hat{F}_{n_2}^{(2)}(x) - F_2(x)| \leq \\ & \leq \sup_{x \in \mathbb{R}} \{ |\hat{F}_{n_1}^{(1)}(x) - F_1(x)| \} + \sup_{x \in \mathbb{R}} \{ |\hat{F}_{n_2}^{(2)}(x) - F_2(x)| \}. \end{aligned}$$

Ou seja,

$$\left| \hat{H}_{n_1, n_2}(x) - H_{n_1, n_2}(x) \right| \leq \sup_{x \in \mathbb{R}} \{ |\hat{F}_{n_1}^{(1)}(x) - F_1(x)| \} + \sup_{x \in \mathbb{R}} \{ |\hat{F}_{n_2}^{(2)}(x) - F_2(x)| \}.$$

Logo

$$\sup_{x \in \mathbb{R}} \{|\hat{H}_{n_1, n_2}(x) - H_{n_1, n_2}(x)|\} \leq \sup_{x \in \mathbb{R}} \{|\hat{F}_{n_1}^{(1)}(x) - F_1(x)|\} + \sup_{x \in \mathbb{R}} \{|\hat{F}_{n_2}^{(2)}(x) - F_2(x)|\}.$$

Aplicando o teorema C.4 (Glivenko-Cantelli) na desigualdade acima obtemos que

$$\lim_{n_1, n_2 \rightarrow +\infty} \lim_{n_2 \rightarrow +\infty} \sup_{x \in \mathbb{R}} \{|\hat{H}_{n_1, n_2}(x) - H_{n_1, n_2}(x)|\} = 0 \text{ q.c.}$$

Portanto o estimador $\hat{H}_{n_1, n_2}$ é fortemente consistente. Como a princípio deveríamos substituir a função de distribuição F_X no produto interno $\langle \cdot, \cdot \rangle_{AD}$ na formulação para uma amostra pela função H_{n_1, n_2} para o caso de duas amostras, mas a única informação que temos são as duas amostras independentes $X^{(1)} = (X_i^{(1)})_{0 \leq i \leq n_1}$ e $X^{(2)} = (X_i^{(2)})_{0 \leq i \leq n_2}$ vamos definir então um versão empírica do produto interno $\langle \cdot, \cdot \rangle_{AD}$ tendo por base o estimador $\hat{H}_{n_1, n_2}$.

$$\langle F, G \rangle_{AD\text{-amostral}} := \int_{\mathbb{R}} F(x)G(x)w(\hat{H}_{n_1, n_2}(x))d\hat{H}_{n_1, n_2}(x) \quad (C.31)$$

Agora que já temos um produto interno é preciso alterar ligeiramente a formulação dos processos empíricos para uma classe um pouco maior de processos chamados de *processos de contraste*. Basicamente um processo de contraste é um processo empírico em que substituímos a função de distribuição conjecturada F_X por uma segunda função de distribuição empírica para uma outra amostra de X independe da primeira.

Definição C.17 (Processos de contraste). Nas mesmas condições da definição C.16 chamamos o processo estocástico $(\mathbb{C}_{n_1, n_2}(x))_{x \in \mathbb{R}}$ definido por

$$\mathbb{C}_{n_1, n_2}(x) := \sqrt{\frac{n_1 n_2}{n_1 + n_2}} (\hat{F}_{n_1}^{(1)}(x) - \hat{F}_{n_2}^{(2)}(x)) \quad (C.32)$$

Com o produto interno em mãos podemos seguir basicamente os mesmos passos que fizemos para definir a estatística de Anderson-Darling para uma amostra. Passamos então para a definição da norma no espaço dos processos de contraste.

$$\|\mathbb{C}_{n_1, n_2}\|_{AD\text{-amostral}} := \sqrt{\langle \mathbb{C}_{n_1, n_2}, \mathbb{C}_{n_1, n_2} \rangle_{AD\text{-amostral}}} \quad (C.33)$$

Com esse objeto podemos definir a estatística que iremos trabalhar nessa subseção.

