



Universidade do Estado do Rio de Janeiro

Centro de Tecnologia e Ciências

Instituto de Matemática e Estatística

Fabio Cordeiro Carvalho

**Probabilidade e Estatística: Um guia acerca da sua importância nas mais
diversas áreas do conhecimento**

Rio de Janeiro

2024

Fabio Cordeiro Carvalho

Probabilidade e Estatística: Um guia da sua importância nas mais diversas áreas do conhecimento.



Dissertação apresentada, como requisito parcial para obtenção do título de Mestre, ao Programa de Mestrado Profissional em Matemática em Rede Nacional (PROFMAT), da Universidade do Estado do Rio de Janeiro.

Orientadora: Prof.^a Dra. Patrícia Nunes da Silva

Rio de Janeiro

2024

CATALOGAÇÃO NA FONTE
UERJ/REDE SIRIUS/BIBLIOTECA CTC/A

C331 Carvalho, Fabio Cordeiro
Probabilidade e estatística: um guia da sua importância nas mais diversas áreas do conhecimento/ Fabio Cordeiro Carvalho. – 2024.
97 f. : il.

Orientadora: Patricia Nunes da Silva
Dissertação (Mestrado Profissional em Matemática em Rede Nacional - PROFMAT) - Universidade do Estado do Rio de Janeiro, Instituto de Matemática e Estatística.

1. Estatística - Teses. 2. Probabilidade - Teses. I. Silva, Patricia Nunes da. II. Universidade do Estado do Rio de Janeiro. Instituto de Matemática e Estatística. III. Título.

CDU 519.2

Patricia Bello Meijinhos CRB7/5217 - Bibliotecária responsável pela elaboração da ficha catalográfica

Autorizo, apenas para fins acadêmicos e científicos, a reprodução total ou parcial desta dissertação, desde que citada a fonte.

Assinatura

Data

Fabio Cordeiro Carvalho

Probabilidade e Estatística: Um guia da sua importância nas mais diversas áreas do conhecimento

Dissertação apresentada, como requisito parcial para obtenção do título de Mestre, ao Programa de Mestrado Profissional em Matemática em Rede Nacional (PROFMAT), da Universidade do Estado do Rio de Janeiro.

Aprovada em 30 de setembro de 2024.

Banca Examinadora:

Prof.^a Dra. Patrícia Nunes da Silva (Orientadora)
Instituto de Matemática e Estatística - UERJ

Prof.^a Dra. Cláudia Ferreira Reis Concordido
Instituto de Matemática e Estatística -UERJ

Prof. Dr. Fernando Antônio de Araújo Carneiro
Instituto de Matemática e Estatística - UERJ

Prof. Dr. Ronaldo da Silva Busse
Universidade Federal do Estado do Rio de Janeiro - UNIRIO

Rio de Janeiro

2024

DEDICATÓRIA

À Deus, pois sem Ele seria impossível a conclusão deste trabalho. “Pois dele, por ele, e para ele são todas as coisas. A ele seja a glória para sempre” (BÍBLIA, Rm, 11, 36)

AGRADECIMENTOS

Gostaria de expressar meus sinceros agradecimentos a todas as pessoas que tornaram possível a realização deste trabalho:

Primeiramente, agradeço a Deus, Rei e Senhor do universo, que me fortaleceu nos momentos de dificuldade e me concedeu a sabedoria necessária para concluir este trabalho.

Aos meus pais, devo tudo à educação e apoio incondicionais que sempre me proporcionaram. Sem o suporte deles, não teria chegado até aqui.

À minha esposa, meu pilar e companheira incansável, que esteve ao meu lado em cada etapa deste processo, oferecendo seu apoio e incentivo valiosos.

Aos meus amigos mestrando, que compartilharam comigo suas experiências e sempre estenderam sua mão amiga nos momentos de necessidade.

À CAPES, pela oportunidade e apoio financeiro fornecidos para a realização deste trabalho. O presente trabalho foi realizado com apoio da Coordenação de Aperfeiçoamento de Pessoal de Nível Superior – Brasil (CAPES) – Código de Financiamento 001.

A todos os professores da UERJ e em especial, a minha orientadora Patrícia Nunes da Silva, minha profunda gratidão por tudo, orientação e ensinamentos e principalmente por acreditar em mim no momento que eu mais precisei.

As barreiras do impossível caem quando você coloca Deus na sua vida.

Billy Graham

RESUMO

CARVALHO, Fabio Cordeiro . Probabilidade e Estatística: um guia de sua importância nas mais diversas áreas do conhecimento 2024. 94 f. Dissertação (Mestrado Profissional em Matemática em Rede Nacional - PROFMAT) – Instituto de Matemática e Estatística, Universidade do Estado do Rio de Janeiro, Rio de Janeiro, 2024

Conceitos matemáticos de probabilidade e estatística são a temática principal deste trabalho. Especificamente, busca-se interligar as definições formais dessas áreas com as suas aplicações no cotidiano. Para tanto, o livro “O Andar do Bêbado”, de Leonard Mlodinow, foi tomado como base por apresentar diversas aplicações do tema no dia a dia. Detalhar, discutir e preencher lacunas nos problemas propostos por Mlodinow proporciona uma visão mais aprofundada do uso da probabilidade e estatística em diversas áreas do conhecimento. Na discussão dos problemas selecionados, além da apresentação dos cálculos matemáticos omitidos, o contexto geral dos problemas foi estendido e aprofundado. O foco principal é mostrar como a probabilidade e a estatística estão presentes em nossa vida e como, por diversas vezes, o “senso comum” leva a erros acerca do acaso. Esses equívocos podem ser de natureza matemática, psicológica ou até mesmo oriundos da forma como os problemas são apresentados. Com isso, esse estudo visa tanto servir como um “guia básico” sobre probabilidade, quanto abordar de forma mais aprofundada os temas, mesclando problemas voltados para o público mais leigo com aqueles que requerem um conhecimento matemático mais avançado.

Palavras-chave: probabilidade; estatística; espaço Amostral; “Andar do Bêbado.”

ABSTRACT

CARVALHO, Fabio Cordeiro. *A guide to its importance in the most diverse areas of knowledge* 2024. 93 f. Dissertação (Mestrado Profissional em Matemática em Rede Nacional - PROFMAT) – Instituto de Matemática e Estatística, Universidade do Estado do Rio de Janeiro, Rio de Janeiro, 2024

Mathematical concepts of probability and statistics are the main theme of this work. Specifically, the aim is to connect the formal definitions of these areas with their applications in everyday life. To this end, the book “The Drunkard’s Walk, by Leonard Mlodinow, was used as a basis because it presents many applications in everyday life. Detailing, discussing and filling gaps in the problems proposed by Mlodinow provides a more depth view of probability and statistics in many areas of knowledge. In the discussion of the selected problems, in addition to the presentation of omitted mathematical calculations, the general context of the problems was extended and deepened. The main focus is to show how probability and statistics are presented in our lifes and how, on many occasions, “common sense” leads to errors about chance. These errors can be of a mathematical or psychological nature or even originate from the way the problems are present. With this, this study aims to serve both as a “basic guideline on probability” and to adress the topics in more depth, combining problems aimed at the more advanced mathematical knowledge.

Keywords: probability; statistics; sample space; Drunkard’s Walk.

LISTA DE FIGURAS

Figura 1 -	Espaço Amostral do experimento 3 bolas em 3 caixas	29
Figura 2 -	Portas da Ilusão.	40
Figura 3 -	Relação da sensibilidade e especificidade com o PPV	57
Figura 4 -	Estimativa dada pelos médicos nos 2 modelos	61
Figura 5 -	Representação da estrutura do DNA	45
Figura 6 -	Ilustração de um campo de beisebol	63
Figura 7 -	Desempenho dos desafiantes em relação a Babe Ruth	65
Figura 8 -	Desempenho dos desafiantes em anos anteriores	67

LISTA DE TABELAS

Tabela 1 -	Pergunta alvo e a sua heurística correspondente.	16
	. . .	
Tabela 2 -	Proporção de nascimento nos 2 hospitais.	19
Tabela 3 -	Relação entre as afirmações sobre Linda e sua chance de ocorrência . . .	24
Tabela 4 -	Relação entre o número de portas abertas e as probabilidades para 10 portas.	42
Tabela 5 -	Relação entre as dezenas apostadas e a probabilidade de vitória	48
Tabela 6 -	Relação entre o resultado dos exames, a sensibilidade e a especificidade . .	56
Tabela 7 -	Probabilidade de o exame acusar um falso positivo ao longo dos anos .	62
Tabela 8 -	Relação entre a taxa de falso positivo e as probabilidades a priori e a posteriori	67

SUMÁRIO

INTRODUÇÃO.....	13
1. PROBABILIDADE E ESTATÍSTICA: UMA VISÃO ATRAVÉS DO VIÉS DA PSICOLOGIA	15
1.1 Heurística da Disponibilidade	17
1.2 Heurística de Representatividade	17
1.2.1 <u>Insensibilidade a dados anteriores sobre a probabilidade de ocorrência de um determinado evento</u>	18
1.2.2 <u>Insensibilidade ao tamanho da amostra</u>	18
1.2.3 <u>Equívocos acerca da probabilidade</u>	20
1.3 Por que sempre nos achamos tão azarados?	20
1.4 O problema das palavras	22
1.4.1 <u>Mas por que as pessoas acham que o conjunto B tem mais elementos?</u>	22
1.5 Será Linda feminista?	24
2. PROBABILIDADE E ESPAÇO AMOSTRAL	26
2.1 Espaço Amostral: Uma breve abordagem histórica	26
2.2 Algumas definições importantes de probabilidade	30
2.2.1 <u>Relação entre dois eventos</u>	30
2.2.2 <u>Propriedades básicas de probabilidade</u>	31
2.3 Problema de D'Alembert	34
2.4 O problema dos pneus	34
2.5 É menino ou menina?	36
2.6 O problema de Monty- Hall	38
3. PROBABILIDADE E JOGOS DE AZAR	43
3.1 Contexto Histórico	43
3.2 Os astrágalos e a Grécia Antiga	44
3.3 Devo estacionar aqui?	45
3.4 Será que é mesmo vantajoso?	46
3.5 O Brasil e os jogos de aposta	47
3.6 A loteria de Virgínia	49
3.7 Problema dos pontos de Pascal: como dividir a aposta de um jogo que não acabou	51
4. PROBABILIDADE E EXAMES CLINICOS	54

4.1	Sensibilidade	55
4.2	Especificidade	56
4.3	Diagnósticos Médicos	58
4.3.1	<u>Mlodinow e o teste de HIV</u>	62
4.4	Exames Antidopings	65
5.	O DNA E O DIREITO	67
5.1	Introdução ao conceito de DNA	67
5.2	Uma correspondência relatada é realmente verdadeira?	69
6.	PROBABILIDADE E OS ESPORTES	77
6.1	O beisebol e o recorde de Babe Ruth	77
6.1.1	<u>A corrida de Maris e Mantle em busca do recorde</u>	79
6.1.2	<u>Mas será que foi isso mesmo que aconteceu ou Maris por puro “acaso do destino” teve uma temporada discrepante do seu desempenho médio?</u>	84
6.2	O futebol e a chance de um time ser campeão	87
	CONCLUSÃO	92
	REFERÊNCIAS	93

INTRODUÇÃO

Na vida cotidiana, muitas vezes os indivíduos se deparam com situações envolvendo incerteza e aleatoriedade, desde aspirar a ganhar na “Mega-Sena” até lidar com imprevistos como engarrafamentos ou filas de supermercado. A compreensão correta da probabilidade desses eventos é essencial, entretanto, frequentemente se cometem equívocos em sua estimativa. A esse respeito, Mlodinow (2009) em sua obra “O Andar do Bêbado”, investiga esses problemas, destacando como o conceito de probabilidade é frequentemente mal interpretado, seja pela falta de conhecimento na área ou por influências dos mecanismos cognitivos do cérebro humano.

Este trabalho, inspirado pela obra de Mlodinow, busca aprofundar a compreensão dos problemas discutidos no livro, fornecendo explicações detalhadas e elucidativas. A abordagem visa não apenas identificar os erros comuns cometidos pelas pessoas ao fazerem julgamentos probabilísticos, mas também compreender os mecanismos matemáticos e cognitivos subjacentes a esses equívocos. É crucial reconhecer que os cérebros humanos são suscetíveis a vieses e heurísticas que podem distorcer a percepção humana da probabilidade. Sendo assim, além de explorar os aspectos matemáticos desses problemas, é necessário examinar também os processos cognitivos envolvidos (BASTOS, 2004). Ao identificar esses mecanismos e propor soluções, pretende-se capacitar os leitores a tomar decisões mais informadas e precisas diante da incerteza e da aleatoriedade. Por meio de uma análise minuciosa, o objetivo é proporcionar aos leitores uma compreensão mais clara da influência da aleatoriedade em suas vidas e como melhor lidar com ela.

Para uma organização mais sistemática dos temas abordados no livro, o presente trabalho será dividido em seis capítulos, cada um dedicado a um tópico específico:

No Capítulo 1, intitulado “Probabilidade e Psicologia”, será explorada a relação entre a aleatoriedade e os vieses cognitivos. Este capítulo não abordará conceitos matemáticos de probabilidade, mas sim como o cérebro humano desenvolve mecanismos que podem distorcer nossa percepção da real chance de ocorrência de eventos.

O Capítulo 2, “Probabilidade e Espaço Amostral”, começará com uma breve análise histórica da evolução do conceito de espaço amostral ao longo do tempo. Em seguida, serão apresentadas as definições clássica e frequentista de probabilidade, juntamente com conceitos fundamentais relacionados à aleatoriedade no cotidiano.

No Capítulo 3, denominado “Probabilidade e Jogos de Azar”, será discutida a relação intrínseca entre o acaso e os jogos de azar, desde tempos antigos até os dias atuais, destacando

como essa relação foi uma das principais motivações para o desenvolvimento da teoria da probabilidade.

O Capítulo 4, “Probabilidade e Exames Médicos”, abordará conceitos relevantes sobre testes clínicos, como especificidade e sensibilidade, e discutirá como a falta de informação tanto do público em geral quanto dos próprios médicos pode levar a uma compreensão equivocada da infalibilidade desses testes.

No Capítulo 5, “DNA e Direito”, será analisado o uso de testes de DNA no contexto jurídico, buscando desmistificar a crença de que esses testes são completamente infalíveis e examinando questões éticas e legais associadas a seu uso.

Por fim, no Capítulo 6, “Probabilidade e Esportes”, serão abordadas questões probabilísticas relacionadas ao mundo dos esportes, como a probabilidade de um time se tornar campeão ou de um recorde ser quebrado.

Cada capítulo oferecerá uma análise aprofundada do tema específico, proporcionando aos leitores uma compreensão mais clara e abrangente da probabilidade e suas aplicações em diversos contextos da vida cotidiana.

No desfecho, a dissertação visa não apenas corrigir equívocos comuns relacionados à probabilidade, mas também fornecer aos leitores ferramentas para uma compreensão mais profunda dos desafios enfrentados ao lidar com eventos aleatórios. Ao destacar a interação entre os aspectos matemáticos e cognitivos, busca-se promover uma abordagem mais holística e esclarecedora sobre a natureza da incerteza em nossas vidas e como melhor navegar por ela.

1. PROBABILIDADE E ESTATÍSTICA: UMA VISÃO ATRAVÉS DO VIÉS DA PSICOLOGIA

Muitos estudos têm sido conduzidos sobre a validade e a consistência dos julgamentos probabilísticos, contudo, há pouca investigação sobre a influência dos mecanismos psicológicos através dos as pessoas avaliam a probabilidade. Ante isso, o presente capítulo se baseia em três obras: “Rápido e Devagar: Duas Formas de Pensar” de Kahneman (2012), “*Availability: A Heuristic for Judging Frequency and Probability*” de Kahneman e Tversky (1973), e o livro “*Judgment under uncertainty Heuristics and Biases*” de Kahneman, Tversky, (1982), nas quais os autores apresentam problemas de natureza probabilística através de uma lente psicológica. A partir desses estudos, será analisado como o cérebro humano desenvolve mecanismos capazes de distorcer a avaliação da probabilidade de ocorrência de um evento específico.

Segundo Kahneman (2012), o cérebro humano opera através de dois tipos de pensamento, denominados Sistema 1 e Sistema 2. As perspectivas de Kahneman (2012) e de Kahneman e Tversky (1973) sobre os mecanismos de julgamento serão apresentadas, com foco particular nos julgamentos de natureza probabilística.

As operações realizadas pelo Sistema 1 são automáticas e rápidas, ocorrendo com pouco ou nenhum esforço consciente; elas geram associações e são frequentemente carregadas de emoções, resultando em um tipo de raciocínio mais intuitivo. Este sistema é responsável por atividades que não exigem grande esforço cognitivo, como dirigir em uma rua vazia, realizar cálculos matemáticos simples, ou perceber que um objeto está mais distante do que outro. Apesar de o Sistema 2 tomar decisões de forma mais lenta que o Sistema 1, suas operações têm a capacidade de resolver tarefas que envolvam um nível de raciocínio mais elaborado, incluindo cálculos complexos. As operações do Sistema 2 são muitas vezes associadas com a experiência subjetiva de atividade, escolha e concentração. É possível citar, contar as ocorrências da letra “a” em uma página de texto, estacionar em uma vaga apertada, verificar a validade de um argumento lógico complexo, como algumas tarefas realizadas pelo Sistema 2. Quando confrontados com problemas de difícil solução, se o cérebro humano não consegue encontrar uma resposta satisfatória rapidamente, o Sistema 1 tende a buscar uma pergunta mais fácil e responder a ela. Esse tipo de raciocínio cognitivo é conhecido como heurística. A heurística é definida como um procedimento simples que auxilia na busca por respostas adequadas, embora geralmente imperfeitas, para questões complexas (KAHNEMAN, 2012).

A Tabela 1 ilustra situações cotidianas, contrastando a decisão ótima (aquela que deveria ser alcançada pelo Sistema 2) e a heurística que o cérebro humano desenvolve para tentar facilitar a resolução do problema:

Tabela 1 – Pergunta Alvo e a sua heurística correspondente.

Pergunta Alvo	Pergunta Heurística
Até que ponto você contribuiria para salvar espécies em risco de extinção?	Até que ponto me emociono quando penso em golfinhos morrendo?
O quanto você está feliz com a sua vida atualmente?	Qual é meu humor neste exato momento?
Qual será a popularidade do presidente daqui a seis meses?	Qual é a popularidade do presidente neste momento?
Como devem ser punidos consultores financeiros que se aproveitam dos aposentados?	Quanta raiva eu sinto quando penso em predadores financeiros?
Esta mulher está concorrendo para a primária. Até onde ela chegará na política?	Esta mulher parece uma vitoriosa na política?

Fonte: KAHNEMAN, 2012.

A Tabela 1 apresenta duas colunas: na primeira, está a pergunta alvo, que é uma avaliação da situação que se procura produzir. Na segunda coluna, está a pergunta heurística, que é a questão que o cérebro humano, através do Sistema 1, utiliza como substituta da pergunta original, quando não consegue resolvê-la. Por exemplo, na terceira linha da tabela, a pergunta alvo “Qual será a popularidade do presidente daqui a seis meses?” é substituída pela pergunta heurística “Qual é a popularidade dele neste momento?”

Ao responder à pergunta usando a heurística, cria-se a ilusão de que se respondeu à pergunta original, quando, na verdade, a popularidade atual do presidente não necessariamente reflete sua popularidade daqui a seis meses, pois esta pode ser influenciada por fatores como crises políticas, econômicas ou escândalos pessoais envolvendo o presidente. Assim, os mecanismos de heurísticas, embora simplifiquem questões complexas, podem levar a erros de julgamento sobre um determinado evento (KAHNEMAN, 2012).

As heurísticas podem ser divididas em quatro tipos principais:

- Heurística de Disponibilidade;
- Heurística de Representatividade;
- Heurística de Ancoragem;
- Heurística de Simulação.

Neste capítulo, será dada ênfase nas heurísticas de disponibilidade e representatividade.

1.1 Heurística da Disponibilidade

A heurística de disponibilidade é um processo mental pelo qual o cérebro humano estabelece a probabilidade de um evento acontecer com base na facilidade com que se lembra de um evento similar do passado. Esse mecanismo foi altamente útil para a sobrevivência humana em tempos ancestrais, quando os seres humanos enfrentavam situações de perigo na selva e precisavam de respostas rápidas, mesmo que por vezes imprecisas. Por exemplo, ao se deparar com um animal perigoso, o acesso rápido a eventos passados similares ajudava a discernir se deveriam ou não correr. Nesses casos, a rapidez na tomada de decisão era essencial para a sobrevivência da espécie (KAHNEMAN & TVERSKY, 1973).

Com isso, podemos observar que a heurística de disponibilidade não é algo necessariamente prejudicial, pois ela pode levar o cérebro humano a acessar situações adversas do passado, permitindo que utilizem essas informações para tomar decisões mais acertadas no presente. Por exemplo, ao lembrar-se de ter chegado atrasado a um compromisso devido ao trânsito, uma pessoa pode decidir sair mais cedo de casa na próxima vez, antecipando mesmo que não haja congestionamento, o que lhe permite chegar com antecedência.

Entretanto, o uso excessivo dessa heurística para lidar com eventos aleatórios pode resultar em problemas, uma vez que as pessoas tendem a confiar incondicionalmente na primeira lembrança que lhes ocorre, na associação mais instantânea que seus cérebros são capazes de produzir, sem considerar outros fatores relevantes para a avaliação probabilística da situação em questão.

1.2 Heurística de Representatividade

Outro tipo de heurística que pode acarretar erros de julgamento sobre a chance de um determinado evento ocorrer é a heurística de representatividade. Esse tipo de heurística é acessado pelo cérebro humano quando confrontado com uma decisão acerca da probabilidade de um evento acontecer, e é dada uma descrição do objeto sobre o qual queremos inferir a probabilidade. Ao avaliar as probabilidades de ocorrência de um evento associado ao objeto, as pessoas julgam o evento como mais ou menos provável, de acordo com sua semelhança com a descrição fornecida. Em outras palavras, o cérebro humano usa estereótipos para determinar a chance de ocorrência de um evento. Esse mecanismo cognitivo é tão forte que,

mesmo quando há informações a priori sobre a probabilidade de ocorrência ou não de um evento, há tendência de ignorar essas informações e se ater apenas ao quão representativo é o evento frente à descrição dada. A seguir, algumas características dessa heurística serão abordadas (KHANEMAN E TVERSKY,1973)

1.2.1 Insensibilidade a dados anteriores sobre a probabilidade de ocorrência de um determinado evento

Um dos fatores que não afetam a representatividade, mas têm influência direta sobre a probabilidade, é a taxa básica de frequência, ou seja, as probabilidades definidas a priori. Para exemplificarmos esta influência, observar-se-á a descrição de Steve, retirada do livro “*Judgment under uncertainty Heuristics and biases*”:

Steve é muito tímido e retraído, invariavelmente prestativo, mas com pouco interesse nas pessoas ou no mundo real. Uma alma mansa e limpa, ele tem necessidade de ordem e estrutura, e uma paixão por detalhes (KHANMEMAN e TVERSKY, 1982, pág.4).

Essa descrição, embora não contenha informações concretas sobre a profissão de Steve, reflete estereótipos que tendem a ser associados mais prontamente a um bibliotecário do que a um fazendeiro. Mesmo que, estatisticamente, haja mais fazendeiros do que bibliotecários em determinada população, o viés da heurística da representatividade leva as pessoas a fazerem inferências com base na representatividade estereotipada da profissão mencionada.

Mesmo quando é fornecida uma proporção exata dos dois grupos (por exemplo, explicitando que em um grupo de 100 pessoas, há 60 fazendeiros), a heurística da representatividade tende a prevalecer sobre essa informação. Assim, as pessoas, mesmo cientes da proporção real de fazendeiros e bibliotecários no grupo, muitas vezes julgam mais provável que Steve seja um bibliotecário do que um fazendeiro, simplesmente porque a descrição se encaixa mais no estereótipo associado à profissão de bibliotecário. Tal fenômeno evidencia como os estereótipos podem influenciar nossas percepções e julgamentos, muitas vezes desconsiderando informações estatísticas ou probabilísticas (KAHNEMAN & TVERSKY, 1982).

1.2.2 Insensibilidade ao tamanho da amostra

Outro fator que não influencia na representatividade é o tamanho da amostra. Consequentemente, ao analisarmos a probabilidade de um conjunto de dados utilizando a heurística da representatividade, o julgamento da probabilidade não irá depender de seu

tamanho. Para ilustrar esta situação, pretende-se analisar uma situação descrita no livro “*Judgment Under Uncertainty: Heuristic and Biases*”:

Uma determinada cidade é servida por dois hospitais. No hospital maior, cerca de 45 bebês nascem a cada dia e, no hospital menor, cerca de 15 bebês nascem a cada dia. Como você sabe, cerca de 50% de todos os bebês são meninos. No entanto, a porcentagem exata varia de dia para dia. Às vezes pode ser maior que 50%, às vezes menor. Por um período de 1 ano, cada hospital registrou os dias em que mais de 60% dos bebês nascidos eram meninos e os dias em que nasceram menos que 60%. (KHANMEMAN e TVERSKY, 1982, p. 6).

A tabela a seguir mostra como estudantes de graduação avaliam em qual dos hospitais ocorrem mais dias com 60% de meninos ou dias com menos que 60% de meninos.

Tabela 2 – Proporção de nascimentos nos 2 hospitais

	Mais que 60%	Menos que 60%
Hospital Maior	12	9
Hospital Menor	10	11
Chance Igual	28	25

Fonte: KHANMEMAN e TVERSKY, 1982.

A partir deste experimento, foi possível observar que a maioria das pessoas tende a acreditar que, independentemente do tamanho do hospital, a chance de o número de meninos ser maior que 60% é a mesma. Este pensamento, embora contrarie a teoria da probabilidade e estatística, que indica que, conforme se aumenta o tamanho da amostra, a chance dela se desviar do valor médio diminui, é justificado pela heurística da representatividade.

Para o cérebro humano, ambos os hospitais são igualmente representativos, ou seja, não há nenhum indício de que um deles teria uma chance maior de se desviar da probabilidade de 50%. Assim, nossa mente julga que a ocorrência de dias "anormais" é a mesma para os dois hospitais. Da mesma forma, a maioria dos estudantes deu a mesma resposta quando confrontados com a pergunta oposta: a chance de o número de meninos ser menor que 60%. No entanto, esta pergunta requer uma análise mais sutil.

Quando se trata do número de meninos menor que 60%, mas próximo do número real de 50%, a chance de ocorrência será maior no hospital em que nascem mais bebês. Isso ocorre porque, com um maior espaço amostral, a probabilidade de o número de bebês se aproximar dos 50% também é maior.

Por outro lado, quando o número de meninos é menor que 60% e está afastado do número real de 50%, a chance de ocorrência será maior no hospital em que nascem menos bebês. Neste caso, com uma quantidade menor de bebês nascendo diariamente, o hospital está mais suscetível a variações amostrais, o que pode resultar em desvios mais significativos do valor médio esperado.

1.2.3 Equívocos acerca da probabilidade

As pessoas possuem a tendência de acreditar que eventos gerados por processos aleatórios representam as características essenciais do processo, mesmo quando a sequência é pequena. Ao observar o lançamento de uma moeda, por exemplo, a sequência *Ca-Co-Ca-Co-Co-Ca* é considerada mais provável de ocorrer do que a sequência *Ca-Ca-Ca-Co-Co-Co*. Esse viés ocorre devido ao processo de heurística de representatividade. Na primeira sequência, a alternância entre caras e coroas representa melhor a aleatoriedade esperada de um lançamento de moeda, enquanto na segunda sequência, o fato de termos três vezes seguidas caras e três vezes seguidas coroas induz as pessoas a pensarem que não se trata de um processo aleatório, pois não representa a “imparcialidade” da moeda.

Uma consequência direta desse pensamento é a chamada falácia do apostador. Nessa falácia, após observar uma longa sequência de vermelhos na roleta, o apostador acredita que existe uma probabilidade maior de o próximo resultado ser preto. Isso ocorre porque a próxima ocorrência da roleta ser preta acarreta numa sequência mais representativa. As pessoas têm a tendência de acreditar na probabilidade como um processo autocorretivo, de modo que qualquer sequência que saia fora dos padrões considerados aleatórios é automaticamente percebida como sendo corrigida para restaurar o equilíbrio.

Após esse breve comentário sobre as heurísticas, serão abordados alguns problemas do livro “O Andar do Bêbado”, que fazem referência a elas, mostrando como cada uma influencia a forma como se lidam com problemas de incerteza.

1.3 Por que sempre nos achamos tão azarados?

Por que as pessoas frequentemente se consideram tão azaradas? Para iniciar esse estudo, pretende-se abordar um problema simples que se relaciona com o cotidiano: as filas de supermercado, sendo notório que estar em uma fila é uma situação estressante para muitos. O desejo comum das pessoas é que o tempo passe rapidamente para que possam dedicá-lo a atividades mais proveitosas. Nesse sentido, “Qual é a probabilidade de escolhermos, entre as

cinco filas disponíveis no mercado, aquela que se revela a mais lenta? (MLODINOW, 2009, p. 25)”.

A resposta para essa pergunta é bem simples e intuitiva: Ao escolher uma das cinco filas disponíveis, existe a possibilidade de selecionar aquela que se revela a mais lenta é de $\frac{1}{5}$. Logo a chance de isso não ocorrer é de $\frac{4}{5}$. Mas por que há a impressão de que sempre se pega a fila mais lenta?

A mente humana constantemente distorce probabilidades. Mecanismos são utilizados pelo cérebro para tentar determinar a probabilidade “verdadeira”, mas muitas vezes simplifica demasiadamente o pensamento, resultando em uma probabilidade subjetiva que se distancia da verdadeira natureza do fenômeno. Conforme Kahneman e Twersky (1973) explicam, no exemplo mencionado, o cérebro busca atalhos para resolver o problema, utilizando a heurística da disponibilidade. Neste tipo de heurística, as pessoas tendem a subestimar ou superestimar probabilidades com base em suas experiências pessoais. Os autores definem a heurística da disponibilidade como a propensão de relacionar um evento específico a fatos que vêm rapidamente à mente.

Para ilustrar este raciocínio, eles analisaram o seguinte problema: “Um clínico escuta de seu paciente que ele está cansado de viver. Como será que ele irá estimar a chance desse paciente cometer suicídio? (KHANMEMAN e TVERSKY, 1973, p. 228)”. Para quantificar esta probabilidade, o médico irá tentar se lembrar de pacientes semelhantes ao que ele conheceu. Às vezes, apenas o caso de um paciente virá a sua mente, possivelmente aquele em que a história for a mais memorável. Neste caso, se os dois casos tiverem muitas similaridades, então o clínico tenderá a achar que aquela situação do passado ocorrerá novamente.

Vale ressaltar que se ocorrer a lembrança de dois ou mais casos relevantes, o cérebro humano os organizará em ordem de similaridade com o caso em estudo. Essas memórias serão categorizadas em três tipos: pacientes com sintomas semelhantes, pacientes que tentaram suicídio e pacientes que tentaram suicídio e tinham sintomas semelhantes.

Para estimar a probabilidade de um paciente cometer suicídio, a classe relevante seria a de pacientes com sintomas semelhantes, e a estatística relevante seriam os casos de tentativas de suicídio dentro dessa classe. Entretanto, a memória tende a superestimar fatos extremos, como tentativas de suicídio. Como resultado, os clínicos tendem a lembrar apenas dos casos de pacientes que tentaram suicídio, buscando estabelecer uma conexão entre eles e analisando se esses pacientes compartilhavam os mesmos sintomas descritos pelo paciente

atual. Isso muitas vezes leva à ignorância dos casos igualmente relevantes de pacientes com sintomas similares, mas que não tentaram suicídio.

Da mesma forma, no caso da fila de mercado, ocorre um erro de julgamento, pois o que vem à mente são os casos em que a fila demora. Quando tudo ocorre conforme o planejado, não é dada a devida importância a essa situação. Os momentos em que a fila do mercado é rápida, quando não há problemas no trabalho, quando o caminho de volta para casa está livre de trânsito, não são gravados no nosso subconsciente. Porém, aquele dia em que tudo deu errado nunca será esquecido: conflito com o chefe, trânsito intenso na volta para casa e ainda uma discussão com a esposa, tudo no mesmo dia! Isso pode levar a pessoa a se sentir a mais azarada do mundo, achando que tudo na sua vida dá errado (a famosa Lei de Murphy). Da mesma forma, o cérebro humano tende a se lembrar sempre das vezes em que a fila demorou muito, ignorando completamente as vezes em que tudo ocorreu de forma rápida.

1.4 O problema das palavras

Neste problema, há uma curiosidade relacionada ao idioma inglês e à tendência das pessoas de cometerem equívocos em problemas envolvendo contagem. Abaixo, pretende-se examinar o problema proposto por Mlodinow (2009):

O que é maior, o número de palavras de seis letras na língua inglesa que têm *n* como sua quinta letra *n*, ou o número de palavras de seis letras na língua inglesa que terminam em *ing*? A maior parte escolhe o grupo terminado em *ing*. Por quê? (MLODINOW, 2009, p. 24).

Primeiramente, o problema será representado através da notação de conjuntos:

$A = \{\text{conjunto das palavras de seis letras em que a quinta letra é o } n\}$

$B = \{\text{conjunto das palavras de seis letras terminadas em } ing\}.$

Através dessa representação, percebe-se que todos os elementos do conjunto *B*, também são elementos do conjunto *A*, pois toda palavra de seis letras que termina em *ing* também é uma palavra de seis letras na qual a quinta letra é a letra *n*. Usando a notação de conjuntos, se poderia afirmar que $B \subset A$.

1.4.1 Mas por que as pessoas acham que o conjunto *B* tem mais elementos?

Isso ocorre porque as pessoas lembram de mais palavras que terminam em "ing" do que palavras genéricas de seis letras em que a quinta letra é "n". Essa tendência é atribuída ao fato de que a terminação "ing" é amplamente utilizada na língua inglesa para formar conjugações verbais no gerúndio, o que cria a ilusão de que existem mais palavras de seis

letras terminando em "ing" do que palavras de seis letras com a quinta letra sendo "n". Isso leva a um erro na avaliação da quantidade de elementos em cada um dos conjuntos apresentados.

Assim como no exemplo anterior, essa avaliação errônea das probabilidades é resultado do uso da heurística de disponibilidade. Apesar de o conjunto de palavras com seis letras terminadas em "ing" ser um subconjunto do conjunto de palavras com seis letras em que a quinta letra é "n", o cérebro humano tem uma capacidade maior de se lembrar das palavras terminadas em "ing" do que das que não terminam dessa forma. Essas palavras estão mais prontamente disponíveis em nossa mente, levando a esse tipo de erro de julgamento.

Outro experimento similar a este, feito por Kahneman e Twersky (1973), será descrito a seguir:

(...) Considere uma palavra da língua inglesa que contenha a letra k.

É mais provável que o k apareça na:

() Primeira Letra

() Terceira Letra (KHANMEMAN e TVERSKY, 1973, p. 211).

A menos que sejam especialistas em língua inglesa, as pessoas não têm como determinar efetivamente qual das duas classes de palavras contém mais elementos. Como isso, é razoável supor que um resultado provável para este experimento seria que aproximadamente metade das pessoas escolhessem a primeira opção e aproximadamente metade escolhessem a segunda. No entanto, o resultado observado foi diferente: a primeira opção foi escolhida por cerca de 66% das pessoas. Isso ocorre porque é mais fácil para o cérebro humano se lembrar de palavras da língua inglesa em que "k" é a primeira letra do que palavras em que "k" é a terceira letra.

Nestes tipos de problemas, não é possível construir todas as sequências possíveis em nosso cérebro; então, o que se faz é um julgamento baseado nas palavras às quais o cérebro consegue acessar, resultando em uma distorção das frequências reais, devido à tendência de lembrar-se de mais palavras que começam com a letra "k" do que aquelas em que "k" é a terceira letra. A este tipo de heurística de disponibilidade, dá-se o nome de disponibilidade por construção.

Embora esses dois problemas sejam bastante análogos, existe uma diferença primordial entre eles. Enquanto no primeiro caso, uma análise mais minuciosa do problema

indica claramente qual dos conjuntos é maior, neste segundo caso, sem um conhecimento prévio, não é possível determinar qual dos conjuntos é maior.

1.5 Será Linda feminista?

Em um outro exemplo utilizado por Mlodinow (2009), ele aborda esta mesma temática. Neste exemplo, ele transcreve um experimento realizado por Kahneman e Tversky (1973), no qual eles descrevem uma mulher chamada Linda e depois apresentam uma tabela em que as pessoas atribuem valores de 1 a 8, de acordo com como elas acham aquele evento provável (sendo 1 o mais provável e 8 o mais improvável).

Imagine uma mulher chamada Linda, de 31 anos de idade, solteira, sincera e muito inteligente. Coursou filosofia na universidade. Quando estudante, preocupava-se profundamente com discriminação e justiça social e participou de projetos contra armas nucleares. Tversky e Kahneman apresentaram esta descrição a um grupo de 88 pessoas e lhes pediram que classificassem as seguintes afirmações numa escala de 1 a 8, de acordo com a sua probabilidade, de modo que 1 representasse a mais provável e 8 a mais improvável (MLODINOW, 2009, p. 20).

Neste experimento, os 88 participantes foram solicitados a fazer uma correspondência entre uma escala variando de 1 a 8 e a probabilidade de cada evento que sabemos variar de 0 a 1. É importante observar que quanto maior o valor escolhido na escala, menos provável se considera o evento analisado.

Tabela 3 – Relação entre as afirmações sobre Linda e sua chance de ocorrência.

Afirmação	Classificação de acordo com a chance de ocorrência
Linda participa do movimento feminista	2,1
Linda é assistente social psiquiátrica	3,1
Linda trabalha em uma livraria e faz aulas de ioga	3,3
Linda é bancária e participa do movimento feminista	4,1
Linda participa da liga pelo voto feminino	5,4
Linda é bancária	6,2
Linda é corretora de seguros	6,4

Fonte: MLODINOW, 2009.

Ao observar a tabela, é possível perceber um erro de avaliação por parte das pessoas. A chance de ocorrência associada a Linda ser bancária é de 6,2, enquanto a dela ser feminista e bancária é de 4,1. Isso significa que, para grande parte das pessoas entrevistadas (aproximadamente 85%, segundo Mlodinow, 2009), Linda é considerada mais provável de ser feminista e bancária do que ser bancária.

Esse equívoco ocorre novamente devido à heurística de representatividade. Ao enfrentar problemas envolvendo o acaso, o cérebro humano tende a analisar a quão representativa é uma situação para todos, ou seja, o quanto o objeto *B* se assemelha à classe *A*. Quanto maior for essa relação de representatividade ou semelhança, mais provável nosso cérebro irá julgar aquele evento.

Apesar de não fazer sentido acreditar que é mais provável que Linda seja feminista e bancária do que seja bancária, a descrição apresentada por Kahneman e Twersky (1982) revela traços estereotipados de uma feminista, tais como lutar contra a discriminação, por justiça social e contra armas nucleares. O cérebro humano tende a assimilar essas características de Linda associadas ao feminismo, levando muitas pessoas a julgarem que há mais chances de ocorrerem ambos os eventos do que um deles, mesmo contrariando regras básicas de probabilidade.

Em outro experimento realizado por Kahneman e Twersky (1973), foi apresentada a um grupo de pessoas uma breve descrição de vários indivíduos, selecionados aleatoriamente de um grupo composto por 100 trabalhadores, alguns dos quais eram engenheiros e outros advogados.

Essas pessoas foram então questionadas sobre a probabilidade de cada indivíduo ser engenheiro ou advogado. Para responder a essa pergunta, elas foram divididas em dois grupos: em um deles, foi informado que havia 30 advogados e 70 engenheiros, enquanto no outro grupo foi dito que havia 30 engenheiros e 70 advogados. Portanto, as chances de um indivíduo ser engenheiro eram significativamente maiores nas condições apresentadas ao primeiro grupo do que nas apresentadas ao segundo.

Todavia, o experimento não refletiu isso. Os julgamentos sobre a probabilidade de um determinado indivíduo ser engenheiro ou advogado foram os mesmos nos dois casos. Isso ocorre porque a avaliação da probabilidade de um indivíduo em particular ser advogado ou engenheiro é acessada pelo cérebro humano através da heurística da representatividade, ou seja, a escolha foi baseada nos estereótipos presentes em cada descrição associada ao perfil de um engenheiro e ao perfil de um advogado, ignorando as probabilidades informadas previamente.

Já quando não são fornecidas as características referentes aos 100 trabalhadores, as pessoas tendem a utilizar a informação acerca da quantidade de engenheiros e advogados, para estimar a probabilidade de um indivíduo ser engenheiro ou advogado, ou seja, para o primeiro grupo, cada indivíduo tinha 30% de chance de ser advogado e para o segundo grupo, cada indivíduo tinha 70% de chance de ser advogado.

Sempre que uma descrição é feita, as probabilidades já conhecidas são ignoradas, mesmo quando a descrição é neutra, ou seja, não apresenta indícios sobre se a pessoa é engenheiro ou advogado. Considere a seguinte descrição: “Dick é um homem de 30 anos, casado sem filhos. Ele possui uma grande motivação e grandes habilidades e promete ser bem-sucedido em sua área de atuação. Ele é bem-querido pelos seus colegas (MLODINOW, 2009, p. 20)”.

Percebe-se que essa descrição não contém elementos indicativos de que Dick é engenheiro ou advogado. Portanto, seria razoável esperar que a probabilidade de ele ser engenheiro fosse igual à proporção de engenheiros no grupo (0,3 no primeiro grupo e 0,7 no segundo). No entanto, o cérebro humano, ao interpretar a descrição, não encontra nenhuma evidência sobre a profissão de Dick. Assim, ignorando as probabilidades anteriores, ele associa que a chance de ele ser engenheiro é de 0,5, ou seja, ele possui a mesma probabilidade de ser advogado ou engenheiro.

Após esta breve análise sobre como fatores psicológicos influenciam nossa percepção, o próximo capítulo introduzirá o conceito de espaço amostral, que é de suma importância para a resolução de diversos problemas relacionados à probabilidade.

2. PROBABILIDADE E O ESPAÇO AMOSTRAL

Um conceito fundamental no estudo de probabilidade é o de espaço amostral. Apesar de ser uma definição relativamente simples, muitos dos equívocos em relação a julgamentos probabilísticos se dão pela sua má interpretação. Muitas vezes, são cometidos erros na contagem de seus elementos, enquanto em outras vezes espaços amostrais não-equiprováveis são considerados como equiprováveis.

Neste capítulo, será feita inicialmente uma abordagem histórica de como o conceito de espaço amostral foi se desenvolvendo ao longo do tempo, correlacionando as estruturas descritas por Von Mises com a definição moderna de espaço amostral. Após este compilado histórico, serão analisados alguns problemas do livro “O Andar do Bêbado”, tentando sempre, além de resolvê-los, explicitar o que leva as pessoas a interpretarem de maneira errônea.

2.1 Espaço Amostral: Uma breve abordagem histórica

Richard Von Mises (1883-1943), um matemático e engenheiro mecânico austríaco, juntamente com R. A. Fisher (1890-1962), são considerados os pais da probabilidade estatística ou empírica. Mises e Fisher calcularam a probabilidade de forma experimental, repetindo diversas vezes um experimento para inferir a sua chance de ocorrência.

Dessa forma, para Von Mises (1928), a probabilidade de um determinado evento ocorrer é considerada como o limite da frequência relativa, considerando uma repetição infinita de experimentos sob as mesmas condições. Considerando “ s ” como o número de vezes em que o evento A ocorre e “ n ” como o número de vezes que o experimento foi observado, Mises define a probabilidade de ocorrência de um determinado evento A , denominada $P(A)$, como sendo:

$$P(A) = \lim_{n \rightarrow \infty} f_n(A), \text{ sendo } f_n(A) = \frac{s}{n}.$$

Ao analisar essa definição de probabilidade proposta por Von Mises, é importante observar que esse limite não está relacionado ao conceito usual de limite de uma sequência. Em análise, pode-se dizer que $P(A) = \lim_{n \rightarrow \infty} f_n$, significa dizer que dado $\varepsilon > 0$, existe n_0 tal que, para todo $n > n_0$, temos que $|f_n(A) - P(A)| < \varepsilon$. Em outras palavras, esta definição nos indicaria que conforme aumentássemos o número de lançamentos, a frequência relativa do evento A teria que estar mais próximo da sua probabilidade, o que nem sempre acontece.

A melhor interpretação que pode ser dada para o limite proposto por Von Mises é a seguinte:

$$P(A) = \lim_{n \rightarrow \infty} f_n, \text{ equivale a dizer que:}$$

Para todo $\varepsilon > 0$ temos que:

$$P(|P(A) - f_n(A)| > \varepsilon) \rightarrow 0, \text{ quando } n \rightarrow +\infty$$

Esta definição, ao contrário da anterior, permite que $|f_n(A) - P(A)|$ exceda qualquer número dado, por maior que seja, desde que seja < 1 (por conta da própria definição de probabilidade que será analisada mais a frente). só que ela nos garante que isto ocorrerá com probabilidades cada vez mais próximas de zero.

Para melhor caracterizar o estudo desses experimentos, Von Mises introduziu alguns conceitos importantes:

- O primeiro deles é o conceito de **Coletivo**. Esse termo denota uma sequência de eventos uniformes ou processos que se diferenciam através de atributos observáveis, como cores, números, entre outros. Por exemplo, todas as rolagens de um dado durante um jogo formam um coletivo, no qual o atributo de cada evento simples é o número que foi rolado no dado;
- Outra definição é a de **Espaço dos Atributos**. O termo atributo é utilizado para designar a característica de cada um dos elementos do evento aleatório. Com base nisso, Von Mises definiu o Espaço dos Atributos como sendo o conjunto de todos os possíveis atributos de um evento simples. Por exemplo, no caso do lançamento de um dado, temos que o Espaço dos Atributos é dado por $\{1,2,3,4,5,6\}$.

A partir disso, Von Mises definiu o espaço dos atributos para um experimento composto. Para isso, deve-se, primeiramente, definir o que seria um experimento composto. Um experimento é dito composto se consiste na realização de n experimentos sucessivos, ou seja, n ensaios independentes do experimento básico. Portanto, o espaço dos atributos desse tipo de experimento é formado por todas as combinações possíveis dos n espaços dos atributos dos experimentos básicos. Sendo assim, se Ω_0 é o espaço dos atributos do experimento básico, então o espaço amostral para o experimento composto pode ser definido da seguinte forma:

$$\Omega_0 = \{ (w_1, \dots, w_i : w_i \in \Omega_0, i = 1, \dots, n) \}, \text{ ou seja, } \Omega_n \text{ é o espaço produto } \Omega_0 \times \Omega_0 \times \dots \times \Omega_0 = \Omega_0^n$$

Em outras palavras, para expressar o espaço dos atributos de um experimento composto, basta fazermos o produto cartesiano do espaço dos atributos do experimento básico, por ele mesmo, n vezes. Para ilustrar este conceito, imagine o lançamento de uma moeda. O evento básico tem como espaço dos atributos $\Omega_0 = \{\text{cara}, \text{coroa}\}$. Se quisermos saber o espaço dos atributos do evento composto “lançamento sucessivo de uma moeda 3 vezes”, basta fazermos o produto cartesiano, ficando com o espaço amostral $\Omega_n = \{(\text{cara}, \text{cara}, \text{cara}), (\text{cara}, \text{cara}, \text{coroa}), (\text{cara}, \text{coroa}, \text{cara}), (\text{cara}, \text{coroa}, \text{coroa}), (\text{coroa}, \text{coroa}, \text{coroa}), (\text{coroa}, \text{coroa}, \text{cara}), (\text{coroa}, \text{cara}, \text{coroa}), (\text{coroa}, \text{cara}, \text{cara})\}$.

Esta definição utilizada por Von Mises (1928) de espaço dos atributos é bem similar à que se utiliza atualmente para espaço amostral. Já a definição clássica de probabilidade foi proposta inicialmente por Bernoulli e Laplace. Nesta definição, é considerado conhecido o conjunto de todos os n resultados possíveis e equiprováveis de um espaço amostral. Se m

deles forem favoráveis a um evento A , então a probabilidade de A acontecer, denotada por $P(A)$ é $P(A) = \frac{m}{n}$. Ao conjunto de todos os resultados possíveis damos o nome de espaço amostral.

Ao analisar essa definição, encontram-se alguns problemas. O primeiro deles é o fato de a própria definição conter o termo “equiprovável” (ter a mesma probabilidade), o que acarreta um problema de recursividade, onde é necessário o entendimento completo da definição de probabilidade para compreendê-la completamente.