Definição C.18 (Estatística de Anderson-Darling para duas amostras). A estatística $T_{n_1, n_2}^{AD}(X^{(1)}, X^{(2)})$ definida por

$$T_{n_1, n_2}^{AD}(X^{(1)}, X^{(2)}) := \|\mathbb{C}_{n_1, n_2}\|_{AD\text{-amostral}}^2 \quad (C.34)$$

é chamada de Anderson-Darling para duas amostras

Assim como no caso para uma amostra temos pela fórmula do processo de contraste $\mathbb{C}_{n_1, n_2}$ que podemos escrever a estatística (C.34) do seguinte modo

$$T_{n_1, n_2}^{AD}(X^{(1)}, X^{(2)}) = \frac{n_1 n_2}{n_1 + n_2} \|\hat{F}_{n_1}^{(1)} - \hat{F}_{n_2}^{(2)}\|_{AD\text{-amostral}}^2.$$

Nessa formulação vemos claramente que a estatística de Anderson-Darling para duas amostras busca comparar a discrepância entre as duas distribuições empíricas segundo uma métrica empírica baseada nos próprios dados disponíveis. Um outro aspecto importante a ser observado é que em

geral a função peso w adotada para o caso de duas amostras é a mesma que no caso anterior, então boa parte das intuições que desenvolvemos para o caso com uma amostra continua valendo inclusive as observações sobre o poder do teste de Anderson-Darling em relação ao poder do Kolmogorov-Smirnov.

Pelos desenvolvimentos feitos até aqui temos que o procedimento de teste para o Anderson-Darling duas amostras é o mesmo que fizemos para o Kolmogorov-Smirnov duas amostras, mas infelizmente ao contrário desse último não é tão simples obter uma expressão para a distribuição assintótica da estatística do Anderson-Darling e portanto somos forçados a trabalhar apenas com aproximações dessa. Felizmente existe boas implementações dessas aproximações nos principais pacotes estatísticos dedicados aos teste da bondade do ajuste.

Referências

- [BBD⁺21] Jeff Bezanson, Ian Butterworth, Nathan Daly, Keno Fischer, Jameson Nash, Tim Holy, Elliot Saba, Mosè Giordano, Stefan Karpinski, and Kristoffer Carls-son. Julia 1.6 highlights. <https://julialang.org/blog/2021/03/julia-1.6-highlights/>, 2021.
- [Ber12] Dimitri P Bertsekas. *Dynamic programming and optimal control: Approximate dynamic programming*. Athena Scientific, 4th edition, 2012.
- [BZ07] John Birge and Gongyun Zhao. Successive linear approximation solution of infinite-horizon dynamic stochastic programs. *SIAM Journal on Optimization*, 18(4):1165–1186, January 2007.
- [con20] GLPK contributors. Gnu linear programming kit, version 5.0, 2020.
- [con22] COIN-OR contributors. coin-or/clp: Release 1.17.7, 2022.
- [dCdFK⁺19] Bernardo Freitas Paulo da Costa, Iago Leal de Freitas, Rafael Benchimol Klaus-ner, Filipe Goulart Cabral, Joari Paulo da Costa, and Débora Dias Jardim Penna. Aceleração de convergência para problemas de otimização estocástica multiestágio. Relatório técnico, Convênio ONS – UFRJ, 2019. Relink – Rela-tório Final.
- [Dow20] Oscar Dowson. The policy graph decomposition of multistage stochastic pro-gramming problems. *Networks*, 76(1):3–23, February 2020.
- [FPdCLdFCdC18] Bernardo Freitas Paulo da Costa, Iago Leal de Freitas, Filipe Goulart Cabral, and Joari Paulo da Costa. Uso de restrições disjuntivas para modelar políti-cas operativas. Relatório técnico, Convênio ONS – UFRJ, 2018. Restrições Disjuntivas – Relatório Final.
- [HH18] Q. Huangfu and J. A. J. Hall. Parallelizing the dual revised simplex method. *Mathematical Programming Computation*, 10(1):119–142, 2018.
- [LL] Artur O. Lopes and Sílvia R. C. Lopes. Introdução aos processos estocásti-cos para estudantes de matemática. <http://mat.ufrgs.br/~alopes/pub3/Livro-Proc-Estocas.pdf>.
- [LR05] Erich L. Lehmann and Joseph P. Romano. *Testing Statistical Hypotheses*. Spring-er, 3 edition, 2005.
- [MT09] Sean Meyn and Richard Tweedie. *Markov Chains and Stochastic Stability*. Cam-bridge Press, 2nd edition, 2009.
- [SD20] Alexander Shapiro and Lingquan Ding. Periodical multistage stochastic pro-grams. *SIAM Journal on Optimization*, 30:2083–2102, August 2020.
- [SDR09] Alexander Shapiro, Darinka Dentcheva, and Andrzej Ruszczyński. *Lectures on stochastic programming: modeling and theory*. SIAM, 2009.
- [Sha03] Jun Shao. *Mathematical Statistics*. Springer, 2 edition, 2003.

- [STdCS13] Alexander Shapiro, Wajdi Tekaya, Joari Paulo da Costa, and Murilo Pereira Soares. Risk neutral and risk averse stochastic dual dynamic programming method. *European journal of operational research*, 224(2):375–391, 2013.