Outro problema é que essa definição não abrange todos os problemas relacionados à probabilidade. Nesse sentido, Feller (1968) lembra que, historicamente, o objetivo original da teoria era descrever um aspecto bastante específico de questões associadas a jogos de azar, e o maior esforço se dirigia ao cálculo efetivo de certas probabilidades. Em diversos problemas, os n resultados possíveis não são equiprováveis. Para solucionar problemas desse tipo, utiliza-se o método proposto por ele, onde o evento do qual se quer calcular a probabilidade é dividido em diversos eventos simples, cada um deles sendo equiprovável. Assim, a chance de ocorrência daquele evento é dada pelo número de eventos simples associados a ele dividido pelo número de elementos do espaço amostral.

Na obra clássica de Feller (1968), “eventos” são definidos como os resultados pensáveis. São chamados de eventos indivisíveis aqueles que não podem ser decompostos em outros eventos. O autor nomeia esses “eventos simples” de pontos amostrais. Essa ideia de ponto amostral é um conceito axiomático e serve como base para as demais definições. Por definição, cada um dos resultados indecomponíveis é representado por um único ponto amostral.

Por exemplo, ao lançar uma moeda, há dois pontos amostrais: cara e coroa. Já o evento de tirar 2 caras e 1 coroa no lançamento de 3 moedas pode ser decomposto em 3 pontos amostrais: (cara, cara, coroa), (coroa, cara, cara) e (cara, coroa, cara).

Ao conjunto formado por todos os pontos amostrais dá-se o nome de espaço amostral. Os possíveis resultados contidos neste espaço amostral são denominados eventos. Por exemplo, ao lançar um dado, alguns exemplos de possíveis eventos são “a face ser um número par” ou “a face ser um número maior do que 4”.

Para ilustrar o conceito de espaço amostral, Feller (1968) propõe o seguinte problema: tem-se que distribuir 3 bolas em 3 caixas. Qual seria o espaço amostral deste experimento?

Figura 1 – Espaço Amostral do experimento 3 bolas em 3 caixas

1. { abc - - }	10. { a bc - }	19. { - a bc }
2. { - abc - }	11. { b a c - }	20. { - b a c }
3. { - - abc }	12. { c ab - }	21. { - - c ab }
4. { ab c - }	13. { a - bc }	22. { a b c }
5. { a c b - }	14. { b - a c }	23. { a c b }
6. { bc a - }	15. { c - ab }	24. { b a c }
7. { ab - c }	16. { - ab c }	25. { b c a }
8. { a c - b }	17. { - a c b }	26. { c a b }
9. { bc - a }	18. { - bc a }	27. { c b a }

Fonte: FELLER, 1968.

Cada uma das bolas é representada por uma letra a , b , c . Como exemplo, a arrumação de número 1 representa que as três bolas estão na primeira caixa, enquanto a arrumação de número 4 representa que as bolas a e b estão na primeira caixa e a bola c na segunda. O espaço amostral deste experimento é composto por todas as 27 organizações possíveis.

A partir desse espaço amostral, é possível definir eventos contidos nele. Por exemplo, o evento A descrito por “alguma caixa está ocupada por duas ou mais bolas” é composto pelos pontos amostrais 1-21, enquanto o evento B “a primeira caixa não está vazia” é composto pelos pontos amostrais 1, 4-15, 22-27.

A partir desses eventos, é possível definir outros eventos relacionados a eles. Por exemplo, o evento C , “ A e B ocorrem simultaneamente”, é composto pelos pontos amostrais 1, 4-15, que são os pontos amostrais que pertencem tanto ao evento A quanto ao evento B . Neste exemplo, como cada um dos pontos amostrais têm a mesma chance de ocorrer, a probabilidade de cada um deles é dada por $\frac{1}{27}$, pois o espaço amostral neste exemplo é composto por 27 elementos.

2.2 Algumas definições importantes de probabilidade

Antes de iniciar a discussão dos problemas deste capítulo, serão definidos alguns conceitos importantes acerca de probabilidade. Essas definições foram baseadas nos conceitos apresentados no livro “*An Introduction to Probability Theory*” (FELLER 1968, págs. 15 a 23).

2.2.1 Relações entre dois eventos

Considere um espaço amostral Ω fixado. Os eventos contidos no espaço amostral serão representados por letras maiúsculas. Denotar-se-á da forma $x \in A$, se x é um ponto

amostral contido no evento A . A partir disso, seguem as seguintes definições (FELLER, 1968):

Definição 1 – Será utilizada a notação $A = \emptyset$ para representar o evento contido no espaço amostral Ω , que não contém nenhum ponto amostral (evento impossível).

Definição 2- O evento que consiste em todos os pontos amostrais que não estão contidos no evento A será denominado complementar de A e representado por A' .

Definição 3- Dados dois eventos A e B , com A e B contidos no espaço amostral Ω , definiremos 2 novos eventos: “ A e B ocorrem simultaneamente” e “ A ou B ocorrem ou ambos ocorrem”. Estes dois eventos serão representados por $A \cap B$ e $A \cup B$, respectivamente. O primeiro deles é composto por todos os pontos amostrais que são comuns a A e B . Se não houver pontos em comum entre A e B , teremos que o evento $A \cap B$ é impossível e, portanto, diz-se que os eventos A e B são mutuamente exclusivos. Já o evento $A \cup B$ significa que pelo menos um dos eventos A ou B ocorre. Ele é composto por todos os pontos do espaço amostral, exceto aqueles em que não pertencem nem ao evento A nem ao evento B .

A partir dessas definições, serão introduzidos alguns conceitos básicos sobre probabilidade.

2.2.2 Propriedades Básicas da probabilidade

Definição 4. Propriedade Fundamental - Dado um espaço amostral discreto e finito Ω com pontos amostrais $E_1, E_2, \dots, E_n$ define-se que cada ponto E_i tem um número associado a ele chamado de probabilidade de E_i , denominado por $P\{E_i\}$. Essa probabilidade será definida como um número não negativo e obedecerá a seguinte relação:

$$P\{E_1\} + P\{E_2\} + \dots + P\{E_n\} = 1.$$

Associa-se a probabilidade 0, a um evento E_i , quando ele é impossível de ocorrer, da mesma forma que associamos a probabilidade 1 a um evento que tem 100% de chance de ocorrência.

É importante salientar, que esta definição não se restringe a apenas a espaços amostrais finitos. Pode-se estender a definição para espaços amostrais infinitos da seguinte forma:

Definição 5 – Considere agora um espaço amostral infinito Ω e enumerável¹, isto é:

$$\Omega = \{w_1, w_2, \dots\}$$

Neste caso, sendo $p_1, p_2, \dots$ uma sequência infinita de números satisfazendo as seguintes condições:

$$p_i \geq 0 \quad \forall i = 1, 2, \dots \quad \text{e} \quad \sum_{i=1}^{\infty} p_i = 1$$

Pode-se fazer a seguinte atribuição $P(w_i) = p_i$, ou seja, é possível associar o número p_i como sendo a probabilidade de o evento w_i ocorrer.

Definição 6- Considere dois eventos A e B , a probabilidade que ambos ocorram ou que pelo menos um deles ocorra é dada por:

$$P(A \cup B) = P(A) + P(B) - P(A \cap B)$$

No caso em que eles forem mutuamente exclusivos, temos que $P(A \cap B) = 0$ e, portanto, a expressão acima pode ser escrita como:

$$P(A \cup B) = P(A) + P(B)$$

Definição 7- A probabilidade $P(A)$ de qualquer evento A contido no espaço amostral Ω é dada pela soma das probabilidades de cada ponto amostral contido nele.

No caso, de cada um dos pontos amostrais serem igualmente prováveis, a probabilidade $P(A)$ será dada pela divisão da quantidade de pontos amostrais pertencentes ao evento A pelo número de pontos amostrais do espaço amostral Ω .

Pretende-se agora observar 2 problemas distintos envolvendo o conceito de espaço amostral. Primeiramente, examinemos um problema que trata de um espaço amostral finito.

Problema 1 - Em um lançamento de um dado, qual é a probabilidade de que a face voltada para cima seja par? Neste exemplo, o evento “face voltada para cima ser par” é composto pelos pontos amostrais $\{2\}, \{4\}, \{6\}$, e o evento “lançamento de um dado” é composto pelos pontos amostrais $\{1\}, \{2\}, \{3\}, \{4\}, \{5\}, \{6\}$. Portanto, de acordo com a definição clássica de probabilidade, como cada um dos pontos amostrais é equiprovável, a probabilidade é calculada por:

¹ Um conjunto diz-se enumerável, quando é finito ou quando existe uma bijeção entre ele e $\mathbb{N}$.

$$\frac{\text{casos favoráveis}}{\text{casos totais}} = \frac{3}{6}$$

Analisa-se agora o seguinte problema:

Problema 2 - Considerando que duas pessoas A e B estejam lançando 2 dados não viciados, o vencedor do jogo será a primeira pessoa a tirar 7 na soma dos dois dados.

Ao observar este problema, é necessário considerar duas etapas: A primeira é determinar a probabilidade de um dos jogadores tirar a soma 7 no lançamento dos 2 dados. Nesta parte, depara-se com um espaço amostral finito, pois é possível listar todas as combinações possíveis dos lançamentos dos dados e inferir a probabilidade.

O grande desafio deste exemplo é que, apesar de ser possível calcular a probabilidade de cada lançamento, não é possível definir em qual rodada o jogo vai acabar, o que gera um espaço amostral infinito.

Primeiramente, deve-se calcular a probabilidade de tirar a soma 7 no lançamento de dois dados. Como em cada dado existem 6 possibilidades de resultado, ao lançar os 2 dados, temos um total de $6 \times 6 = 36$ possibilidades. Como o evento “soma das faces igual a 7” é composto pelos pontos amostrais $\{(1,6), (2,5), (3,4), (4,3), (5,2), (6,1)\}$ temos que sua chance de ocorrência é de $\frac{6}{36} = \frac{1}{6}$. A partir disso e das definições 1 e 2, é possível inferir também que a probabilidade de não sair soma 7 é dada por:

$$1 - \frac{1}{6} = \frac{5}{6}$$

Pretende-se agora calcular a probabilidade do participante A ganhar, sabendo que é ele quem inicia o jogo. Para o participante A se tornar vencedor, ele pode ganhar na primeira, segunda, terceira rodada e assim por diante.

1ª Rodada- A probabilidade já foi calculada e é de $\frac{1}{6}$.

2ª Rodada – Para ele ganhar na segunda rodada, tanto ele quanto o seu oponente devem ter “falhado” na primeira rodada, o que acontece em cada caso, com probabilidade de $\frac{5}{6}$ e ele acertar a sua segunda jogada, o que acontece com chance de $\frac{1}{6}$. Logo, a probabilidade de ele ganhar na segunda rodada é dada por:

Raciocinando de forma análoga, temos que a probabilidade de ele ganhar na terceira rodada é de $\left(\frac{5}{6}\right)^2 \cdot \frac{1}{6}$.

Logo, conclui-se que na n -ésima rodada, a probabilidade de vitória do jogador A é de $\left(\frac{25}{36}\right)^{n-1} \cdot \frac{1}{6}$.

Com isso, pela definição 2, a probabilidade do jogador A ganhar é dada por $\sum_{n=1}^{\infty} \left(\frac{25}{36}\right)^{n-1} \cdot \frac{1}{6}$, que é uma série geométrica convergente, cuja soma S_n é dada por:

$$S_n = \frac{\frac{1}{6}}{1 - \frac{25}{36}} = \frac{6}{11}.$$

Com isso, a chance do jogador B ganhar será de $\frac{5}{11}$. Deve-se observar que a chance do jogador A ganhar é maior do que a do jogador B ; isto se dá pela vantagem que ele tem ao começar o jogo.

É possível perceber que a probabilidade de vitória na rodada atual é menor do que a na rodada anterior, até chegarmos a um ponto em que ela possa ser tão pequena, quanto quisermos. Esse resultado é oriundo do seguinte teorema:

Se $\sum a_n$ é uma série convergente, então $\lim_{n \rightarrow \infty} a_n = 0$, relativo ao comportamento do termo geral de uma série convergente

2.3 Problema de D'Alembert

O primeiro problema do livro “Andar do Bêbado” que será analisado trata de um equívoco cometido pelo matemático D'Alembert. O problema, embora seja relativamente simples, apresenta um erro comumente cometido pelas pessoas ao resolverem problemas envolvendo espaços amostrais não equiprováveis. O problema analisado por ele é o seguinte: “Qual é a probabilidade de sair uma cara no lançamento de duas moedas”? (MLODINOW, 2009, p. 40).

Quando são lançadas duas moedas, existem três possibilidades para o número de caras: 0, 1 ou 2 caras. Baseado nisso, como são três possibilidades, D'Alembert concluiu que a chance de sair uma cara era de $\frac{1}{3}$, o que não corresponde à verdadeira probabilidade. O erro cometido foi o de utilizar a definição clássica de probabilidade em um espaço amostral no qual os eventos não são equiprováveis. Para resolver adequadamente este problema, utilizaremos a ideia proposta por Feller (1968), listando todos os eventos simples contidos no espaço amostral. Com isso, o espaço amostral é composto pelos quatro eventos simples, listados a seguir: {(cara, cara), (coroa, coroa), (cara, coroa), (coroa, cara)}, dentre os quais,

dois deles satisfazem a condição de sair duas caras. Sendo assim, a probabilidade que estamos procurando é de $\frac{2}{4} = \frac{1}{2}$.

Esta técnica de decompor eventos não equiprováveis em eventos mais simples de forma a termos um espaço amostral equiprovável é bastante útil, conforme veremos nos problemas a seguir.

2.4 O problema dos pneus

Quem nunca precisou faltar a uma prova? Seja por motivo de doença ou por não estar preparado para fazer aquela avaliação, todos um dia já necessitaram se ausentar de um exame acadêmico. Em muitas dessas vezes, ao tentarem se justificar, contam fatos que não parecem ser muito verossímeis.

No problema a seguir, o autor relata a situação de dois alunos que faltaram ao exame final e deram uma desculpa esfarrapada. Estariam eles falando a verdade? Veremos no problema a seguir, que foi retirado da coluna Ask Marilyn (1996).

Meu pai ouviu esta história no rádio. Na Universidade Duke, dois alunos receberam nota máxima em química ao longo de todo o semestre. Mas na véspera da prova final, os dois foram a uma festa em outro estado e não voltaram para a prova final. Disseram ao professor que um pneu do carro havia furado e perguntaram se poderiam fazer uma prova de segunda chamada. O professor concordou, escreveu uma segunda prova e mandou os dois para salas separadas. A primeira pergunta valia 0,5 ponto. A segunda pergunta valia 9,5 pontos e dizia qual era o pneu? Qual é a probabilidade de eles darem a mesma resposta? (MLODINOW, 2009, p. 42).

Mlodinow relata que ao indagar seu pai sobre a chance de os alunos acertarem a resposta, ouviu como resposta $\frac{1}{16}$. Esta resposta está claramente equivocada. Mas então, qual será a resposta correta? E o que leva a este tipo de raciocínio falacioso?

O argumento que induz a este tipo de pensamento é o seguinte: “Como cada aluno tem 4 pneus, a chance de ele escolher um pneu específico é de $\frac{1}{4}$. Como são dois alunos, a chance de eles escolherem o mesmo pneu é de $\frac{1}{16}$. (MLODINOW, 2009, p. 42)”. Neste caso, apesar do uso correto da regra da probabilidade da multiplicação, a solução do problema está errada. Com isso, pretende-se analisar a seguir duas formas de resolver corretamente este problema.

Na primeira delas, será feito uso da solução dada pelo pai de Mlodinow (2009). Para isso, será completada a solução proposta por ele. O erro desta solução se dá pelo fato de que nela é considerada apenas a probabilidade de ambos os alunos escolherem como resposta um determinado pneu, por exemplo, o pneu dianteiro esquerdo. Para ajustarmos esta probabilidade, será considerado todos os pontos amostrais contidos no espaço amostral tais que os dois pneus são iguais. São eles: {(dianteiro esquerdo, dianteiro esquerdo), (dianteiro direito, dianteiro direito), (traseiro esquerdo, traseiro esquerdo), (traseiro direito, traseiro

direito)). Através da solução do pai de Mlodinow, vê-se que cada uma destas soluções tem probabilidade de $\frac{1}{16}$. Como são 4 possibilidades, a probabilidade de os dois pneus serem iguais será dada por $4 \times \frac{1}{16} = \frac{4}{16} = \frac{1}{4}$.

Outra possibilidade de resolução deste problema é listar todos os casos possíveis e os casos favoráveis (ambos responderem o mesmo pneu), e calcular a probabilidade através da razão entre os casos favoráveis e os casos possíveis. Como cada um deles tem 4 escolhas acerca de qual dos pneus furou, temos $4 \times 4 = 16$ possibilidades diferentes de escolha, dentre as quais 4 são favoráveis. Logo, como cada uma das escolhas é equiprovável, a probabilidade de ambos terem escolhido o mesmo pneu é de $\frac{4}{16} = \frac{1}{4}$ (MLODINOW, 2009).

As duas soluções, apesar de distintas, possuem algumas similaridades. Na primeira, é feito o cálculo da probabilidade de ambos os pneus serem iguais e depois é calculado em quantos casos isso ocorre, enquanto na segunda, primeiro é analisado em quantos casos ocorre o evento "pneus iguais" e só depois é calculada a probabilidade de ocorrência deste evento.

2.5 É menino ou menina?

Muitos pais ficam curiosos sobre o sexo do seu futuro filho, inclusive alguns deles querem conhecê-lo antes do parto. Imagine pais de gêmeos? Será que será um casal ou um menino e uma menina? O problema a seguir aborda alguns cálculos probabilísticos acerca deste tema. Ele ilustra como uma concepção errônea do espaço amostral pode alterar o cálculo da probabilidade. Observe o problema a seguir: “Suponha que uma mãe está grávida de gêmeos fraternos e quer saber qual é a probabilidade de nascerem duas meninas, dois meninos ou um menino e uma menina” (MLODINOW, 2009, p. 41).

Para resolver o problema, basta enumerar os elementos do espaço amostral. Como temos dois filhos, as combinações possíveis são: (menino, menina), (menino, menino), (menina, menino), (menina, menina).

Portanto, as probabilidades são dadas por:

- **Dois meninos**

Tem-se apenas um elemento que satisfaz essa exigência. Logo, a probabilidade de isso ocorrer é de $\frac{1}{4}$, ou 25% .

- **Duas meninas**

Raciocinando de forma análoga ao caso de 2 meninos, concluímos que a probabilidade de isso acontecer também é de $\frac{1}{4}$, ou 25% .

- **Um menino e uma menina**

Já neste caso, há 2 casos favoráveis, (menino, menina) ;(menina, menino), logo a probabilidade da ocorrência deste caso é de $\frac{2}{4} = \frac{1}{2}$, ou 50%.

Neste problema das duas filhas, geralmente é feita uma pergunta adicional: qual é a chance, dado que um dos bebês seja uma menina, de que ambas sejam meninas?

Essa nova informação altera drasticamente o cálculo da probabilidade, pois a restrição de um dos filhos ser menina, altera a quantidade de elementos do espaço amostral. Neste tipo de problema, é necessário reescrevê-lo e calcular a probabilidade em cima do novo espaço amostral. Com isso, há o novo espaço amostral formado por: {(menina, menina), (menino, menina), (menina, menino)} (MLODINOW, 2009).

Como destes 3 casos, em apenas um deles temos uma situação favorável, a probabilidade pedida é de $\frac{1}{3}$.

A seguir analisa-se uma outra variação deste problema, formulada por Mlodinow (2009), na página 79 do livro “O Andar do Bêbado”. Em uma família com duas crianças, qual é a probabilidade, se uma delas for uma menina chamada Flórida, de que ambas sejam meninas?

A princípio, pode-se imaginar que o nome de uma das meninas se chamar Flórida não interfere na probabilidade, tornando este problema com resolução idêntica ao anterior. Mas, ao analisarmos o espaço amostral desta variante com maior cuidado, perceberemos que ele não é o mesmo em relação à questão anterior (MLODINOW, 2009).

Antes de enumerar os possíveis resultados para este experimento, é possível descartar a possibilidade das duas meninas se chamarem Flórida, pois de acordo com Mlodinow (2009, p. 83), a taxa de registro de meninas com o nome Flórida era de 1:1000000, além de ser bastante estranho os pais darem nomes iguais aos seus filhos. Com base nisso, representaremos *MF*, como menina chamada Flórida e *MNF*, como menina que não se chama Flórida. Com isso, ficamos com o seguinte espaço amostral:

$$\{(MF, MNF), (MF, Menino), (Menino, MF), (MNF, MF)\}$$

Tem-se, então, que o espaço amostral é composto por 4 elementos, dos quais 2 deles se encaixam no que o problema pede. Portanto, a probabilidade de ocorrência deste evento é de $\frac{2}{4} = \frac{1}{2}$, ou 50%.

Nos problemas resolvidos acima, é comum ocorrerem alguns erros acerca do conceito de espaço amostral. O primeiro deles se dá no exemplo 1. Neste caso, como se tem três possibilidades (os dois serem meninos, os dois serem meninas, ou termos um menino e uma menina), o senso comum acaba por levar a acreditar que cada uma das possibilidades tem probabilidade igual a $\frac{1}{3}$.

O problema com esse tipo de raciocínio é a não percepção de que o espaço amostral de interesse não é equiprovável, portanto, não é possível a aplicação direta da definição clássica de probabilidade. Para resolver esse tipo de problema, uma estratégia é descrever o espaço amostral de forma que ele se torne equiprovável e, em seguida, calcular a probabilidade.

No segundo exemplo, ocorre uma restrição do espaço amostral, impondo que um dos filhos seja menina. Nesse caso, como um dos filhos é menina, é razoável imaginar que existe igual possibilidade de termos um menino e outra menina, ou seja, ambos os casos têm a chance de $\frac{1}{2}$, de ocorrer. Contudo, o problema desse tipo de abordagem é que em nenhum momento o problema diz que o primeiro filho é a menina, apenas que um deles é menina. Com isso, teríamos 3 ao invés de 2 possibilidades, das quais uma delas satisfaz à condição apresentada no problema.

No terceiro exemplo, houve a introdução de um elemento que, a princípio, não alteraria a probabilidade. Nesse caso, eliminamos a chance de ambas se chamarem Flórida (pois é bastante improvável) e analisamos as possibilidades. Além disso, outro erro recorrente é não considerar separadamente os pontos amostrais (MF, MNF) e (MNF e MF), o que acarreta erros na enumeração do espaço amostral. Esse acréscimo do nome de uma das meninas provocou um aumento do espaço amostral de 3 para 4 elementos, pois agora o ponto amostral (menina, menina) se subdividiu em dois (MLODINOV, 2009).

2.6 -O problema de Monty-Hall

Esse problema foi proposto pela coluna Ask Marilyn, em setembro de 1990, baseado em um programa de auditório chamado “*Let’s Make a Deal*”, transmitido entre os anos de 1963 e 1976, apresentado por Monty-Hall. Este quadro foi feito também aqui no Brasil, com o nome de Porta dos Desesperados, ancorado pelo humorista Sergio Mallandro. A relevância

desta questão é tão grande que ela é utilizada em diversos cursos de probabilidade, sendo inclusive retratada no cinema, com o filme *Quebrando a Banca* (2008), de Robert Luketic. A seguir temos o problema retirado do livro *O Andar do Bêbado*.

Suponha que os participantes de um programa de auditório recebam a opção de escolher uma dentre 3 portas: atrás de uma delas há um carro, atrás da outra há cabras. Depois que um dos participantes escolhe uma porta, o apresentador que sabe o que há atrás das três portas, abre uma das portas não escolhidas, revelando uma cabra. Então ele diz para o participante: Você gostaria de mudar a sua escolha para a porta fechada? Para o participante é vantajoso fazer a troca? (MLODINOW, 2009, p. 35).

A resolução deste problema é ao mesmo tempo simples e intrigante. O senso comum nos leva a crer que indifere caso se realize a troca ou não, pois após ao apresentador eliminar uma das portas, vão restar duas portas, portanto a probabilidade de cada uma delas ser a porta com o prêmio é de 50%. Marilyn contradiz este pensamento, afirmando que a troca de porta seria algo vantajoso e daria uma porcentagem de aproximadamente 66% de vitória. Ela foi duramente criticada pelos seus leitores, que estavam revoltados com o seu “erro”, sendo inclusive alvo de críticas por professores universitários da área de matemática. Como ela poderia ter errado algo tão simples e intuitivo? A resposta é não. Ela não errou. 92% dos americanos estavam errados, inclusive professores que lecionavam cursos sobre probabilidade. A seguir, iremos mostrar de forma simples, o motivo de Marilyn estar correta (MLODINOW, 2009).

Para ilustrar esta situação vamos supor que o prêmio esteja na porta de número 1. Os casos em que o prêmio estiver nas portas 2 e 3 são explicados de forma análoga. Com isso, pretende-se analisar três casos: O primeiro, a pessoa escolheu a porta de número 1 (a premiada), o segundo a pessoa escolheu a porta de número 2 e no terceiro a pessoa escolheu a porta de número 3 (MLODINOW, 2009).

No **Caso 1**, em que a pessoa escolheu a porta de número 1:

- Como o prêmio está na mesma porta que o jogador escolheu, o apresentador mostra a porta de número 2 ou 3, revelando uma cabra. Após isso, ele perguntará ao candidato se deseja trocar ou não de porta.
- Se o jogador decidir manter a porta escolhida, ele **GANHA**.
- Se o jogador decidir trocar de porta, ele **PERDE**.

No **Caso 2**, em que a pessoa escolheu a porta de número 2:

- O apresentador terá de mostrar obrigatoriamente a porta de número 3, revelando uma cabra. Após isso, perguntará ao participante se ele deseja trocar ou manter a porta.

Com base na escolha, se tem o seguinte:

- Se o jogador decidir manter a porta escolhida, ele PERDE.
- Se o jogador decidir trocar a porta escolhida, ele GANHA.

No **Caso 3**, em que a pessoa escolheu a porta de número 3:

- O apresentador terá de mostrar obrigatoriamente a porta de número 2, revelando uma cabra. Após isso, perguntará ao participante se ele deseja trocar ou manter a porta.

Com base na escolha, conclui-se o seguinte:

- Se o jogador decidir manter a porta escolhida, ele PERDE.
- Se o jogador decidir trocar a porta escolhida, ele GANHA.

Após analisar os 3 casos percebemos que o jogador só ganha mantendo a escolha inicial em um deles (caso 1). Logo, a probabilidade de vitória mantendo a escolha inicial é de $\frac{1}{3}$, ou 33,3%. Se ao invés disso, o jogador mudar a porta que escolheu inicialmente, ele ganhará em dois dos casos (caso 2 e caso 3). Logo, a sua probabilidade de vitória é de 2, ou 66,6%. A figura 2, ilustra de forma resumida este raciocínio. A Figura 2 ilustra este raciocínio.

Figura 2 – Portas da Ilusão



Fonte : O Autor

Portanto, demonstra-se com um argumento bem simples, sem precisar de uma matemática rebuscada, que Marilyn estava certa.

- **Mas, o que leva o participante a cometer este erro?**

O pensamento errado ocorre porque o participante considera as probabilidades após o apresentador mostrar a porta com a cabra, enquanto a probabilidade relevante é aquela no início do jogo, quando ele tinha 33,3% de chance de escolher a porta certa e 66,6% de chance de escolher a errada. No entanto, a chance que deve ser estimada é a de permanecer entre uma porta e outra, a qual erroneamente é atribuída como 50%. Além disso, o fato de o apresentador ter eliminado apenas uma porta e mantido a outra, cria a ilusão de que a escolha inicial é irrelevante. Se houvesse 10 portas ao invés de 3 e o apresentador eliminasse 8 delas, talvez a intuição nos levasse a trocar de porta com mais frequência, o que seria uma boa ideia, visto que, de forma análoga ao problema anterior, a primeira escolha nos dá uma probabilidade de $\frac{1}{10}$, em ganhar e, portanto, a troca de porta de dá uma chance de $\frac{9}{10}$, (ou seja, 9 vezes maior).

Mas e se ao invés de abrir as 8 portas, o apresentador abrisse uma porta só? O que seria melhor, manter a escolha inicial, trocar para uma das 8 portas restantes, ou seria indiferente?

Para o exemplo de 10 portas, o jogador tem uma chance de $\frac{1}{10}$, em acertar a porta com o prêmio em sua primeira tentativa e, portanto, a chance de o prêmio estar em uma das outras 9 portas é de $\frac{9}{10}$. Quando o apresentador, que sabe onde está o prêmio, abre uma das portas, restam apenas 8 portas além da que o jogador escolheu. Como a chance de o prêmio estar em uma porta que não seja a primeira é de $\frac{9}{10}$ e todas as portas têm igual chance de ser a premiada, a probabilidade de o apostador encontrar o prêmio ao fazer a troca é de $\frac{9}{10} \div 8 = \frac{9}{80} = 11,25\%$, que é superior a probabilidade de vitória caso ele mantenha a escolha inicial (neste caso a chance de vitória seria de 10%)

De forma análoga, caso o apresentador elimine 2 portas (restando assim 7 portas além da escolhida pelo jogador), a probabilidade de vitória caso o jogador escolha trocar para uma das portas remanescentes será dada por $\frac{9}{10} \div 7 = \frac{9}{70}$, ou aproximadamente 12,85%. Para este caso de 10 portas, pode-se concluir que a probabilidade de vitória caso o jogador troque de porta vai depender de quantas portas o apresentador irá abrir. Caso o apresentador abra m portas, restaram $10 - m - 1$ escolhas para o jogador caso ele queira trocar de porta e, portanto, a chance de acerto será de $\frac{9}{10} : (10 - m - 1) = \frac{9}{10(10-m-1)}$.

É possível perceber que, à medida que aumenta o número de portas abertas pelo apresentador, o denominador da fração que representa a probabilidade de vitória ao trocar de porta diminui, e, portanto, a probabilidade aumenta. Com isso, como a probabilidade de vitória já é maior ao trocar de porta, quando o apresentador abre apenas uma porta, denota-se que a troca é sempre vantajosa.

A tabela 4, ilustra as probabilidades de vitória de acordo com o número de portas que o apresentador abriu:

Tabela 4 – Relação entre o número de portas abertas e as probabilidades para 10 portas

Número de portas abertas	Probabilidade de vitória mantendo a escolha	Probabilidade de vitória trocando a escolha
1	$\frac{1}{10} = 10\%$	$\frac{9}{80} = 11,25\%$
2	$\frac{1}{10} = 10\%$	$\frac{9}{70} \cong 12,85\%$
3	$\frac{1}{10} = 10\%$	$\frac{9}{60} = 15\%$
4	$\frac{1}{10} = 10\%$	$\frac{9}{50} = 18\%$
5	$\frac{1}{10} = 10\%$	$\frac{9}{40} = 22,5\%$
6	$\frac{1}{10} = 10\%$	$\frac{9}{30} = 30\%$
7	$\frac{1}{10} = 10\%$	$\frac{9}{20} = 45\%$

Fonte: O autor

Com isso, é possível generalizar este problema para o caso de m portas em que o apresentador abra n delas, com $n < m - 1$. Como são m portas, a probabilidade de o jogador acertar na sua primeira escolha é de $\frac{1}{m}$ e, portanto, a probabilidade de a porta escolhida não ser a do prêmio é de $\frac{m-1}{m}$. Como o apresentador irá abrir n portas, restará ao jogador $(m - n - 1)$ opções caso ele deseje fazer a troca, e, portanto, a probabilidade de vitória neste caso será dada por $\frac{m-1}{m(m-n-1)}$.

Sendo assim, para demonstrar que esta probabilidade é sempre maior que a probabilidade de se ganhar mantendo a escolha original, basta mostrar que $\frac{m-1}{m(m-n-1)} > \frac{1}{m}$, ou em outras palavras que $\frac{m-1}{m-n-1} > 1$, ou ainda que $m - 1 > m - n - 1$, ou seja $n > 0$, que é sempre verdade. Com isso, conclui-se que a chance de vitória será sempre maior trocando a porta, independentemente de quantas portas o apresentador abriu.

3. PROBABILIDADE E JOGOS DE AZAR

O presente capítulo explora o fascinante mundo da probabilidade e sua aplicação em jogos de azar. Ao longo deste capítulo, serão utilizados alguns conceitos fundamentais de probabilidade já discutidos, como espaços amostrais, eventos, e cálculo de probabilidades, além de sua relevância no contexto de jogos de azar. Utilizaremos também o conceito de valor esperado ou esperança matemática. Serão analisados diversos exemplos e casos práticos, revelando como a compreensão dos princípios probabilísticos podem influenciar estratégias de jogo e tomadas de decisão. Ao final do capítulo, os leitores terão adquirido uma sólida compreensão das bases matemáticas por trás dos jogos de azar e como a probabilidade pode ser utilizada para uma melhor compreensão das chances de vitória neste tipo de jogo.

3.1 Contexto Histórico

Apesar do estudo sistemático da probabilidade ser uma prática mais contemporânea, desde a antiguidade existiam questões relacionadas a ela. Os gregos, apesar de possuírem um vasto conhecimento em geometria, ainda eram leigos acerca da influência da aleatoriedade em suas vidas. Segundo Mlodinow (2009), isso se deve ao misticismo profundamente enraizado na sociedade da época, o que dificultava um entendimento mais profundo das teorias de probabilidades. Um exemplo disso era um jogo de apostas com ossos tridimensionais de carcaça de animais, chamados astrágalos, no qual os gregos buscavam a orientação dos oráculos. O resultado do jogo era interpretado como uma manifestação da vontade dos deuses, tornando irrelevante um melhor entendimento do acaso naquela época.

Por outro lado, os romanos adotaram uma abordagem mais pragmática da matemática. Enquanto os gregos se dedicavam a teoremas da geometria, os romanos priorizavam a resolução de problemas cotidianos, como calcular o comprimento de um rio enquanto enfrentavam inimigos na margem oposta. Essa mentalidade, aliada a outros traços culturais, levou os romanos a considerarem mais cuidadosamente as questões relacionadas ao acaso. Cícero, um dos precursores da ideia de probabilidade na antiguidade, utilizou argumentos probabilísticos para desafiar a crença de que o sucesso nas apostas era obra dos deuses. Mesmo sem compreender as leis que regem o acaso, Cícero argumentava que, ao jogar muitas vezes, um jogador eventualmente teria uma jogada de sorte, como a jogada de Vênus, por mera casualidade (MLODINOW, 2009).

Essa abordagem de Cícero pode parecer trivial à luz do conhecimento matemático atual, mas na época representou uma significativa quebra de paradigma. Seu pensamento

contribuiu para o avanço da sociedade romana em diversas áreas, especialmente no campo do direito. Assim, mesmo em uma era onde o misticismo dominava o pensamento, as sementes da compreensão probabilística foram plantadas, eventualmente dando frutos em sociedades futuras (MLODINOW, 2009).

Apesar disso, apenas no século XVI, Cardano escreveu o primeiro livro sobre probabilidade, mais uma vez associando-a aos jogos de azar. Em 1516, mudou-se para Pavia para estudar medicina e, necessitando de dinheiro para financiar seus estudos, começou a apostar em jogos de azar. Com seus ganhos, economizou mais de mil coroas, o suficiente para pagar sua faculdade e ingressar na universidade em 1520.

Pouco depois, Cardano começou a desenvolver sua teoria sobre apostas. Seu livro “Os Jogos de Azar” abordou gamão, cartas e astrágalos, introduzindo conceitos importantes em probabilidade, como o espaço amostral. Embora tenha analisado apenas espaços equiprováveis e conter alguns erros na análise das probabilidades, é inegável sua relevância como o primeiro trabalho bem-sucedido na tentativa de compreender a incerteza.

Mesmo nos dias de hoje, as crenças e superstições estão profundamente enraizadas na sociedade. As pessoas atribuem ocorrências de eventos aleatórios a objetos de sorte e realizam rituais antes de empreender atividades importantes em suas vidas. Ao longo deste capítulo, será apresentada uma coleção de problemas relacionados a jogos de azar, desde a Grécia antiga até os dias atuais, analisando a visão predominante das pessoas sobre a probabilidade.

3.2 Os astrágalos e a Grécia Antiga

Como observado, o principal “jogo de azar” na Grécia antiga era o jogo de astrágalos. Apesar da forte associação com os deuses, este jogo apresentava algumas peculiaridades matemáticas intrigantes. A jogada de Vênus, considerada a melhor jogada possível, era extremamente rara. Nessa jogada, todos os quatro astrágalos caíam com lados diferentes voltados para cima. Estudiosos modernos, conforme descrito por Mlodinow (2009), apontam que, devido à anatomia dos ossos, cada lado não possuía a mesma probabilidade de cair voltado para cima. Dois dos lados tinham uma probabilidade de 10% de cair com a face para cima (denominados L_1 e L_2), enquanto os outros dois lados tinham uma probabilidade de 40% de cair com a face para cima (denominados L_3 e L_4).

Com base nisso, é possível calcular a probabilidade de uma jogada de Vênus aplicando a regra do produto:

$$\frac{10}{100} \times \frac{10}{100} \times \frac{40}{100} \times \frac{40}{100} = \frac{16}{10000}$$

Neste ponto, é crucial analisarmos com cuidado, pois o resultado obtido até agora não representa a probabilidade que buscamos. O cálculo realizado aborda apenas uma das combinações possíveis (por exemplo, L_1, L_2, L_3, L_4), mas há outras possibilidades, como L_2, L_3, L_1, L_4 ou L_1, L_4, L_3, L_2 . É necessário observar que existem 4 escolhas para a primeira posição, 3 para a segunda, 2 para a terceira e 1 para a quarta. Portanto, há $4 \times 3 \times 2 \times 1 = 24$ combinações possíveis a serem consideradas. Portanto, a probabilidade de se fazer uma jogada de Vênus é de $24 \times \frac{16}{10000} = \frac{384}{10000}$.

Apesar de os gregos a considerarem esta jogada com uma manifestação dos deuses na vida do jogador, analisar-se-á a seguir que ela não é a mais improvável de se acontecer. A jogada mais rara é todos caírem com o lado L_1 para cima ou todos caírem com o lado L_2 para cima. Estas duas sequências só ocorrem em 2 combinações (L_1, L_1, L_1, L_1) e (L_2, L_2, L_2, L_2). Cada uma dessas probabilidades é dada por:

$$\frac{10}{100} \times \frac{10}{100} \times \frac{10}{100} \times \frac{10}{100} = \frac{1}{10000}$$

Vale ressaltar que mesmo que todos os lados fossem equiprováveis, a combinação de cair todos os astrágalos no mesmo lado ainda seria a jogada mais rara, pois neste caso todas as combinações saem com a mesma chance. Essa jogada na qual as faces voltadas para cima serem todas iguais, ocorrem em 4 combinações, a saber (L_1, L_1, L_1, L_1), (L_2, L_2, L_2, L_2), (L_3, L_3, L_3, L_3) e (L_4, L_4, L_4, L_4). Cada uma delas teria a chance de $\frac{1}{4} \times \frac{1}{4} \times \frac{1}{4} \times \frac{1}{4} = \frac{1}{256}$. Cada uma das combinações possíveis para a jogada de Vênus teria a mesma probabilidade. No entanto, como já destacado, existem 24 combinações que levam à jogada de Vênus, resultando em uma probabilidade seis vezes menor de ocorrer a jogada com lados iguais do que a jogada de Vênus.

3.3 Devo estacionar aqui?

O problema apresentado levanta a questão da tomada de decisões na vida cotidiana, ressaltando como muitas vezes o ser humano age por impulsividade sem considerar completamente as consequências das ações cometidas individualmente. Embora não esteja diretamente relacionado aos jogos de azar, ele ilustra como decisões aparentemente simples podem envolver riscos significativos. A esse respeito, destaca-se o seguinte:

Ao questionar o custo de estacionar em um determinado local, o autor menciona um exemplo em que o valor aparentemente baixo de 25 centavos é subvertido pela probabilidade de receber uma multa de US\$ 40 aproximadamente uma vez a cada 20 vezes que se estaciona lá. Isso demonstra como avaliar os riscos envolvidos em uma decisão pode ser crucial, mesmo em situações aparentemente triviais (MLODINOW, 2009, p. 58).

Para resolver esse problema, é necessário introduzir o conceito de esperança matemática, que desempenhará um papel fundamental na abordagem de outros problemas neste capítulo.

Diz-se que X é uma variável aleatória discreta, se os seus possíveis valores puderem ser listados como $x_1, x_2, \dots$. Essa variável pode ser finita ou infinita. No caso dela ser finita, esta lista possui um valor final x_n , e no caso dela ser infinita ela continua indefinidamente.

Definição 8: Seja X uma variável aleatória discreta que assume os valores $\{x_1, x_2, \dots, x_n\}$, com probabilidades $p_1, p_2, \dots, p_n$, então define-se a esperança ou valor esperado de X , da seguinte maneira:

$$E(X) = \sum_{i=1}^n x_i p_i$$

Após estas definições pode-se concluir que ao avaliar o custo do estacionamento não devemos considerar somente o valor de 25 centavos. Utilizando a definição acima de esperança matemática, temos que o valor pago será de $0,25 + \frac{1}{20} \times 40 + 0 \times \frac{19}{20} = \text{US\$ } 2,25$ pois além do valor fixo de \$ 0,25, 1 em cada 20 vezes ele tem que pagar a multa de US\$40. Este problema apesar de ser bem simples, ilustra bem o conceito de esperança matemática, o qual será aplicada nos problemas a seguir.

3.4 Será que é mesmo vantajoso?

Muitas pessoas têm o desejo de se tornarem milionárias, e uma das maneiras consideradas “fáceis” para isso é apostar em jogos de loteria. Existe uma ideia difundida de que um dia se ganhará na loteria e se poderá parar de trabalhar. No entanto, sabe-se que as chances de ganhar nesse tipo de sorteio são bastante baixas, mesmo assim muitos continuam a jogar. Os problemas que serão apresentados a seguir abordam um pouco desse universo dos jogos de azar, destacando o fato de que eles são projetados para que os apostadores percam dinheiro. O primeiro exemplo dessa dinâmica está descrito na página 58 do livro “O Andar do Bêbado”, que descreve um sorteio “gratuito”, onde bastava enviar uma carta para concorrer a um grande prêmio em dinheiro.

Um sorteio anunciado recentemente pelo correio oferecia um prêmio máximo de US\$55 milhões. Bastava mandar uma carta com a inscrição. Não havia limite para o número de inscrições feitas por cada pessoa, mas cada uma delas tinha que ser enviada separadamente. Os patrocinadores esperavam cerca de 200 milhões de inscrições, pois as letras miúdas indicavam esta probabilidade de ganhar o prêmio era de $\frac{1}{200000000}$ (MLODINOW, 2009, p. 58).

Pode-se utilizar o conceito de esperança matemática para mensurar o valor “justo”² do envio da carta. O prêmio do concurso é estipulado em US\$ 55 milhões e, segundo os organizadores, a probabilidade de vitória é de $\frac{1}{200000000}$. É possível calcular um valor “justo” para o envio da carta, multiplicando o prêmio pela probabilidade de vitória da pessoa. Logo este valor deveria ser de $\frac{1}{200000000} \times 55000000 = 0,275$, ou aproximadamente US\$ 0,27. Este valor reflete o que o apostador esperaria pagar, caso toda a arrecadação fosse revertida para o prêmio final (MLODINOW, 2009).

É claro que a maioria dos sorteios não funciona desta forma, sendo, portanto, necessário que o valor pago pelo envio da carta seja maior que US\$ 0,27. Com isso, segundo Mlodinow (2009), os grandes ganhadores foram os próprios correios, pois descontando investimentos em propaganda e o valor pago de US\$ 55 milhões, arrecadaram US\$ 80 milhões.

3.5 Brasil e o jogo de apostas

No Brasil, o jogo de apostas mais famoso é a Mega Sena. Nele, são sorteadas 6 dezenas dentre um total de 60. Portanto, a quantidade de jogos correspondentes à aposta mínima (escolher 6 números dentre os 60) é dada por:

$$C_{60,6} = \frac{60!}{54!6!} = 50.063.860$$

Ou seja, a probabilidade de uma pessoa ganhar com a aposta mínima é de $\frac{1}{50063860}$. É possível então, a partir disso, calcular o valor justo para o bilhete. Para tal, deve-se considerar um prêmio de 39 milhões, que foi o prêmio do sorteio 2595, ao qual foi realizado no dia 24/05/2023. Utilizando o conceito de valor esperado, multiplicaremos o valor do prêmio pela probabilidade de vitória fazendo um jogo simples. Com isso, o valor justo do bilhete será dado por:

² O termo valor justo é utilizado no texto com o significado do valor a ser pago de modo a toda a arrecadação ser revertida em premiação.

$$\frac{1}{50063860} \times 39000000 = 0,78$$

É claro que o valor de um jogo simples é bem maior que este, à época custando R\$ 5,00. Para tentar maximizar a probabilidade de ganhar, o jogador pode jogar mais números. A tabela 5, mostra a quantidade de dezenas jogadas, valor da aposta, probabilidade de acertar as 6 dezenas e o valor “justo” das apostas (considerando o mesmo prêmio de 39 milhões):

Tabela 5 – Relação entre as dezenas apostadas e a probabilidade e vitória (24/05/2023)

Dezenas	Valor da Aposta	Probabilidade (1 em)	Valor Justo
7	R\$ 35,00	7.151.980	R\$ 5,45
8	R\$ 140,00	1.787.995	R\$ 21,81
9	R\$ 420,00	595.998	R\$ 65,43
10	R\$ 1050,00	238.399	R\$ 163,59
11	R\$ 2310,00	108.363	R\$ 359,90
12	R\$ 4620,00	54.182	R\$ 719,79
13	R\$ 8580,00	29.175	R\$ 1.336,76
14	R\$ 15.015,00	16.671	R\$ 2.339,39

Fonte: LOTERIAS CAIXA, 2023.

Nesta tabela, é observada a relação entre o número de dezenas jogadas, o valor a ser pago pela aposta, a chance de vitória e o "valor justo" que deveria ser pago por cada bilhete. Para calcular a chance de vitória em cada um dos casos, basta pensar em quantas sequências de 6 números podem ser formadas em cada uma das apostas. Por exemplo, ao fazer a aposta com sete dezenas, podem ser formadas 7 sequências com 6 dezenas ($C_{7,6}$), portanto, a chance de vitória será $7 \times \frac{1}{50063860}$. Da mesma forma a chance de ganhar apostando 8 dezenas é dada por $28 \times \frac{1}{50063860}$ pois se pode formar 28 sequências de 6 dezenas com esta aposta

Generalizando, em uma aposta de n dezenas, a probabilidade de vitória é dada por $C_{n,6} \times \frac{1}{50063860}$. Já o valor justo foi baseado em uma premiação de R\$ 39.000.000,00 (prêmio do concurso 2595, realizado em 24/05/2023). Para encontrá-lo, basta multiplicar o prêmio oferecido pela probabilidade de vitória do apostador.

- **Mas e se o “valor justo” fosse maior que o valor do bilhete, o que aconteceria?**

Geralmente isso acontece aqui no Brasil na Mega Sena da Virada. A Mega- Sena da virada do ano de 2022 pagou um prêmio de R\$ 541.969.966,29 para quem acertasse as 6 dezenas, ou seja, o “valor justo” pelo bilhete seria de $\frac{541969.966,29}{50063860}$, o que resulta aproximadamente em R\$ 10,82, valor bem maior que os R\$ 5,00 cobrados por cada bilhete. Na prática, isso significa que o apostador gastaria menos dinheiro apostando em todas as combinações possíveis de 6 dezenas do que o valor da premiação. Neste caso, para jogar todas as combinações possíveis, o apostador gastaria $R\$ 5,00 \times 50063860 = R\$ 250.319.300,00$ conseguindo assim aproximadamente quase R\$310.000.000,00 de lucro. Mas quais são os problemas disso?

- Primeiro, o processo de preenchimento dos jogos é bastante complicado. Um apostador não teria tempo hábil para fazê-lo sozinho, sendo necessário que um grupo de investidores realize essa operação.
- Segundo, mesmo que fosse possível fazê-lo, a Mega Sena da Virada tem um apelo muito grande entre as pessoas, o que significa que muitas apostas são feitas. Com isso, a chance de duas ou mais pessoas ganharem o prêmio máximo é considerável. A título de curiosidade, de 2009 a 2022, nenhuma pessoa ganhou o prêmio sozinha. Em 2009, houve 3 ganhadores, e em 2018, foi registrado um recorde de 52 ganhadores. Isso implica que uma pessoa que apostasse em todos os jogos na Mega Sena do ano de 2022 só teria lucro se ganhasse o prêmio sozinha ou no máximo dividido com alguém. Na realidade, se ela conseguisse jogar todos os jogos, teria um grande prejuízo, pois o concurso teve um total de 5 ganhadores, o que resultaria em um prejuízo de cerca de R\$ 142.000.000,00, considerando que cada ganhador recebeu aproximadamente R\$ 108.000.000,00.

3.6 A loteria de Virgínia

O problema a seguir retrata uma situação semelhante àquela que ocorre nas Mega-Senas da virada aqui no Brasil, ou seja, é possível jogar todas as combinações com um custo menor que o valor do prêmio. Será analisado que, ao contrário das megasenas da virada, nesta loteria existem peculiaridades que tornam possível para um apostador fazer todos os jogos e obter lucro.

Em 1992, alguns investidores de Melbourne, na Austrália, notaram que a loteria de Virgínia violava esse princípio. A loteria consistia em escolher 6 números de 1 a 44. Se encontrarmos um triângulo de Pascal com esse número de linhas, veremos que existem 7.059.052 maneiras de selecionar 6 números de um grupo de 44 possíveis. O prêmio máximo da loteria era de US\$ 27 milhões, e incluindo o segundo, o terceiro e o quarto prêmio \$27.917.561 (MLODINOW, 2009, p. 59).

De acordo com Mlodinow (2009), o prêmio para quem acertasse as 6 dezenas era de 27 milhões de dólares. Com isso, é possível concluir que um “preço justo” para o bilhete seria de: $\frac{27000000}{7059052}$, ou seja, o valor justo de um bilhete para a loteria de Virgínia seria aproximadamente 3,95 dólares. Porém, a loteria estipulou que o preço do bilhete seria de apenas 1 dólar, permitindo assim que com um investimento de 7.059.052 dólares fosse possível jogar todas as combinações possíveis. Como destacado, esse tipo de operação enfrenta dois problemas principais: a impossibilidade de uma pessoa fazer todos os jogos possíveis e a chance de várias pessoas acertarem as seis dezenas e o prêmio máximo ter que ser dividido.

- O primeiro problema foi resolvido reunindo um grande grupo de 2500 investidores, onde cada um deveria estar disposto a investir cerca de 3000 dólares. Para maximizar a quantidade de apostas feitas, eram realizadas tanto apostas presenciais quanto online;
- Em relação ao segundo problema, havia uma diferença importante em relação à Mega-Sena da virada. Esta não era uma aposta diferenciada, com um número maior de apostadores do que as outras (como é o caso da Mega-Sena da virada), e por isso não havia indícios que levassem ao grupo de investidores a crer que este concurso teria uma probabilidade maior de incidência de prêmios compartilhados do que os demais sorteios. Para embasar esta teoria, foi feito um levantamento dos últimos 170 sorteios, de onde o grupo pôde fazer uma análise acerca da quantidade de ganhadores em cada um deles.

Com isso, foi observado que em 170 concursos, em 120 não houve ganhador (eles ganhariam sozinhos), em 40 houve 1 ganhador (eles dividiriam o prêmio pela metade) e em apenas 10 houve dois ganhadores (dividiriam o prêmio em três partes). Se eles ganhassem sozinhos cada investidor ganharia cerca de $\frac{2700000}{2500} = 10800$ dólares (um lucro de \$ 7800,000). Caso eles ganharem o prêmio dividido com outra pessoa, o montante que o grupo iria ganhar seria de US\$ 5400,00, lucrando assim \$ 2400,00.

No cenário que o prêmio fosse dividido para três pessoas, cada investidor ganharia US\$ 3600,00, gerando um lucro de \$ 600,00. Apenas se o prêmio fosse dividido para 4 pessoas eles teriam prejuízo, pois cada investidor ganharia US\$ 2200,00, gerando assim um

prejuízo de US\$ 800,00 para cada um deles. Como essa possibilidade nunca havia acontecido nos últimos 170 sorteios, o grupo decidiu que a operação valia a pena. Assim, utilizando o conceito de esperança matemática, é possível inferir o lucro esperado de cada investidor da seguinte forma: $\frac{120}{170} \times 7800 + \frac{40}{170} \times 2400 + \frac{10}{170} \times 600$, gerando um lucro médio de US\$ 5500,00.

No fim das contas, por questões de erros de logística, os investidores acabaram por realizar apenas cerca de 5 milhões de apostas, o que poderia colocar em risco toda a operação. No entanto, a sorte estava do lado deles. A aposta vencedora estava com eles e, melhor ainda, era a única. Após a loteria descobrir o “golpe”, ela se recusou a pagar o prêmio, mas após alguns embates jurídicos os investidores ganharam a causa, recebendo assim o seu prêmio.

3.7 Problema dos pontos de Pascal: Como dividir uma aposta de um jogo que não acabou

Pascal, assim como a maioria dos matemáticos da época, iniciou seus estudos sobre a aleatoriedade por meio de jogos de azar. Seu interesse por jogos de azar (e subsequentemente pela teoria da probabilidade) foi impulsionado por sua saúde frágil. Ele sofria de fortes dores de barriga, dificuldade de engolir e manter a comida no estômago, além de fraqueza e dores de cabeça constantes. Após várias tentativas de tratamentos médicos sem sucesso, Pascal, aos 24 anos, desesperado, procurou uma nova junta de médicos que lhe propôs um tratamento inusitado: evitar todo trabalho mental contínuo e buscar ao máximo qualquer oportunidade de se distrair (MLODINOW, 2009).

Logo após isso, Pascal começou a passar seu tempo na companhia de outros jovens ricos que viviam de renda (o pai de Pascal morreu quando ele tinha 24 anos, deixando-lhe uma farta herança). Em uma dessas festas, Pascal conheceu Antoine Goumband, que tinha o título de nobreza de Chevalier de Mére. Este amigo propôs o seguinte problema:

Era o ano de 1654. A questão que De Mére levou a Pascal era conhecida como o problema dos pontos: suponha que você e outro jogador estejam participando de um jogo no qual ambos têm a mesma chance de vencer e o vencedor será o primeiro a atingir um certo número de pontos. Em um determinado momento, o jogo é interrompido quando um dos jogadores está na liderança. Qual é a maneira mais justa de dividir o dinheiro apostado? (MLODINOW, 2009, p.51).

A forma mais justa de dividir o dinheiro apostado é distribuir a premiação proporcionalmente à probabilidade que cada um deles tem de ganhar o jogo. Para realizar essa análise, é necessário observar o quanto falta para cada um dos jogadores ganhar e fazer o cálculo da probabilidade baseado nisso.

Neste exemplo, dois jogadores de xadrez estão jogando uma partida melhor de três, com um deles ganhando de 1-0. Dois apostadores apostaram 600 reais cada, cada um em um

dos jogadores diferentes. Caso o jogo seja interrompido neste momento, quanto cada apostador deverá receber?

A primeira observação que deve ser feita acerca desse problema é que, como não foi especificado, cada jogador tem a mesma chance de vencer uma partida em específico. Outro ponto a salientar é que, ao calcular a probabilidade, deve-se presumir que as 2 partidas finais foram realizadas, mesmo que o jogador que está ganhando de 1-0 vença a segunda partida. Feito isso, é necessário observar as combinações que dão a vitória a cada um dos jogadores. Para facilitar a nomenclatura, o jogador que está à frente será chamado de jogador 1 e o jogador que está perdendo será chamado de jogador 2.

Combinações que levem o jogador 1 à vitória:

- Ganhar os dois jogos;
- Ganhar o segundo e perder o último;
- Perder o segundo e ganhar o último.

Combinações que levam o jogador 2 a vitória:

- Ganhar os dois últimos jogos

A partir disso, observa-se que das 4 possibilidades, o jogador 1 ganha em 3 delas, enquanto o jogador 2 ganha em apenas uma delas. Logo o montante total da aposta (R\$600,00) deve ser dividido da seguinte forma: $\frac{3}{4}$ para o apostador que investiu seu dinheiro no jogador 1 e $\frac{1}{4}$ no jogador que investiu seu dinheiro no jogador 2, sendo assim eles iriam receber, respectivamente R\$ 450,00 e R\$ 150,00.

Outro problema ao qual se pode aplicar a ideia de Pascal é o seguinte: Em 1996, o Atlanta Braves venceu o New York Yankees nos dois primeiros jogos finais do campeonato de beisebol dos Estados Unidos, o qual é decidido através de uma melhor de 7 jogos (o primeiro time que ganhar 4 jogos é declarado campeão). Considerando este cenário de vitória dos Braves nos dois primeiros jogos, qual deveria ser um pagamento justo numa aposta feita no Yankees?

Para resolver este problema, será utilizado o mesmo princípio que foi utilizado no problema dos pontos de Pascal, ou seja, devemos inferir a probabilidade de cada time ser campeão e fazer a divisão da aposta de forma proporcional a esses valores.

Como o Atlanta venceu os dois primeiros jogos, ele será campeão se acontecer uma dessas possibilidades nos seus 5 jogos restantes:

- **O Atlanta vencer os 5 jogos restantes:**

Isso ocorre de $C_{5,5}$ maneiras diferentes, ou seja, temos apenas 1 possibilidade de isso ocorrer.

- **O Atlanta vencer 4 jogos e perder 1 dentre os jogos restantes:**

Isso ocorre de $C_{5,4}$ maneiras diferentes, ou seja, de 5 maneiras distintas

- **O Atlanta vencer 3 jogos e perder 2 jogos**

Isso ocorre de $C_{5,3}$ maneiras diferentes, ou seja, de 10 maneiras distintas

- **O Atlanta vencer 2 jogos e perder 3 jogos**

Isso ocorre de $C_{5,2}$ maneiras diferentes, ou seja, 10 casos distintos.

Somando as possibilidades, é possível observar que o Atlanta sairia campeão em 26 delas. Como são 5 jogos e cada um deles tem 2 resultados possíveis (vitória ou derrota do Atlanta), temos um total de $2^5 = 32$ possíveis resultados. Com isso, tem-se uma probabilidade de vitória de Atlanta de $\frac{26}{32} = 0,8125$, ou 81,25% isso quer dizer que se alguém apostou no Atlanta, essa pessoa deve ficar com 81,25% do montante, enquanto quem apostou no New York Yankees deve ficar apenas com 18,75%. Isso reflete a probabilidade de cada time ser campeão, com base no cenário em que o Atlanta Braves venceu os dois primeiros jogos.

Neste tipo de problema, ambos os times têm a mesma chance de ganhar um jogo em específico. Suponha agora um problema bem similar a este, com apenas uma alteração: o time de Atlanta tem 40% de chance de ganhar um jogo em específico, enquanto o time de Nova Iorque tem 60% de ganhar um jogo em específico. Este modelo é mais próximo da realidade, pois ele leva em conta não apenas as possibilidades para os jogos restantes, mas também as forças dos times. Pretende-se analisar como esta mudança afetará a probabilidade final.

A cada vitória do Atlanta deverá ser atribuído 0,4 (probabilidade de 40%), enquanto a cada derrota do Atlanta 0,6 (probabilidade de 60%) e então aplicando a probabilidade do produto, é possível definir a probabilidade de cada time se sagrar campeão.

- **5 vitórias do Atlanta**

- $0,4 \times 0,4 \times 0,4 \times 0,4 \times 0,4 = 0,01024 = 1,024\%$

Vimos no exemplo anterior que essa combinação só ocorre em 1 caso. Com isso, observa-se que no exemplo anterior existe apenas uma combinação para este caso, resultando em uma probabilidade final de 1,024%.

- **4 Vitórias do Atlanta**

- $0,4 \times 0,4 \times 0,4 \times 0,4 \times 0,6 = 0,01536 = 1,536\%$

Observou-se que há 5 combinações possíveis para este resultado. Portanto, o resultado deve ser multiplicado por 5, resultando em 7,68% .

- **3 Vitórias do Atlanta**

- $0,4 \times 0,4 \times 0,4 \times 0,6 \times 0,6 = 0,02304 = 2,304\%$

Foi observado que existem 10 possibilidades para este caso. Portanto, o total resulta em 23,04%.

- **2 Vitórias do Atlanta**

- $0,4 \times 0,4 \times 0,6 \times 0,6 \times 0,6 = 0,03456 = 3,456\%$

Nesse caso, também haverá 10 possibilidades, resultando em um total de 34,56%. Ao somar essas probabilidades, obtém-se uma probabilidade de vitória para o time de Atlanta de 66,30%, a qual é substancialmente inferior à do caso anterior.

4. PROBABILIDADE E EXAMES CLÍNICOS

O presente capítulo explora a interseção entre probabilidade e exames clínicos, destacando a importância dos conceitos probabilísticos no contexto da medicina e da saúde. Ao longo deste capítulo, serão abordadas questões fundamentais, como a interpretação de

resultados de testes diagnósticos, a avaliação da precisão e acurácia desses exames, bem como o impacto da probabilidade na tomada de decisões médicas. Com isso, compreender que a aplicação da probabilidade em exames clínicos é essencial para uma prática médica informada e eficaz, permitindo uma análise mais precisa e uma tomada de decisão fundamentada em evidências.

Ao longo deste capítulo, serão exploradas as relações entre o conceito de probabilidade e os exames clínicos, examinando tanto os aspectos médicos, como os exames para detecção de doenças, quanto os relacionados aos testes clínicos para controle de dopagem, analisando a interpretação dos dados desses exames tanto por especialistas da área quanto por indivíduos submetidos a eles.

Um desafio significativo na interpretação dos testes clínicos é a compreensão limitada da população em geral sobre o "falso positivo". Embora as taxas de falso positivo divulgadas pelos laboratórios sejam geralmente baixas, isso pode levar à crença de que esses testes são infalíveis. O falso positivo representa a proporção de pessoas que não têm a doença (ou substância proibida, no caso de exames antidoping), mas que o teste acusa como positivo (GIGERENZER, 2002).

Entretanto, para avaliar verdadeiramente a eficácia de um teste ou exame, é crucial considerar a pergunta reversa: dado um teste positivo, qual é a probabilidade de a pessoa realmente ter a doença ou ter usado a substância ilícita? Essa probabilidade é chamada de probabilidade a posteriori. Para essa análise, serão definidas duas características fundamentais dos exames clínicos: sensibilidade e especificidade, com base na terminologia médica, mas aplicáveis também aos testes de controle de dopagem. Esses conceitos serão delineados com base no artigo "Sensibilidade e Especificidade" (BASTOS, 2004).

4.1. Sensibilidade

A sensibilidade de um teste diagnóstico refere-se à sua capacidade de detectar corretamente os indivíduos que realmente possuem a condição em questão. Por exemplo, se em um grupo de 100 pessoas doentes, o teste consegue identificar 70 delas como positivas, então sua sensibilidade é de 70%. Esse índice é fundamental para avaliar a eficácia de um exame na detecção de uma doença específica. No entanto, é importante ressaltar que a sensibilidade não representa a taxa de acurácia do teste, e sim a sua capacidade de identificar corretamente os casos positivos (CRISTIANO, 2017).

Por outro lado, a probabilidade complementar à sensibilidade é conhecida como falso negativo. Isso se refere à proporção de pessoas que não têm a doença, mas que testam positivo no exame. Em outras palavras, são os casos em que o teste indica erroneamente a presença da condição, resultando em um diagnóstico incorreto. Essa métrica é crucial para entender a especificidade do teste, ou seja, sua capacidade de identificar corretamente os casos negativos. Sendo assim, a sensibilidade é uma medida essencial para avaliar a confiabilidade de um teste diagnóstico, pois indica sua habilidade em detectar adequadamente os casos positivos. Compreender essa métrica, juntamente com a probabilidade de falso positivo, é fundamental para uma interpretação precisa dos resultados dos exames clínicos e para uma tomada de decisão informada em relação ao tratamento e monitoramento da saúde dos pacientes (GIGERENZER, 2002)

4.2 Especificidade

A especificidade de um teste diagnóstico representa sua capacidade de identificar corretamente os verdadeiros negativos, ou seja, a proporção de pacientes que não possuem a doença e que o teste corretamente detecta como negativos. Por exemplo, se em um grupo de 100 pessoas sem a doença, o teste indica negativo para 70 delas, então sua especificidade é de 70%. Essa métrica é crucial para avaliar a precisão do teste em descartar corretamente os casos negativos, contribuindo para uma interpretação confiável dos resultados (GIGERENZER, 2002).

Por outro lado, a probabilidade complementar à especificidade é conhecida como falso positivo. Isso se refere à proporção de pacientes que não possuem a doença, mas que testam positivo no exame. Em outras palavras, são os casos em que o teste detecta incorretamente a presença da condição, resultando em um diagnóstico positivo incorreto. Compreender a especificidade, juntamente com a sensibilidade do teste, é essencial para uma avaliação abrangente de sua eficácia na detecção e exclusão de uma determinada condição médica.

A Tabela 6, adaptada do “*Calculated Risks: How to know When numbers deceive you*”, sintetizam de forma simples estes conceitos de sensibilidade e especificidade:

Tabela 6 – Relação entre o resultado dos exames, a sensibilidade e a especificidade.

Resultado do Teste	doença	
	Sim	Não
Positivo	sensibilidade	taxa de falso positivo
Negativo	taxa de falso negativo	especificidade

Fonte: GIGERENZER, 2002.

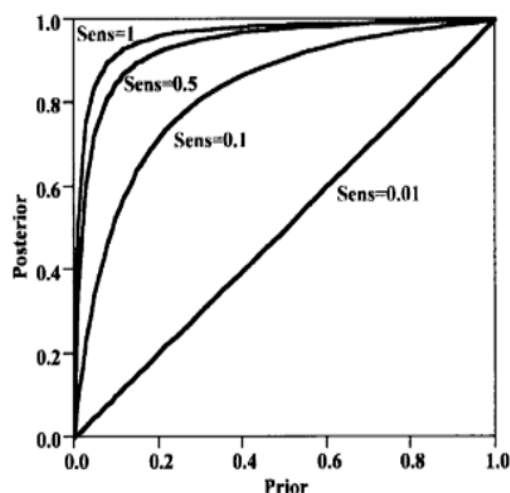
Outra informação crucial para uma análise completa dos resultados dos testes é a taxa de incidência da doença na população, conhecida como probabilidade a priori. Esse índice, juntamente com os valores da especificidade e sensibilidade, permite medir a eficácia do teste, ou seja, a probabilidade de um resultado positivo indicar corretamente a presença da doença, denominada *Positive Predictive Value (PPV)*.

Com isso, é importante ressaltar que a especificidade e a sensibilidade são características intrínsecas de cada teste e não dependem da probabilidade a priori. No entanto, em populações com maior probabilidade de incidência da doença, mais indivíduos serão identificados pelo teste.

Conforme Bastos (2004), os exames clínicos raramente conseguem ter alta especificidade e sensibilidade simultaneamente. Geralmente, é necessário fazer uma escolha sobre qual atributo é mais importante em um determinado contexto. Por exemplo, é crucial que testes diagnósticos tenham alta especificidade, pois se deseja minimizar ao máximo os falsos positivos, uma vez que um diagnóstico positivo para certas doenças implica em tratamentos muitas vezes dolorosos e traumáticos. Por outro lado, em ambientes de atendimento primário, é preferível usar testes com alta sensibilidade, pois é mais importante não deixar de detectar doenças, mesmo que isso resulte em falsos positivos. Nesses casos, os pacientes serão encaminhados para especialistas que poderão identificar e resolver esses falsos positivos posteriormente.

A Figura 3, ilustra a relação entre as probabilidades a priori e a posteriori (ou PPV), para diferentes valores de sensibilidade. Em todos os casos, a especificidade foi fixada em 99%.

Figura 3 – Relação da sensibilidade e especificidade com o PPV



Fonte: GIGERENZER, 2002.

A partir da análise do gráfico, é possível inferir que se a taxa de probabilidade a priori for muito grande a eficácia do exame também será, pois em uma população com alta taxa de incidência de uma doença, é mais provável que um exame positivo realmente corresponda a verdade. De forma análoga, um aumento da sensibilidade do teste, também acarretará um aumento do PPV, visto que este aumento irá ocasionar um número maior de pacientes doentes detectados pelo teste

4.3 Diagnósticos Médicos: Será que são realmente confiáveis?

O imaginário comum sobre resultados de exames clínicos é geralmente de que eles são infalíveis. Quando um paciente recebe um diagnóstico médico, raramente reflete sobre a confiabilidade do exame realizado; a reação instintiva é aceitar o diagnóstico e iniciar o tratamento sugerido pelo médico (CRISTIANO, 2017).

De acordo com Gigerenzer (2002), isso ocorre porque os pacientes não são informados sobre a real probabilidade de um resultado positivo indicar que eles realmente têm a doença. Além disso, os próprios médicos muitas vezes não conseguem interpretar as informações de um exame e relacioná-las com a incidência da doença na população. Para ilustrar isso, Gigerenzer conduziu uma pesquisa com 24 médicos de duas grandes cidades da Alemanha, com uma média de 14 anos de experiência profissional. Esses médicos foram questionados da seguinte forma:

A probabilidade de uma mulher desenvolver câncer de mama é de 0,8%. Se ela tiver câncer de mama, a chance de a mamografia indicar um resultado positivo é de 90%. Caso ela não tenha câncer de mama, existe uma probabilidade de 7% de que ela ainda obtenha um teste positivo. Imagine uma mulher que tenha dado positivo em uma mamografia, qual é a probabilidade de ela realmente ter câncer de mama? (GIGERENZER, 2002, p. 51).

Para resolver esse tipo de problema, o Teorema de Bayes pode ser utilizado. Este teorema indica a probabilidade de um evento A ocorrer, dado que outro evento B já ocorreu, o que é exatamente o que se busca neste caso, pois deseja-se encontrar a probabilidade de um paciente realmente estar doente, dado que o resultado do exame foi positivo.

Segundo o teorema, a probabilidade do evento A ocorrer, dado que o evento B ocorreu, é dada por:

$$P(A|B) = \frac{P(B|A) \cdot P(A)}{P(B)}$$

Passando para a notação do exame clínico, tem-se os seguintes aspectos:

$P(A)$ –proporção de incidência de doença na população, ou seja, a probabilidade a priori.

$P(B)$ - probabilidade de o resultado do exame dar positivo. Para calcularmos esta probabilidade, deve-se considerar dois casos: o primeiro é o caso de a pessoa estar doente e o teste dar positivo (sensibilidade) e o segundo caso é o da pessoa não ter a doença e mesmo assim o teste indicar um resultado positivo (1- especificidade)

$P(B/A)$ - probabilidade de o resultado do exame ser positivo, caso o paciente realmente tenha a doença, ou seja, a sensibilidade do exame clínico.

$P(A/B)$ - probabilidade de o paciente estar realmente doente, dado que o resultado do exame foi positivo. Esse resultado é o que se chama de *PPV* e ele é o indicador da eficácia do exame médico.

Com isso, aplicando o teorema de Bayes, calcula-se o *PPV* da seguinte maneira:

$$PPV = \frac{sensxprior}{sensxprior + (1 - prior) \times (1 - esp)}$$

Logo, utilizando os dados disponíveis, é possível concluir que a probabilidade de uma pessoa ter câncer de mama, dado que um resultado foi positivo é:

$$PPV = \frac{0,9 \times 0,08}{0,08 \times 0,9 + (0,992) \times (0,7)} \approx 9,4 \%$$

O problema com essa abordagem é sua natureza técnica, que exige um conhecimento específico sobre probabilidades, levando a interpretações equivocadas por parte dos médicos, conforme observado por Gigerenzer (2002). No experimento conduzido por ele, não houve consenso entre os médicos sobre a probabilidade real de o paciente estar doente, com estimativas variando de 1% a 90%. Um terço dos médicos interpretou erroneamente a sensibilidade do teste como sendo o Valor Preditivo Positivo (VPP) do teste. Em outras

palavras, eles entenderam a sensibilidade como a porcentagem de casos em que o teste detecta a doença, quando na verdade a pergunta a ser respondida era inversa: se o paciente testou positivo, qual é a probabilidade de ele ter a doença?

Dos 16 médicos restantes, 8 estimaram probabilidades entre 50% e 80%, enquanto os outros 8 estimaram probabilidades inferiores a 10%, com metade deles estimando uma probabilidade de 1%. Apenas dois médicos estimaram corretamente a probabilidade, que é cerca de 9,4% (GIGERENZER, 2002).

Gigerenzer (2002) sugere que uma possível explicação para esse erro de análise por parte dos médicos é a forma como os dados sobre a incidência da doença, a sensibilidade e a especificidade do exame foram apresentadas. O autor argumenta que o cérebro humano tem dificuldades em assimilar informações expressas em termos de probabilidade, o que pode explicar a superestimação da probabilidade correta de o paciente realmente ter a doença.

Para melhorar a interpretação dos dados, ele propõe que as informações sejam apresentadas na forma de frequência. Para testar essa hipótese, ele conduziu um novo experimento no qual as informações foram apresentadas em valores absolutos em vez de probabilidades, trabalhando com um universo de 1000 mulheres e apresentando as porcentagens como quantidades de pacientes, permitindo assim que os médicos analisassem a situação de forma mais direta (GIGERENZER, 2002).

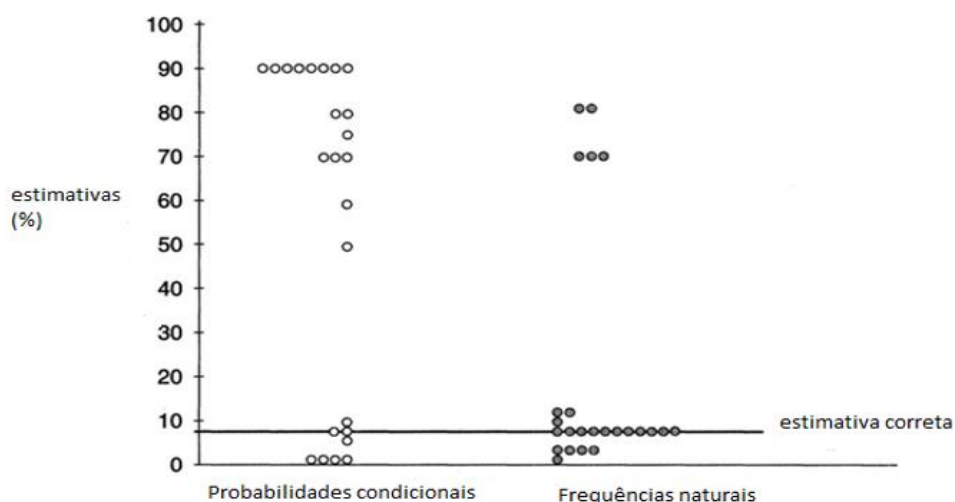
Oito em cada 1000 mulheres têm câncer de mama. Destas 8, 7 são detectadas pelo exame. Das 992 mulheres que não possuem câncer de mama, 70 são detectadas como positivas pela mamografia. Com base nisso, qual é a chance de uma paciente que tenha testado positivo, realmente tenha a doença? (GIGERENZER, 2002, p. 61).

Apesar de ser exatamente a mesma pergunta, a forma com que foi apresentada a questão deixa ela muito mais simples. Aqui se consegue facilmente ver que das 77 pessoas que testaram positivo (70+7) apenas 7 delas realmente estavam doentes. Com isso, se pode tranquilamente inferir que a probabilidade desejada é dada por $\frac{7}{77} \cong 9,1\%$. Esta forma de apresentar os dados, embora seja em essência a mesma ideia contida na solução utilizando a Regra de Bayes (razão entre os casos positivos que estavam realmente doentes e todos os casos positivos), ela é bem mais simples para a compreensão tanto dos médicos como dos pacientes.

A Figura 4 ilustra de forma clara como a apresentação dos dados é crucial para uma interpretação precisa da eficácia de um teste clínico. Ela demonstra como uma mudança na forma de apresentar os dados pode significativamente aumentar a precisão do diagnóstico e da comunicação do mesmo ao paciente, resultando em um maior número de médicos capazes de

avaliar corretamente a eficácia do exame médico. Sendo assim, para Gigerenzer (2002), em vez de atribuir a falta de habilidade dos médicos em calcular probabilidades, é necessário reformular a maneira como os dados são comunicados a eles.

Figura 4 – Estimativas das probabilidades dadas pelos médicos de acordo com a forma que os dados foram colocados



Fonte: GIGERENZER, 2002.

Outro aspecto que pode ser analisado é como esses resultados positivos se comportam ao longo de repetidos exames. Este estudo é extremamente relevante, pois a mamografia é feita, geralmente anualmente. Para calcular a probabilidade de uma mulher que não possui a doença receber pelo menos um diagnóstico positivo ao longo dos anos, é possível utilizar a probabilidade complementar. Dado que há uma probabilidade de 7% de obter um resultado positivo em pacientes sem a doença, isso implica que há uma probabilidade de 93% de receber um resultado negativo corretamente a cada ano.

Ao considerar que a mulher realiza a mamografia anualmente ao longo de vários anos, a probabilidade de receber um resultado negativo corretamente em todos os anos é de 0.93 elevado ao número de anos. Para obter a probabilidade de receber pelo menos um resultado positivo ao longo dos anos, se pode subtrair essa probabilidade de 1 (a probabilidade

complementar), uma vez que queremos encontrar a probabilidade de pelo menos um evento positivo ocorrer.

Aqui está a tabela mostrando a probabilidade de receber pelo menos um diagnóstico positivo erroneamente ao longo dos anos:

Tabela 7 – Probabilidade de o exame acusar um falso positivo ao longo dos anos

Quantidade de Exames	Probabilidade de pelo menos um falso positivo
1	$1 - (0,93) = 0,07 = 7\%$
2	$1 - (0,93)^2 = 1 - 0,86 = 0,14 = 14\%$
5	$1 - (0,93)^5 = 1 - 0,69 = 0,31 = 31\%$
10	$1 - (0,93)^{10} = 1 - 0,48 = 0,52 = 52\%$
15	$1 - (0,93)^{15} = 1 - 0,33 = 0,67 = 67\%$

Fonte: O autor

A partir desta tabela, vê-se que com apenas 10 anos de mamografia, mais da metade das mulheres sadias já passaram por um diagnóstico de falso positivo. Ao incrementar um pouco mais para 15 anos, $\frac{2}{3}$ das mulheres já positivaram o exame.

Com isso, é importante salientar que essas probabilidades são baseadas em exames de rotina, feitos em mulheres aparentemente sadias, ou seja, que não possuam nenhuma evidência ou sintoma da doença. Caso já exista algum outro indicativo prévio da presença da doença, o resultado positivo da mamografia será muito mais relevante em termos de confiabilidade da presença da doença.

4.3.1 Mlodinow e o teste de HIV

No ano de 1989, Leonard Mlodinow fez uma ligação para o seu médico depois de ter seu seguro de vida recusado. O médico afirmou categoricamente que ele havia testado positivo para HIV, informando que havia uma probabilidade de 99,9% de chance de ele estar infectado. No entanto, Mlodinow acabou sendo um caso de falso positivo. Esse tipo de diagnóstico errôneo é mais comum do que se imagina. Os fragmentos a seguir são trechos transcritos do livro *“Calculated Risks: How to Know When Numbers Deceive You”*, contendo depoimentos de pacientes que testaram positivo para o HIV e depois descobriram que o resultado estava equivocado.

- **Caso David**

O Chicago Tribune publicou uma carta em 5 de março de 1993, intitulada “Um falso teste de HIV causou 18 meses de inferno”, acompanhada de uma resposta. O remetente, David, descreveu sua experiência angustiante após receber um resultado positivo para HIV em um teste de rotina realizado em março de 1991. A notícia devastadora o levou a um estado de profunda depressão e considerações suicidas, mas com o apoio de sua família e amigos, decidiu reagir.

Após se mudar para a Califórnia em busca de tratamento especializado, descobriu, surpreendentemente, que novos testes realizados por um médico local apresentavam resultados negativos, contradizendo o diagnóstico inicial. David expressou gratidão por sua saúde, mas salientou os 18 meses de sofrimento causados pela falsa identificação do vírus. Ele instou os médicos a serem mais diligentes e encorajou os leitores a buscarem segundas opiniões médicas, enquanto ele mesmo planejava continuar realizando testes de HIV a cada seis meses, embora sem o mesmo pânico anterior.

- **O pesadelo de Susan**

Durante uma consulta médica de rotina em um hospital da Virgínia, em meados da década de 1990, Susan, uma mãe solteira de 26 anos, foi submetida a um exame para o HIV. Embora não usasse drogas por via intravenosa e não se considerasse em risco de contrair o vírus, o teste realizado algumas semanas depois resultou positivo, o que na época equivalia a um diagnóstico terminal. A notícia abalou Susan profundamente, causando perturbação e choque. A divulgação de seu diagnóstico provocou o isolamento social, levando colegas a

evitá-la e resultando na perda de seu emprego. Eventualmente, ela se mudou para uma casa de recuperação para pacientes infectados pelo HIV, onde, em meio a esse ambiente, tomou decisões arriscadas, incluindo o envolvimento em relações sexuais desprotegidas.

Preocupada com a saúde de seu filho de 7 anos, Susan adotou medidas extremas de precaução, como evitar beijá-lo e manipular sua comida, embora isso tenha causado sofrimento emocional. Após meses vivendo sob o diagnóstico terminal, Susan desenvolveu bronquite, levando um médico a solicitar um novo teste de HIV. Apesar de sua descrença inicial, o teste resultou negativo. Posteriormente, ao investigar, descobriu-se que o resultado original do teste de sangue de Susan foi trocado inadvertidamente com o de um paciente HIV positivo no momento em que os dados foram inseridos no computador do hospital.

O erro não apenas mergulhou Susan em um falso desespero, mas também deu falsas esperanças à outra paciente. A falta de informação sobre a possibilidade de resultados falsos positivos nos testes de HIV deixou Susan vulnerável, levando-a a viver nove meses sob a sombra de um diagnóstico incorreto e terminal. Como resultado, ela processou seus médicos e obteve um acordo generoso, que incluiu a compra de uma casa. Além disso, Susan abandonou o uso de drogas e experimentou uma transformação espiritual após esse pesadelo que mudou sua vida.

Estes dois relatos destacam a importância crucial de os profissionais de saúde serem mais cautelosos ao comunicar resultados positivos para pacientes, especialmente em casos de doenças graves como câncer e HIV/AIDS, onde há estigma e desinformação sobre a forma de contágio. Como evidenciado nas cartas, diagnósticos equivocados podem acarretar sérios transtornos psicológicos, levando alguns pacientes à depressão e até mesmo ao suicídio. Essas consequências poderiam ser evitadas se os pacientes fossem informados sobre as reais chances de um teste estar incorreto.

Com isso, é essencial que os médicos tenham um entendimento mais claro da verdadeira taxa de eficácia dos exames, como observado na questão do exame de câncer de mama. Isso requer uma compreensão precisa da incidência da doença no grupo específico ao qual o paciente pertence. No caso de Leonard Mlodinow, por exemplo, o médico cometeu um erro ao afirmar que o exame tinha uma chance de 99,9% de estar correto, sem considerar a taxa de incidência da doença no grupo demográfico de Mlodinow. Considerando esses dados, a probabilidade real de um resultado positivo indicar que o paciente está infectado é muito menor do que o indicado pelo médico. Com isso, como a incidência da doença no grupo ao qual Mlodinow está inserido (homens brancos, heterossexuais e não usuários de drogas), era

de $\frac{1}{10000}$, e utilizando a taxa de falso positiva informada pelo médico (0,1%), pode-se inferir que em um grupo de 10000 pessoas, 10 testariam positivo sem ter a doença, ou seja neste mesmo grupo, 11 pessoas testariam positivo, das quais apenas uma teria a doença. Logo, a taxa de eficácia do exame é dada por $\frac{1}{11} \approx 9\%$, bem menor que a informada pelo médico.

É importante ressaltar que essas probabilidades variam de acordo com o grupo estudado. Por exemplo, de acordo com Gigerenzer (2002), a taxa de incidência de HIV em homens homossexuais nos Estados Unidos é de 1,5%. Portanto, ao analisar o caso de Mlodinow utilizando os dados desse grupo demográfico específico, é possível calcular a probabilidade correta de um resultado positivo indicar uma infecção real. Com isso, a probabilidade de um exame positivo realmente indicar a doença é de $\frac{150}{160} \approx 94\%$.

4.4 Exames Antidopings

No contexto esportivo, os exames de controle de dopagem têm dois objetivos principais. Primeiramente, visam verificar se os atletas estão fazendo uso de substâncias que possam melhorar sua performance, como anabolizantes e hormônios. Isso permite que as organizações antidoping mantenham a integridade e a equidade nas competições esportivas. Em segundo lugar, esses exames têm o propósito de coibir o uso de drogas ilícitas, como maconha, cocaína e heroína, que são prejudiciais à saúde e contradizem a ideia de uma vida saudável associada à prática esportiva.

O processo de exame consiste na coleta de urina do atleta, escolhido aleatoriamente por sorteio. A urina é dividida em dois frascos, sendo o primeiro considerado para análise inicial. Em caso de resultado positivo, o atleta tem o direito de solicitar a análise do segundo frasco, conhecido como contraprova, para confirmação ou não do resultado de doping. É fundamental destacar que todas as características inerentes aos testes clínicos também se aplicam aos exames *antidopings*.

Para analisar a eficácia do exame, é necessário considerar não apenas a taxa de falso positivo, mas também outros fatores. Por exemplo, ao investigar o caso da atleta Mary Slaney, campeã mundial em 1983 nos 1.500m e 3.000m, pega no exame antidoping por uso de Testosterona ao tentar se classificar para as Olimpíadas de Atlanta em 1996, é preciso calcular a probabilidade de inocência ou não da atleta.

Com base nos dados do artigo “*Inference About Testosterone Abuse*”, que considera um exame antidoping com sensibilidade de 50% e especificidade de 99%, além de uma taxa

de incidência de doping entre atletas de 10%, pode-se calcular o Valor Preditivo Positivo (PPV). Em um grupo de 1000 atletas, 10% estariam utilizando doping, dos quais o exame detectaria apenas 50%. Por outro lado, dos 900 atletas inocentes, 1% seria erroneamente acusado. Essas informações permitem uma análise mais precisa da situação da atleta e da eficácia do exame antidoping.

Com isso, entre um grupo de 1000 atletas, 59 seriam pegos no exame antidoping. Desses 59, 50 são realmente culpados, mas 9 são inocentes. Portanto, a probabilidade de um atleta pego no exame antidoping ser realmente culpado é de $\frac{50}{59}$, o que resulta em aproximadamente 84,7%, o que apesar de ser uma taxa alta, não tem a mesma sensação de justiça da probabilidade de 99%. Com isso, ao analisar milhares de atletas, inevitavelmente, a chance de algum deles ser inocente é relativamente alta.

Este resultado apresenta uma taxa de inocência de 15,3%. É claro, que se ao analisar-se a contraprova, e o resultado do teste continuar dando positivo, a probabilidade de inocência cai drasticamente. Para calcular esta taxa, basta aplicar a probabilidade do produto:

$$\frac{15,3}{100} \times \frac{15,3}{100} = \frac{234,09}{10000} = 2,34\%.$$

A partir disso, concluímos que o *PPV* para este caso é de 97,66%. Esse percentual, embora menor do que os 99% que muitos imaginam ser a eficácia do exame, pode ser considerado uma margem razoável de acerto.

Outro aspecto a ser considerado no caso de Stanley é a métrica utilizada nos exames antidopings para a medição dos níveis de testosterona, o que complica ainda mais a análise dos resultados e questiona sua equidade. Segundo Berry e Chastain (2004), não há consenso sobre os níveis aceitáveis de testosterona devido à grande variação na concentração de testosterona na urina, especialmente entre as mulheres, devido a fatores como menstruação, gravidez e uso de pílulas anticoncepcionais.

Apesar disso, desde as Olimpíadas de 1980, foi estabelecido que a concentração aceitável de testosterona (razão entre o nível de testosterona e epitestosterona) seria $T/E \leq 6,0$. No entanto, de acordo com Catlin, citado por Berry e Chastain, o nível aceitável seria $T/E \leq 15,0$, o que resultaria, segundo esse critério, na absolvição de Stanley, pois a proporção de testosterona que levou à sua suspensão era de 9,5 e 11,6 na prova e na contraprova, respectivamente.

Neste caso específico, em que não há consenso na comunidade científica, a escolha do parâmetro permitido de testosterona é feita de forma arbitrária. No final do século passado, o

atletismo enfrentava uma crise grave de doping, o que exigia a adoção de “padrões rígidos” nos testes antidoping, mesmo que isso levasse à condenação de vários atletas inocentes. Por outro lado, se a proporção proposta por Catlin fosse adotada, muitos casos de uso de substâncias proibidas passariam despercebidos, resultando em falsos negativos. Portanto, é necessário encontrar um valor intermediário que equilibre a sensibilidade e especificidade do teste, garantindo sua confiabilidade.

5. O DNA E O DIREITO

Neste capítulo, será dado seguimento ao estudo sobre testes clínicos, com foco no teste de DNA aplicado na área do direito. Segundo Bonnaccorso (2005), a utilização do exame de DNA em processos jurídicos foi de extrema importância para a medicina forense³. Com a descoberta desse método, os legistas tornaram-se capazes de elucidar casos judiciais com maior clareza, analisando o material biológico contido no DNA. Ao contrário dos testes de sorologia, que não possuíam o poder discriminatório necessário para esclarecer crimes em que havia um grande número de suspeitos, o teste de DNA proporcionou uma ferramenta mais eficaz para esse fim.

5.1 Introdução ao conceito de DNA

A descoberta dos loci hiper variáveis por Alec Jeffreys em 1984 marcou o início de uma revolução genética, que resultou na criação dos primeiros testes de DNA, um ano mais tarde, em 1985. Esses loci hiper variáveis, também conhecidos como minissatélites, foram fundamentais para o progresso da medicina forense, uma vez que produziam uma espécie de impressão digital do DNA (BONNACCORSO, 2005).

Para continuar, algumas definições básicas sobre a estrutura do DNA serão fornecidas.

Genoma: Todas as diferentes longas cadeias de DNA em cada uma de suas células – é único para quase todos (a única exceção são gêmeos idênticos, trigêmeos...).

³ A medicina forense é o ramo da medicina que atrela conhecimentos médicos e jurídicos a fim de elucidar casos criminais.

Repetições em Tandem: São padrões de repetição sequencial de bases nitrogenadas. São utilizadas como marcadores moleculares em análises de parentesco e em análises forenses.

Essas bases nitrogenadas podem ser de 2 tipos:

- Purinas, que englobam a adenina (A) e a guanina (G)
- Piridinas, que englobam a citosina (C) e timina (T)

Exemplo de repetição em tandem: GAG-GAG-GAG-GAG.

Um perfil de DNA é simplesmente uma contagem dessas pequenas repetições em tandem em locais específicos do genoma. A partir da análise dessas repetições, é feita uma comparação entre o material genético proveniente da vítima e dos possíveis suspeitos. A Figura 5 ilustra como é a estrutura do DNA.

Figura 5 – Representação da estrutura do DNA



Fonte: <https://br.freepik.com/fotos/dna>

Com isso, faz-se necessário esclarecer que o teste de DNA, apesar de ser uma poderosa ferramenta, não é infalível, como muitas pessoas, incluindo alguns especialistas da área criminal, acreditam. De acordo com Koehler (1993), esta crença se dá pelo fato conhecido que “Duas pessoas não podem ter o mesmo DNA, exceto se eles forem gêmeos idênticos”. Este fato é refutado através de dois motivos principais:

- Primeiramente, ao considerar aproximadamente 13 locais diferentes no genoma em testes de DNA altamente precisos, é importante notar que não é logicamente impossível que

duas pessoas não relacionadas compartilhem o mesmo DNA por mero acaso. Embora essa possibilidade seja rara, ainda é possível, especialmente em testes de DNA que examinam menos regiões.

- A taxa de falso positivo também é uma preocupação, abrangendo erros humanos durante a realização do teste de DNA, incluindo erros na manipulação das amostras e trocas de amostras.

Para obter uma análise mais abrangente da eficácia do DNA em tais casos, Koehler (1993) propõe uma sequência lógica na qual a crença na culpa do réu aumenta a cada etapa avançada do processo de análise do DNA.

Seguindo essa sequência lógica, para compreender melhor o teste de DNA, é necessário considerar as seguintes perguntas:

1. A correspondência de DNA reportada é verdadeira?
2. Se a correspondência for verdadeira, o réu é realmente a fonte desse DNA?
3. Se o réu for a fonte do DNA, será que ele é realmente o autor do crime?

A análise de cada aspecto relacionado a essas perguntas visa destacar os erros analíticos presentes em cada uma delas.

5.2 Uma correspondência relatada é realmente verdadeira?

Segundo Koehler (1993), uma correspondência de DNA é considerada verdadeira se as características identificadas nos vestígios e nas amostras dos suspeitos coincidirem de fato. Embora alguns especialistas acreditem que seja impossível ocorrer uma correspondência de DNA falsa, erros podem ocorrer devido a falhas no processo de análise do DNA. Esses erros podem ser tanto técnicos, como falhas enzimáticas ou concentrações anormais de sal, quanto humanos, como contaminação da amostra ou erro na identificação dos testes de DNA.

Para avaliar a confiabilidade de um teste de DNA, o ideal seria realizar testes de proficiência “às cegas”, nos quais os laboratórios não saberiam que estão sendo testados e as condições seriam similares às encontradas na realidade. Entretanto, os testes são geralmente agendados com antecedência, o que pode distorcer os resultados, pois os técnicos tendem a ser mais cautelosos na análise do DNA. Além disso, o material genético utilizado para esses testes é selecionado para estar preservado, o que nem sempre reflete a realidade (KOEHLER, 1993).

Apesar dos problemas mencionados, Koehler (1993) relata que testes de proficiência revelaram uma incidência de erro maior do que muitas pessoas imaginam. Em estudos

realizados em 1987 e 1988, foram detectados diversos erros. Em um dos estudos, 50 amostras foram enviadas a 3 laboratórios, que juntos detectaram 113 correspondências de DNA, das quais 2 eram falsas, resultando em uma taxa de erro de 1,7%. Em outro estudo mais recente, 38 laboratórios relataram 75 correspondências de DNA, das quais 3 eram falsas, correspondendo a uma taxa de erro de 4%. Esses estudos desafiam a crença na infalibilidade dos testes de DNA defendida por alguns especialistas jurídicos.

Essas informações geralmente são omitidas nos tribunais, que consideram a probabilidade de existir duas amostras com o mesmo DNA, mas não a probabilidade de erro técnico ou humano, levando os juízes a desprezarem a possibilidade de um falso positivo ao analisar a culpabilidade ou inocência do réu. Segundo o autor, os promotores muitas vezes confundem a probabilidade de correspondência aleatória⁴ com a probabilidade de o réu ser realmente culpado. A análise da culpabilidade requer dados adicionais para ser mais precisa (KOEHLER, 1993).

Ao considerar a probabilidade de correspondência aleatória como a probabilidade de culpa do réu, negligenciamos outros indícios que existiam antes da realização do exame de DNA, além da taxa de falso positivo. Isso implica deixar de lado o conceito de "probabilidade a priori", que se refere à convicção de que o suspeito era culpado antes do exame de DNA ser realizado. Embora determinar esse valor às vezes seja complexo, ele não pode ser ignorado ao avaliar o peso do resultado positivo de um exame de DNA. Dentre os fatores que podem influenciar essa probabilidade, estão a quantidade de suspeitos do crime e outras evidências colhidas durante o processo criminal, como testemunhas oculares e informações obtidas por meio de interrogatórios.

A Tabela 8 analisa como a probabilidade a posteriori é calculada com base em estimativas da probabilidade a priori, da taxa de falso positivo e da probabilidade de correspondência aleatória. Nessa tabela, Thompson (2003) estabelece a relação entre a probabilidade a posteriori e a probabilidade *a priori*, a probabilidade de correspondência aleatória e a taxa de falso positivo. Antes de apresentar a Tabela 8, uma breve exposição é feita sobre como Thompson estimou essas grandezas.

Para classificar as probabilidades *a priori*, Thompson utilizou como referência as seguintes informações:

⁴ Uma correspondência aleatória ocorre quando duas pessoas possuem o mesmo perfil de DNA, ou seja, eles são coincidentes em cada alelo que o exame investigado pelo exame de DNA.

- **Probabilidades a priori 2:1** – as outras evidências são fortes, mas não suficientes para a condenação do réu. Estatisticamente, um terço dos suspeitos de estupro são excluídos por exames de DNA. Com isso, essa probabilidade a priori de 2:1 reflete um caso típico deste tipo de crime.
- **Probabilidades a priori 1:10 e 1:100** – é utilizada quando as outras evidências indicam uma baixa probabilidade que o suspeito é realmente a fonte do DNA. Geralmente isso ocorre quando a “match” ocorre em um “*DNA Dragnet*” (neste caso a polícia testa pessoas próximas a cena do crime apenas pela proximidade sem nenhuma suspeita da participação delas no crime).
- **Probabilidades a priori 1:1000** – não possui praticamente nenhuma evidência a não ser o teste de DNA. Este caso pode ocorrer quando o suspeito é retirado de um banco de dados com milhares de pessoas.

Quanto a Probabilidade de Correspondência Aleatória, Thompson faz a seguinte estimativa (THOMPSON, 2003):

- Um para um bilhão (10^{-9}) – acontece quando os laboratórios conseguem uma amostra correspondente a da vítima em 10 ou mais loci;
- Um para um milhão (10^{-6}) – Isso ocorre quando um número menor de loci é analisado, ou quando o laboratório possui apenas uma parte de uma das amostras, ou ainda quando há a mistura de DNAs de mais de uma pessoa;
- Um para mil (10^{-3}) – Geralmente é resultado de testes menos eficazes, tais como DQ- Alpha/polymaker, particularmente quando envolvem amostradas misturadas.

No que tange a questão dos falsos positivos, vimos que que testes realizados nas décadas de 80 e 90, apontam que a taxa de falso positivo gira em torno de 1,5 a 4,7%. Como não tivemos uma grande quantidade de testes ao longo do tempo, o autor utiliza uma porcentagem mais conversadora, fazendo uma estimativa de 1% de falsos positivos. Além disso, o autor considera duas outras taxas de falso positivo, a taxa zero de falso positivo, a qual sabemos que irreal (apesar de alguns especialistas da área do direito ainda afirmarem ser possível) e a taxa de 1: 10000, que é a probabilidade de falso positivo combinada quando são feitos dois exames ($\frac{1}{100} \times \frac{1}{100}$) (THOMPSON, 2003).

Feita essas observações, iremos agora analisar os dados contidos na tabela:

Tabela 8 – Relação entre a probabilidade *a priori*, probabilidade de correspondência aleatória, taxa de falso positivo e a probabilidade *a posteriori*.

Probabilidade à Priori	Probabilidade de Correspondência Aleatória	Probabilidade de Falso Positivo	Probabilidade à Posteriori
2:1	10 ⁻⁹	0	2 000 000 000
2:1	10 ⁻⁹	0.0001	20 000
2:1	10 ⁻⁹	0.001	2000
2:1	10 ⁻⁹	0.01	200
2:1	10 ⁻⁶	0	2 000 000
2:1	10 ⁻⁶	0.0001	19 802
2:1	10 ⁻⁶	0.001	1998
2:1	10 ⁻⁶	0.01	200
2:1	10 ⁻³	0	2000
2:1	10 ⁻³	0.0001	1818
2:1	10 ⁻³	0.001	1001
2:1	10 ⁻³	0.01	182
1:10	10 ⁻⁹	0	100 000 000
1:10	10 ⁻⁹	0.0001	1000
1:10	10 ⁻⁹	0.001	100
1:10	10 ⁻⁹	0.01	10
1:10	10 ⁻⁶	0	100 000
1:10	10 ⁻⁶	0.0001	990
1:10	10 ⁻⁶	0.001	100
1:10	10 ⁻⁶	0.01	10
1:10	10 ⁻³	0	100
1:10	10 ⁻³	0.0001	91
1:10	10 ⁻³	0.001	50
1:10	10 ⁻³	0.01	9
1:100	10 ⁻⁹	0	10 000 000
1:100	10 ⁻⁹	0.0001	100
1:100	10 ⁻⁹	0.001	10
1:100	10 ⁻⁹	0.01	1
1:100	10 ⁻⁶	0	10 000
1:100	10 ⁻⁶	0.0001	99
1:100	10 ⁻⁶	0.001	10
1:100	10 ⁻⁶	0.01	1
1:100	10 ⁻³	0	10
1:100	10 ⁻³	0.0001	9
1:100	10 ⁻³	0.001	5
1:100	10 ⁻³	0.01	1
1:1000	10 ⁻⁹	0	1 000 000
1:1000	10 ⁻⁹	0.0001	10.0
1:1000	10 ⁻⁹	0.001	1.0
1:1000	10 ⁻⁹	0.01	0.1
1:1000	10 ⁻⁶	0	1000
1:1000	10 ⁻⁶	0.0001	9.9
1:1000	10 ⁻⁶	0.001	1.0
1:1000	10 ⁻⁶	0.01	0.1
1:1000	10 ⁻³	0	1.00
1:1000	10 ⁻³	0.0001	0.91
1:1000	10 ⁻³	0.001	0.50
1:1000	10 ⁻³	0.01	0.09

Fonte: THOMPSON, 2003.

Para deduzir a fórmula utilizada por Thompson para calcularmos a probabilidade de culpa do réu, partiremos do Teorema de Bayes, descrito abaixo:

$$\frac{P(S|R)}{P(S'|R)} = \frac{P(S)}{P(S')} \times \frac{P(R|S)}{P(R|S')}$$

Aonde:

s – 0 esperma veio do suspeito

s' - o esperma não veio do suspeito

R - Correspondência reportada

Inicia-se a análise distinguindo uma correspondência reportada (a qual se denota por R), de uma correspondência verdadeira a qual denotaremos por M . A partir disso, faremos as seguintes definições:

M : o suspeito e o esperma têm uma correspondência de DNA compatíveis

M' : o suspeito e o esperma não têm uma correspondência de DNA compatíveis.

É claro que não é possível saber se M ou M' é verdadeiro, pois a única coisa que temos é o resultado do DNA dado pelo laboratório, que pode estar errado.

Analisa-se primeiramente o numerador da fração $\frac{P(R|S)}{P(R|S')}$. Utilizando a definição de probabilidade condicional, se escreve $P(R|S)$ como $\frac{P(R \cap S)}{P(S)}$, aonde $P(R \cap S)$ é a probabilidade de o esperma vir do suspeito e o exame dar positivo. Com isso, é possível escrever $P(R \cap S)$ da seguinte maneira:

$$P(R \cap S) = P(R \cap S \cap M) \cup P(R \cap S \cap M')$$

Como os dois conjuntos são disjuntos, tem-se o seguinte:

$$P(R \cap S) = P(R \cap S \cap M) + P(R \cap S \cap M') .$$

Utilizando novamente a definição de probabilidade condicional, obtemos seguintes equações:

$$P(R \cap S \cap M) = P(R|M \cap S).P(M|S).P(S)$$

$$P(R \cap S \cap M') = P(R|M' \cap S).P(M'|S).P(S)$$

Como se está considerando $P(R|S)$, isso implica que o esperma já veio do suspeito e com isso, considera-se $P(S) = 1$.

Com isso, tem-se o seguinte:

$$P(R|S) = P(R|M \cap S).P(M|S) + P(R|M' \cap S).P(M'|S)$$

O raciocínio para o denominador $P(R|S')$ é feito de forma análoga e com isso escreve-se o seguinte:

$$\frac{P(R|S)}{P(R|S')} = \frac{P(R|M \cap S).P(M|S) + P(R|M' \cap S).P(M'|S)}{P(R|M \cap S').P(M|S') + P(R|M' \cap S').P(M'|S')}$$

Acerca desta expressão, é possível fazer algumas observações:

- $P(M|S) = 1$, pois se o esperma realmente vier do suspeito (S), ele obrigatoriamente teria um DNA correspondente.
- Com isso, há também que $P(M' \cap S) = 0$, pois não é possível que o esperma venha do suspeito e ele não tenha um DNA correspondente.
- Também será considerado que $P(R \cap M)$ é independente de S e S' , ou seja, que a probabilidade de uma correspondência relatada seja realmente correta não seja afetada pela probabilidade de correspondência aleatória. Com isso, pode ser escrito que:

$$P(R|M \cap S) = P(R|M \cap S') = P(R|M)$$

A partir destas observações, é possível escrever a seguinte equação:

$$\frac{P(R|S)}{P(R|S')} = \frac{P(R|M)}{P(R|M).P(M|S') + P(R|M).P(M'|S')}$$

Na equação acima:

- $P(R|M)$ é a probabilidade de o laboratório reportar uma correspondência caso o suspeito e o esperma tenham DNA compatíveis. Como a chance de obter um falso negativo em um exame de DNA é quase nula, consideraremos para efeito de cálculo que $P(R|M) = 1$.
- $P(M|S')$ é a probabilidade de correspondência aleatória, a qual denomina-se de *PCA*.
- De forma análoga, $P(M'|S')$ é o complementar da probabilidade de correspondência aleatória, a qual denotaremos por $1 - PCA$.

- Por fim, temos que $P(R|M')$ é a probabilidade de o laboratório reportar uma correspondência de DNA falsa, ou seja, é a taxa de falso positivo do exame, a qual denotaremos por PF .

Com isso, ao substituir na fórmula, obtém-se a seguinte equação:

$$\frac{P(R|S)}{P(R|S')} = \frac{1}{PCA + [TFP \cdot (1 - PCA)]}.$$

Através desta tabela, podem ser inferidas algumas informações importantes:

A primeira informação derivada da tabela é que, quando a probabilidade a priori é de 2:1, um resultado positivo no teste de DNA reflete uma probabilidade a posteriori consideravelmente alta de culpa do réu. Mesmo no cenário mais favorável possível para o réu, com uma probabilidade de correspondência aleatória de 1:1000 e uma taxa de falso positivo de 1:100, ainda resultaria em uma possibilidade de 182:1 contra o réu (THOMPSON, 2003).

Quando a probabilidade é de 1:10, as probabilidades a posteriori se tornam mais variadas, dependendo das probabilidades de falso positivo e de correspondência aleatória (variando de 1000:1 contra o réu até 9:1 contra o réu). Apesar de todas as probabilidades apontarem contra o réu, em pelo menos um cenário, o nível de certeza esperado para a condenação do réu não é atingido. Há um consenso de que a probabilidade justificativa para uma condenação deve ser pelo menos superior a 90%. (THOMPSON, 2003), ou seja, uma probabilidade de 9:1. Nesse sentido, a análise demonstra que, para uma probabilidade de correspondência aleatória de 10 – 3 e uma taxa de falso positivo de 1:100, esse limiar seria alcançado (THOMPSON, 2003).

Todavia, quando a probabilidade a priori é de 1:100, surgem os primeiros problemas graves acerca das condenações utilizando testes de DNA. Ao examinar os dados da tabela para essa probabilidade, é observado que em seis casos a probabilidade a posteriori é inferior ou igual a 90%, tornando-se insuficiente para a condenação do réu. Por fim, quando a probabilidade a priori é de 1:1000, em oito casos são geradas probabilidades a posteriori inferiores a 90% (THOMPSON, 2003).

Através disso, é possível inferir dois aspectos vitais no processo de DNA (THOMPSON, 2003):

- O primeiro deles é que falsos positivos existem e influenciam na probabilidade *a posteriori*. Viu-se na tabela que sempre ao colocar a probabilidade de falso positivo em 0, a

probabilidade *a posteriori* fica extremamente alta. Isso é de extrema relevância, pois essa informação não é levada em consideração pelos juízes;

- O segundo é que a probabilidade *a priori* é extremamente relevante para analisar a relevância do resultado do DNA. Quando temos evidências que o réu seja culpado, mas ainda temos dúvidas, o teste de DNA pode ser a prova final que era necessária para a sua condenação. Em contrapartida, se já tivermos grandes evidências de que o réu seja inocente, o exame de DNA não pode se sobrepor às evidências que temos.

Um exemplo de um caso em que isto ocorre foi o caso da condenação de Timothy Durham⁵, condenado a 3100 anos de prisão por conta do estupro de uma menina de 11 anos.

Durham foi considerado culpado pelo júri baseado em dois aspectos bem relevantes: O primeiro foi o teste de DNA recolhido do esperma da cena do crime, que era compatível com o esperma de Durham. O segundo ponto, foi o depoimento da vítima, uma menina de 11 anos que o reconheceu como sendo o homem que a estuprou.

Estas duas provas combinadas, seriam no geral, mais que suficientes para a condenação de qualquer réu. O que o juiz não levou em consideração ao dar a sentença do caso era o fato de que Durham possuía onze testemunhas que serviam de álibi, atestando que ele estava em Dallas, durante o período relevante para a análise do crime.

O que ocorreu neste caso foi uma superestimação das provas contra o acusado em detrimento aos álibis. Isto se deu provavelmente ao fato da crença que o exame de DNA ser uma prova cabal, irrefutável. Afinal como vimos, probabilidade de duas pessoas possuírem o mesmo DNA varia de 10^{-3} a 10^{-9} , que é bem baixa. Mas como a Tabela 8 mostra, para termos dados mais precisos acerca da probabilidade de o réu ser realmente culpado, temos que levar em conta além desse fator, a taxa de falso positivo, bem como a probabilidade *a priori* (que neste caso eram bem favoráveis a Durham por conta dos 11 álibis).

O que ocorreu neste caso foi uma interpretação errônea por parte do juiz do depoimento do promotor. O teste de DNA foi na verdade inconclusivo quanto à certeza de que aquele DNA pertencia realmente a Durham, indicando que ele estava em uma parcela pequena da população que se enquadrava nas características de DNA observadas.

Mesmo assim, ao confrontar essas provas com os álibis apresentados no processo jurídico contra o acusado, o mais plausível teria sido determinar uma nova coleta de DNA e adiar o anúncio da sentença, ao invés da decisão de condenação do réu. Após a condenação

⁵ Exemplo retirado da página 30 do livro *O Andar do Bêbado* (Mlodinow, 2009).

em 1993, foi realizado um novo teste de DNA, que indicou uma "forte evidência" de que "talvez ele não fosse culpado", ao que a promotoria rebateu alegando que esse novo teste não era suficiente para sua absolvição, visto que a vítima o havia reconhecido.

Durham só foi solto em dezembro de 1996, quase 5 anos depois de ser preso e mais de 3 anos após sua condenação, após exames de DNA excluírem qualquer ligação entre seu DNA e o do estuprador. Isso só foi possível devido ao uso de marcadores de DNA, que não existiam em 1993, possibilitando assim a comprovação de sua inocência.

6. A PROBABILIDADE E OS ESPORTES

Neste capítulo, serão abordadas questões relacionadas à aleatoriedade no mundo dos esportes. É notório que, no que concerne a este tema, todos acham ser profundamente conhecedores, desde torcedores, que mesmo não entendendo absolutamente nada de probabilidade, comemoram quando ouvem que seu time tem 90% de chance de ser campeão, até dirigentes que contratam um atleta que fez um grande sucesso na temporada anterior, mesmo não tendo dados estatísticos suficientes para embasar tal contratação.

O objetivo principal deste capítulo é mostrar como esses conceitos podem ser aplicados em esportes como o beisebol e o futebol, visando esclarecer ao leitor de onde são oriundas essas probabilidades, bem como instruir sobre as armadilhas que análises estatísticas e probabilísticas mal fundamentadas podem acarretar.

O capítulo será dividido em dois momentos fundamentais:

- O primeiro trata da odisseia de Maris e Mantle na busca da quebra do recorde de Babe Ruth de maior número de Home-Runs em uma única temporada. Este problema será subdividido em duas subseções. A primeira tratará da probabilidade de quebra de recorde de cada um deles, restando alguns jogos para o fim da temporada. A segunda tratará de um fenômeno chamado de Regressão à Média, que visa analisar a interligação entre um desempenho "acima da curva" com a real capacidade do atleta;
- No segundo momento, será abordado o uso da probabilidade aplicada a campeonatos de futebol. Aqui, buscar-se-á elucidar como são feitos cálculos envolvendo este esporte, como a probabilidade de um time ser campeão em um determinado campeonato de futebol, utilizando um modelo probabilístico como base.

6.1 O beisebol e o recorde de Babe Ruth

Era o ano de 1961. Eu mal tinha aprendido a ler, mas ainda me lembro das caras de Maris, do New York Yankees e de seu parceiro de equipe mais famoso que ele, Mantle, na capa da revista Life. Os dois jogadores estavam envolvidos em uma corrida histórica para igualar ou quebrar o celebrado recorde de Babe Ruth, estabelecido em 1927, de em um só ano, fazer 60 home runs (MLODINOW, 2009, p. 16).

O mundo dos esportes é fascinado pela quebra de recordes, seja no futebol, no atletismo ou em esportes menos conhecidos, como o beisebol. Todo atleta almeja a glória de superar seus limites e eternizar seu nome na história do esporte. Este desejo não era diferente para Roger Maris e Mickey Mantle. Na temporada de 1961 da MLB (*Major League Baseball*), ambos atuavam pelo New York Yankees e estavam tendo um desempenho extraordinário, o que despertou nos torcedores a expectativa de que o recorde de home runs de Babe Ruth, estabelecido em 1927, pudesse ser batido.

Para compreender melhor essa disputa, é importante entender como funciona um jogo de beisebol. A Figura 6 a seguir ilustra um campo de beisebol.

Figura 6 – Ilustração de um campo de beisebol



Fonte: <https://www.todoestudo.com.br/educacao-fisica/beisebol>

O jogo de beisebol é disputado por 9 jogadores em cada time e consiste em 9 entradas. Durante cada entrada, os times alternam entre as posições de ataque e defesa. Para iniciar o jogo, o arremessador lança a bola em direção ao home plate, onde estão o rebatedor e o catcher. O objetivo principal do arremessador é acertar a bola na zona de strike, uma área imaginária entre os joelhos e o torso do rebatedor. Se o rebatedor sofrer três strikes, ele é eliminado e vai para o fim da fila do seu time.

Se o rebatedor acertar a bola, ele deve correr em direção à primeira base antes que um jogador de defesa consiga tocá-lo com a bola enquanto ele corre ou antes que um jogador de

defesa pegue a bola no ar. Se um jogador de defesa tocar o rebatedor com a bola enquanto ele corre ou pegar a bola no ar, o rebatedor é eliminado. Se a equipe defensora conseguir eliminar três jogadores da equipe atacante, as equipes trocam de posição: a equipe defensora passa a ser a atacante e vice-versa.

Para marcar pontos, um jogador deve percorrer todas as bases e retornar ao *home plate*. A equipe marca um ponto cada vez que um jogador completa o percurso e retorna ao home plate. No final das 9 entradas, a equipe com mais pontos vence o jogo.

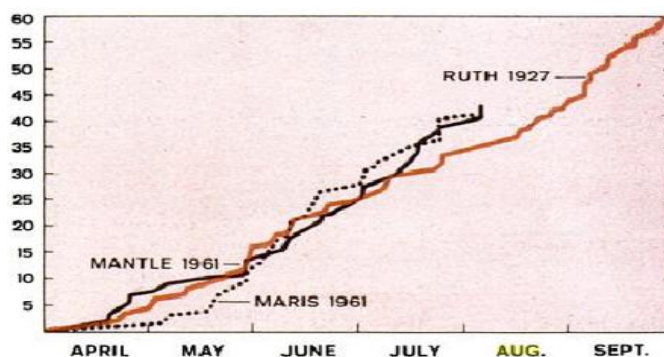
- **Mas, o que seria um Home Run?**

No contexto do beisebol, um Home Run é considerado a jogada perfeita. Consiste no rebatedor lançar a bola para fora da área de jogo. Com isso, o rebatedor e todos os outros jogadores do time atacante que estiverem em alguma base podem avançar automaticamente pelas quatro bases, marcando assim um ponto cada.

6.1.1 A corrida de Maris e Mantle em busca do recorde

No ano de 1961, a revista Life, em sua publicação do mês de agosto, fez um breve histórico dos competidores, analisando as chances de cada um deles quebrar o recorde de Babe Ruth. Àquela altura, já haviam se passado 110 dos 154 jogos da temporada. O gráfico da Figura 7 mostra o desempenho de cada um dos atletas após os primeiros 110 jogos da temporada, fazendo um comparativo entre as performances de Mantle e Maris e a performance do recordista Babe Ruth ao longo de toda a temporada.

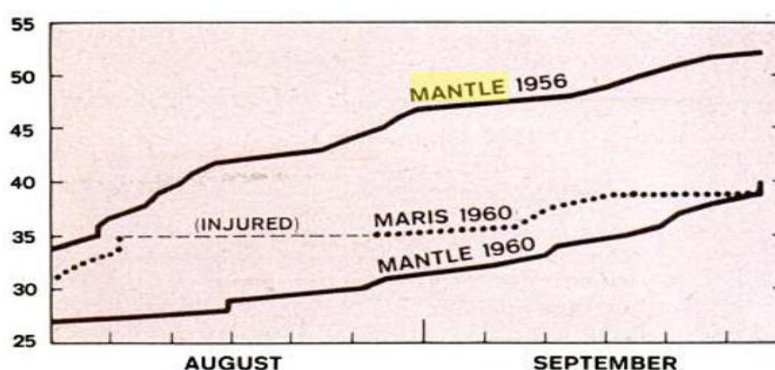
Figura 7 – Desempenho dos desafiantes em relação a Babe Ruth



Fonte: Revista Life, 1961.

A partir do gráfico, pode-se observar que em agosto de 1961, tanto Mantle quanto Maris estavam à frente do recorde de Babe Ruth em número de home runs. Isso gerava a expectativa de que o recorde poderia ser quebrado.

Figura 8 – Desempenho dos desafiantes em anos anteriores a 1961.



Fonte: Revista Life, 1961.

Da análise do gráfico acima, é possível inferir que Mantle na temporada de 1956 e Mantle e Maris no ano de 1960 não tiveram fins de temporada tão ascendentes quanto na temporada de quebra de recorde de Babe Ruth (conforme mostrado no gráfico 1).

A partir desses dados, a revista calculou a probabilidade de cada um deles bater o recorde de Babe Ruth, baseando-se tanto no fato de estarem à frente no número de home runs quanto no desempenho menos destacado no final da temporada, em comparação com Ruth. Para estimar essa probabilidade, foi utilizado um modelo probabilístico chamado de distribuição de Bernoulli.

Para aplicar esse modelo probabilístico, é necessário ter uma variável que admita apenas dois valores: sucesso ou fracasso. Nesse caso, a variável estudada é uma variável de Bernoulli, pois em cada rebatida que um jogador de beisebol faz, existem apenas duas possibilidades: fazer o *home run* (sucesso) ou não o fazer (fracasso).

Antes de analisar as chances de cada um dos desafiantes, é importante definir esse modelo e algumas propriedades inerentes a ele.

- **Distribuição de Bernoulli**

Se uma variável aleatória assume somente os valores **0** e **1**, com probabilidades $1 - p$ e p , ou seja, se sua distribuição for dada por:

x	0	1
P(X=x)	$1-p$	p

Então neste caso confirma-se que X é uma variável de Bernoulli.

Estes valores **0** e **1** são denominados de sucesso ou fracasso, mas não no sentido em que se utiliza rotineiramente. O sucesso é a probabilidade de determinado evento ocorrer e o fracasso é a probabilidade deste mesmo evento não ocorrer. Um exemplo de evento que é modelado pela distribuição de Bernoulli é um lançamento de uma moeda.

É possível definir este evento da seguinte maneira:

$$X = \begin{cases} 1, & \text{se der coroa} \\ 0, & \text{se der cara} \end{cases}$$

Neste caso a distribuição de Bernoulli é dada por:

x	0	1
P(X=x)	$1/2$	$1/2$

É evidente que não é uma condição necessária o experimento ser equiprovável para ser considerado uma distribuição de Bernoulli. A única restrição que se aplica a esse tipo de distribuição é que só podem existir dois resultados possíveis. Portanto, em experimentos com múltiplos resultados possíveis, como o lançamento de um dado, não é viável utilizar a distribuição de Bernoulli.

- **Distribuição Binomial**

Se um experimento consiste em n repetições independentes e sendo a probabilidade de sucesso em cada repetição constante e igual a p , então dizemos que este é um experimento binomial.

Se uma variável aleatória X corresponde ao número de sucessos em n repetições de um experimento binomial, sendo p a probabilidade de sucesso em cada uma dessas repetições então dizemos que X tem distribuição binomial com parâmetros n e p , e sua função de distribuição de probabilidades é dada por:

$$P(X = k) = \binom{n}{k} p^k (1 - p)^{n-k}, k = 0, 1, 2, \dots, n$$

Este modelo binomial foi utilizado pela revista Life para estimar a probabilidade de cada um dos atletas quebrarem o recorde de home runs. Com isso, utiliza-se a fórmula da distribuição binomial para estimar a chance de cada um deles bater o recorde de Ruth.

$$P(X = k) = \binom{n}{k} p^k (1 - p)^{n-k}, k = r, r + 1, \dots, n$$

Aonde:

n – é o número de rebatidas que cada um terá até o fim da temporada

r – Quantidade de home runs que falta para cada um deles quebrar o recorde

p – Probabilidade de um determinado jogador executar um home run em uma rebatida

$1 - p$ – Probabilidade de um determinado jogador não fazer este home run.

O único ajuste realizado na fórmula da distribuição é que a variável k varia não de zero até n , mas sim de x até n , onde x representa o número mínimo de home runs necessários para quebrar o recorde. Isso ocorre porque há um interesse apenas no valor da probabilidade de fazer x ou mais home runs, pois é esse caso que garante a quebra do recorde.

Para estimar o valor da probabilidade p de cada um deles bater o recorde, é necessário inferir a probabilidade de cada um deles acertar um home run em uma rebatida. Para isso, será utilizada a informação do número de rebatidas e de home runs que cada um deles realizou até o jogo de número 110.

Inicialmente, serão feitas duas estimativas. A primeira considerará que eles precisariam quebrar o recorde em 154 jogos, enquanto a segunda suporá que eles poderiam quebrar o recorde até o fim da temporada, no jogo de número 162. Essa distinção é feita devido à divergência entre os membros da confederação americana de Beisebol em 1961 sobre se o recorde deveria ser batido em 154 jogos (número de jogos que Babe Ruth fez em 1927) ou se seria homologado mesmo que fosse quebrado após o jogo 154.

A primeira análise será feita considerando a necessidade de quebra de recorde até o jogo 154. Primeiramente, será analisado o desempenho de Mickey Mantle. Mantle teve um total de 378 rebatidas nos 110 primeiros jogos da temporada⁶, com um total de 43 home runs até o presente momento. Com isso, estima-se que a probabilidade de Mantle executar um home run em uma determinada rebatida é de $\frac{43}{378} = 0,11375$, ou de 11,37%. Como ele fez 378 rebatidas em 110 jogos, concluindo que o número de rebatidas por jogo é de 3,43 rebatidas.

Como naquela altura ainda restavam 44 jogos, Mantle teria ainda uma média de 151,2 rebatidas (valor que arredondado para 151). Logo, como o recorde anterior de Babe Ruth era de 60 home runs, Mantle teria que fazer mais 18 home runs nas 151 rebatidas que ele ainda teria no restante da temporada.

Logo, ao substituírmos os valores $p = 0,1137$, $1 - p = 0,8863$, $n = 151$ e $r = 18$, na distribuição binomial, ficaremos com: $\sum_{k=18}^{151} \binom{151}{k} (0,1137)^k (0,8863)^{151-k} = 0,4531097$, o que dá uma porcentagem de quebra de recorde de aproximadamente 45,31%.

Agora pretende-se analisar o desempenho de Roger Maris. Maris teve um total de 403 rebatidas nos 110 primeiros jogos da temporada, com um total de 41 home runs até o presente momento. Com isso, estima-se que a probabilidade que Mantle de executar um home run em uma determinada rebatida é de $\frac{41}{403} = 0,1017$, ou de 10,17%. Como ele fez 417 rebatidas em 110 jogos, concluímos que o número médio de rebatidas por jogo é de 3,66 rebatidas. Como naquela altura ainda restavam 44 jogos, Maris teria ainda uma média de 161,2 rebatidas (valor que arredondaremos para 161). Logo, como o recorde anterior de Babe Ruth era de 60 home runs, Mantle teria que fazer mais 20 home runs nas 161 rebatidas que ele ainda teria no restante da temporada.

Logo, ao substituírmos os valores $p = 0,1017$, $1 - p = 0,8983$, $n = 161$ e $r = 20$, na distribuição binomial, ficaremos com a seguinte distribuição binomial: $\sum_{k=20}^{161} \binom{161}{k} (0,1007)^k (0,8983)^{161-k} = 0,203995$. Ou seja, Maris tinha àquela altura uma probabilidade de quebra de recorde de 20,40% .

Com isso, pode-se inferir que a probabilidade de quebra de recorde seria de:

$$0.4531 + 0.2039 - 0.4531 \times 0.2039 = 0.5646, \text{ ou } 56,46\%.$$

⁶ Estas estatísticas usadas para fazermos a estimativa estão disponíveis no site <https://www.baseball-almanac.com/>.

Outra análise que pode transcorrer se dá pela probabilidade de quebra do recorde, considerando o total de 162 partidas. Considerando este número de jogos, faltariam 52 partidas a serem disputadas. Com isso, Mantle teria uma média de $3,43 \times 52 = 178,36$ (arredondaremos para 178). Com isso, utilizando os valores descritos no caso anterior, a probabilidade de Mantle quebrar o recorde seria dada por:

$\sum_{k=18}^{178} \binom{178}{k} (0,1137)^k (0,8863)^{178-k} = 0,7351777$, ou aproximadamente 73,52%.

Já Maris, teria uma média de $3,66 \times 52 = 190,32$ (arredondaremos para 190). Com isso, ao utilizarmos os dados anteriores, pode-se inferir a probabilidade de quebra de recorde de Maris é dada por $\sum_{k=20}^{190} \binom{190}{k} (0,1017)^k (0,8983)^{190-k} = 0,4703297$, ou aproximadamente 47,03%.

A partir disso, é possível concluir que a probabilidade de quebra de recorde, por um dos 2 competidores, neste caso é dada por:

$$0,7351 + 0,4703 - 0,7351 \times 0,4703 = 0,8596 = 85,96\%.$$

Por fim, uma terceira análise pode ser realizada tomando como base as performances dos últimos jogos disputados por Mantle e Maris. Para isso, será analisado o desempenho de Mantle e Maris no mês de julho e nos primeiros jogos do mês de agosto (até o jogo 110). Supondo que o recorde deva ser batido em 154 jogos, durante esse período foram disputados 38 jogos, nos quais Mantle realizou 131 rebatidas, das quais 18 resultaram em Home Runs. Com isso, conclui-se que a probabilidade de acerto de home runs, baseado nesse recorte temporal, é calculada $\frac{18}{131} = 0,1374$. Como ele rebateu 131 vezes nos 38 jogos, temos que a sua média de rebatidas por jogo é de 3,44, o que resulta em aproximadamente 152 rebatidas nos 44 jogos restantes.

Com isso, a probabilidade de quebra de recorde neste caso é dada por: $\sum_{k=18}^{152} \binom{152}{k} (0,1374)^k (0,8626)^{152-k} = 0,7844 = 78,44\%$, o que é uma probabilidade bem maior que os 45,31% calculados anteriormente, apontando assim para um viés de crescimento de Mantle.

Analogamente, o mesmo cálculo será realizado para Maris. Durante esse mesmo período de 38 jogos, Maris rebateu 141 vezes, das quais 14 se transformaram em Home Runs. Com isso, a probabilidade de acerto de Home Run nesse período é calculada por:

$\frac{14}{141} = 0,099 = 9,99\%$. Como ele rebateu 141 vezes em 38 jogos, temos que a sua média de rebatidas por jogo é de 3,710, o que resulta em aproximadamente 163 rebatidas nos 44 jogos restantes. Com isso, a probabilidade de quebra de recorde é dada por:

$$\sum_{k=20}^{163} \binom{163}{k} (0,099)^k (0,901)^{152-k} = 0,1865 = 18,65\%$$

om isso, temos que, para a probabilidade de pelo menos um dos 2 quebrarem ok recorde é dada por:

$$0,7844 + 0,1865 - 0,7844 \times 0,1865 = 0,8246 = 82,46\%$$

Apesar de todos esses dados apontarem um enorme favoritismo para Mantle quebrar o recorde de Ruth, não foi isso o que realmente aconteceu. Mantle sofreu uma lesão que o impediu de bater o recorde. Por outro lado, Maris, contra todos os prognósticos, rebateu seu 61º Home Run no jogo de número 160, eternizando assim o seu nome na história do beisebol.

6.1.2 Mas será que foi isso mesmo que aconteceu ou Maris por puro “acaso do destino” teve uma temporada discrepante do seu desempenho médio?

Para estudar esse fenômeno, Mlodinow (2009) desenvolveu um modelo matemático que considera não apenas a habilidade do jogador, mas também outros fatores que podem influenciar sua performance, como vento, sol, iluminação do estádio e qualidade do arremesso dos outros jogadores. Ao longo de uma temporada, esses fatores tendem a se equilibrar, mas ocasionalmente podem não se anular. A simulação realizada por Mlodinow mostrou que esses eventos não são tão raros quanto se poderia supor.

O desempenho de Maris na temporada anterior à quebra do recorde foi de 1 home run a cada 14,7 rebatidas, uma média similar às quatro temporadas seguintes à quebra do recorde. Por meio de simulações computacionais, Mlodinow constatou que, na maioria das temporadas simuladas, o desempenho de Maris permanecia próximo de sua média, com algumas variações leves para mais ou para menos. A chance de ele bater o recorde puramente por acaso foi calculada como sendo de 1 em cada 32 temporadas, um valor surpreendentemente alto.

Após a quebra do recorde, Maris não se tornou um jogador ruim ou sucumbiu à pressão; ele simplesmente voltou ao seu desempenho médio. Esse fenômeno é conhecido

como regressão à média, que é a tendência de uma variável aleatória retornar à sua média após períodos de variação extrema, seja positiva ou negativa.

Para ilustrar esse conceito, uma situação descrita no livro "Rápido e Devagar: Duas Formas de Pensar" é mencionada.

Durante uma conversa sobre psicologia de treinamento com instrutores de voo da força aérea israelense, Kahneman foi interrompido por um dos instrutores, que compartilhou sua observação de que elogiar os cadetes por uma manobra perfeita muitas vezes resultava em desempenho inferior na próxima tentativa, enquanto criticar um erro frequentemente levava a uma execução melhor na tentativa seguinte. Isso ilustra como a recompensa pelo sucesso e a punição pelo fracasso podem produzir resultados opostos, evidenciando a tendência de regressão à média (KAHNEMAN, 2012, p. 35).

A afirmação do instrutor está correta, uma vez que observou que após uma falha, os pilotos tendem a realizar uma manobra melhor, enquanto cada erro é seguido por uma melhoria na execução subsequente. No entanto, o equívoco de análise reside na interpretação das causas dessas melhorias ou pioras. Cada piloto possui um desempenho médio em suas manobras, que pode variar ao longo do tempo, mas não entre uma manobra e outra (KAHNEMAN, 2012).

No exemplo apresentado, a variação extrema no desempenho ocorre devido à aleatoriedade. Um piloto que por “sorte” realiza uma manobra excepcional provavelmente terá uma performance mais próxima de sua média na próxima tentativa, resultando em uma aparente piora. Da mesma forma, um piloto que falha em uma manobra provavelmente terá uma performance mais próxima de sua média na tentativa seguinte, resultando em uma aparente melhoria (KAHNEMAN, 2012).

É comum que as pessoas busquem padrões em diversos aspectos da vida, desde a natureza até desempenhos esportivos ou estimativas de bilheteria de filmes. A mente humana tende a procurar causas e explicações para fenômenos observados. Entretanto, muitas vezes, o que ocorre são simplesmente eventos aleatórios. Em qualquer área de atuação, pode haver momentos em que os desempenhos estão muito acima ou muito abaixo da capacidade real, devido a fatores aleatórios.

A ocorrência do fenômeno de regressão à média é frequentemente interpretada como uma relação de causa e efeito. Por exemplo, a pressão sobre Maris pode ser vista como a causa de seu declínio de desempenho, ou os gritos dos instrutores podem ser considerados como a causa da melhoria de desempenho dos cadetes. Embora esses fatores possam influenciar o desempenho, tendemos a superestimá-los e a desconsiderar a natureza aleatória

desses processos. Esse tipo de pensamento pode levar a escolhas erradas, baseadas principalmente em desempenhos excepcionais.

Embora nem sempre seja possível coletar uma quantidade infinita de dados sobre um profissional em determinada área, é importante fazer isso na medida do possível ou, pelo menos, ser capaz de avaliar os riscos associados a decisões baseadas em desempenhos excepcionais.

6.2 O futebol e a chance de um time ser campeão

É comum encontrar informações sobre cálculos de probabilidade relacionados a campeonatos de futebol, onde expressões como "Time A tem 80% de chance de ser campeão" ou "time B tem 90% de chance de ser rebaixado" são frequentemente utilizadas por torcedores engajados. Nesta seção, será abordado como esses cálculos de probabilidade são realizados.

Um dos principais desafios dessa análise é lidar com a grande quantidade de possibilidades que esses cálculos podem gerar. Isso coloca esse tipo de análise em um espectro completamente diferente de problemas mais simples, como o exemplo anterior, onde havia uma estimativa de um número relativamente pequeno de possibilidades em relação ao número de rebatidas dos jogadores.

Por exemplo, calcular a chance de um time ser campeão faltando 15 rodadas seria inviável devido à quantidade absurda de combinações de resultados possíveis para os jogos restantes. Em um campeonato com 20 times e 10 jogos por rodada, teríamos 150 jogos restantes, totalizando 3150 combinações possíveis. Mesmo com um computador ultramoderno, levaria um tempo impossível para gerar todas essas combinações.

Vale ressaltar que, mesmo se fosse possível processar todas essas combinações, ainda assim haveria um erro, pois as probabilidades de vitória, empate e derrota de um determinado time não são iguais. Elas variam de acordo com diversos fatores, como a qualidade do elenco, o fator casa, e a qualidade do adversário. Portanto, é necessário construir um modelo que leve em consideração essas características e seja computacionalmente viável para melhor avaliar essas probabilidades.

Para realizar essa abordagem, em vez de analisar cada um dos resultados possíveis, o computador simula os resultados dos jogos do torneio, o que é computacionalmente mais eficiente e rápido. Em questão de minutos, é possível gerar milhares de resultados de torneios.

Com base nessas simulações, é possível calcular a probabilidade de um time ser campeão faltando 15 rodadas, dividindo o número de vezes que o time foi apontado como campeão pelo número total de simulações realizadas. Por exemplo, se em 10.000 simulações o Flamengo terminou em primeiro lugar em 6.874 delas, estima-se que a probabilidade do Flamengo ser campeão gira em torno de 68,74%, sendo importante ressaltar que quanto maior o número de simulações, mais precisa será essa estimativa.

No entanto, é crucial que o computador não faça essa simulação de forma aleatória, pois os resultados gerados não teriam aplicação prática. Para isso, é necessário "ensinar" o computador a simular cada um dos jogos, de modo que em cada partida, ele receba dois vetores que caracterizam os times A e B, denotados por $V_A = (PV_A, PE_A, PD_A)$ e $V_B = (PV_B, PE_B, PD_B)$, aonde PV_A, PE_A, PD_A , representam, respectivamente a probabilidade de vitória, empate e derrota do time A, assim como, de forma análoga PV_B, PE_B e PD_B representam a probabilidade de vitória, empate e derrota do time B. Cada um desses vetores é chamado de vetor força de um time.

A partir desses vetores, o computador pode simular cada jogo de forma realista, utilizando o seguinte vetor: $P_{AXB} = (\frac{PV_A+PD_B}{2}, \frac{PE_A+PE_B}{2}, \frac{PD_A+PV_B}{2})$. Por exemplo, suponha-se que o vetor força do time A para esse jogo seja (0,6; 0,3; 0,1), enquanto o vetor força do time B para este jogo seja (0,5, 0,1, 0,4). O Vetor probabilidade P_{AXB} para este confronto será (0,5 ; 0,2; 0,3. As coordenadas deste vetor representam respectivamente, a probabilidade de vitória de A, empate e vitória de B.

Para simular o resultado de uma partida, o computador divide o intervalo real [0,1] em três partes de acordo com as probabilidades de vitória, empate e derrota do time A. Em seguida, ele sorteia um número real entre 0 e 1 e associa esse número ao resultado da partida, seguindo os intervalos correspondentes às probabilidades pré-definidas pelo vetor força.

No início de cada campeonato, é atribuído um vetor força para cada time, levando em consideração diversos fatores, como investimento na temporada, qualidade do elenco, desempenho em campeonatos anteriores e o fator casa. Para cada time, são adotados dois vetores de força: um para jogos em casa (PC) e outro para jogos fora de casa (PF). Após cada rodada, esses valores são atualizados de acordo com um algoritmo baseado no modelo probabilístico escolhido.

Para ilustrar esse processo, considera-se o modelo proposto por Lima et al. (2012) no seu artigo "Probabilidades para o esporte" para a Universidade Veiga de Almeida.

Por exemplo, em uma partida entre Flamengo e Palmeiras no estádio do Maracanã, onde o Flamengo joga como mandante e o Palmeiras como visitante, utilizamos os vetores de força PC e PF, representando as probabilidades de vitória, empate e derrota do Flamengo em casa e do Palmeiras fora de casa, respectivamente. Se a simulação apontar o Flamengo como vencedor, o vetor PC do Flamengo é modificado para aumentar a probabilidade de vitória em casa e diminuir as probabilidades de empate e derrota, conforme uma fórmula específica de atualização:

$$PCN = \frac{p.PCA + r.(1,0,0)}{p+r}, \text{ onde temos:}$$

PCN – é o vetor probabilidade do Flamengo em jogos em casa, atualizado após a vitória nesse jogo.

PCA- é o vetor probabilidade do Flamengo anterior a esta vitória.

p- é o peso que damos ao passado, ou seja, o quanto o modelo adotado quer que essa vitória altere a força do time. Esse valor representa o quão sensível se quer que a vitória do time seja. Por exemplo, consideremos hipoteticamente o caso em que $p = 0$ ou seja, passado não tem importância alguma. Neste caso, ao substituirmos $p = 0$ na fórmula com $\frac{r(1,0,0)}{r} = (1,0,0)$

Basicamente, a probabilidade de vitória do Flamengo passaria a ser de 100%, ignorando toda a probabilidade pré-estabelecida. Da mesma forma, se $p = 1$, a vitória acarretará uma mudança bem pequena deste vetor, pois como dessa forma o multiplicador p em cima do *PCA* irá enfatizar muito a probabilidade de vitória anterior ao jogo, bem como o “ r ” no denominador aumentará o seu valor, fazendo que assim também diminua a *PCN*.

r- é o rendimento do Palmeiras fora de casa. Este valor é calculado da seguinte forma:

$$r = \frac{\text{pontos conquistados fora de casa pelo Palmeiras}}{\text{pontos disputados pelo Palmeiras fora de casa}}$$

Na fórmula para atualização do vetor força após uma derrota, é utilizado $(1-r)$ em vez de r , como na fórmula anterior. Isso ocorre porque, nesse caso, aplica-se o raciocínio contrário. Se o time perdeu para um adversário com alto desempenho, essa derrota não deve

ter muita influência no vetor força, enquanto uma derrota para um time com baixo desempenho deve influenciá-lo mais.

Quanto ao resultado de empate, a atualização dos vetores força de cada time depende do desempenho do adversário. Por exemplo, se o Flamengo e o Palmeiras empatarem em uma simulação, a atualização do vetor força do Flamengo em casa é determinada pelo desempenho do Palmeiras fora de casa. Se esse desempenho for menor que 0,5, o empate terá um peso negativo no vetor força, considerado como “meio empate, meio derrota”. Se o desempenho do Palmeiras for maior que 0,5, o empate terá um peso positivo no vetor força do Flamengo, considerado como “meio empate, meio vitória”. A fórmula para esse cenário é analisada separadamente, começando com o caso em que o rendimento do Palmeiras é menor que 0,5.

$$PCN = \frac{p.PCA + (1-2r).(0;0,5;0,5) + 2r(0,1,0)}{p+1}$$

Acerca desta fórmula, é possível fazer algumas observações:

- As variáveis desta fórmula são as mesmas que já foram definidas para o caso vitória do Flamengo;
- O vetor $(0; 0,5; 0,5)$ representa o que foi anteriormente colocado. Como o desempenho do Palmeiras é ruim fora de casa, esse resultado é considerado, metade empate, metade derrota. Este termo é multiplicado pelo fator $(1 - 2r)$. Como $0 \leq r \leq 0,5$, este fator varia de 0 até 1. Caso $r=0$, $(1 - 2r)$ seria igual a 1, ou seja, como o Flamengo perdeu de um time que não havia pontuado. Com isso, o segundo termo seria maximizado, e, portanto, este empate teria um peso muito grande. Caso r seja igual a 0,5, o fator $(1-2r)$ seria igual a zero, e com isso este termo não teria nenhuma influência no novo vetor força. É possível destacar, neste caso, que o empate “não seria nem bom, nem ruim”.

O vetor $(0,1,0)$ ao contrário do vetor anterior, não altera a componente “derrota” do vetor força do time, apenas a componente empate, ou seja, pode-se dizer que este termo representa o “empate puro”. Esse vetor, está sendo multiplicado pelo fator $2r$. Como, temos que $0 \leq r \leq 0,5$, este fator varia de 0 até 1. Se $r = 0$, esse termo se anula, enquanto se $r = 0,5$ esse termo atinge o seu valor máximo.

Percebe-se que o segundo e o terceiro termos do denominador são opostos entre si. Quanto menor for o rendimento r do Palmeiras, maior será o segundo termo e menor o

terceiro termo. Quanto maior for o rendimento do Palmeiras, menor será o segundo termo e maior o terceiro.

Assim, a fórmula consegue equilibrar os empates contra times "bons" e times "ruins", de modo que a probabilidade de derrota no próximo jogo aumente para os times que empataram contra times ruins ($0 \leq r < 0,5$), permaneça a mesma para times medianos ($r = 0,5$) e diminua em caso de empate com times bons ($r < 0,5$).

Para o último caso, é necessário fazer um ajuste na fórmula, de modo que os fatores que multiplicam os vetores fiquem entre 0 e 1. Além disso, o vetor do segundo termo deve ser alterado para (0,5; 0,5; 0), pois o empate agora indica um aumento na probabilidade de vitória do time, devido ao empate com um time extremamente qualificado. Assim, a fórmula ajustada fica da seguinte forma:

$$PCN = \frac{p.PCA + 2r.(0,5;0,5;0) + 2(1-r).(0,1,0)}{p+1}$$

CONCLUSÃO

No livro “O Andar do Bêbado” de Mlodinow (2009), uma tentativa é feita para aproximar os conceitos de acaso da população leiga. A obra é conhecida por sua riqueza em exemplos, proporcionando aos leitores familiaridade com vários aspectos da probabilidade que, de outra forma, poderiam ser desconhecidos. Seu foco principal reside na maneira equivocada como a probabilidade e os conceitos relacionados ao acaso são percebidos pelo público em geral e, muitas vezes, até por especialistas em áreas como matemática, como demonstrado no problema de Monty-Hall, bem como em outras áreas, como ilustrado nos Capítulos 4 e 5, onde médicos e juízes falham na interpretação de exames clínicos e testes de DNA, respectivamente.

O trabalho atual explorou os diversos problemas apresentados no livro "O Andar do Bêbado", buscando preencher lacunas deixadas por Mlodinow. Para isso, foram abordados três aspectos principais considerados incompletos. O primeiro aspecto abordado é a resolução dos problemas apresentados no livro. Apesar da linguagem acessível ao público leigo, muitos problemas carecem de uma solução matemática detalhada, o que dificulta a compreensão por parte daqueles menos familiarizados com a matemática. Desta forma, o estudo procurou explicitar de forma mais detalhada as nuances de cada problema, buscando analisar outros aspectos associados a eles, como é feito, por exemplo, no problema dos astrágalos e no problema das portas.

Um segundo aspecto do trabalho consistiu em dar uma ênfase maior aos temas abordados pelo livro "O Andar do Bêbado". Embora o trabalho tenha como objetivo principal adotar uma linguagem acessível para diversos níveis de familiaridade com a matemática, foi necessária uma abordagem mais robusta em alguns pontos. Por exemplo, no primeiro capítulo, uma análise mais abrangente foi realizada em relação aos mecanismos que levam o cérebro a cometer erros na análise da probabilidade.

Outro aspecto importante, menos enfatizado por Mlodinow, foi a questão: "O que se pode fazer para uma melhor análise da probabilidade?" Isso levanta duas reflexões relevantes: em primeiro lugar, por vezes as informações simplesmente não são comunicadas de forma clara para as pessoas. Isso foi evidenciado no Capítulo 5, sobre os exames clínicos. Quando as informações sobre mamografia foram apresentadas em formato percentual, a maioria dos médicos cometeu equívocos ao inferir a probabilidade real de uma pessoa estar doente após receber um diagnóstico positivo para a doença. Contudo, quando Gigerenzer conduziu o

experimento apresentando os dados em valores absolutos, houve um aumento significativo na capacidade dos médicos de interpretar corretamente o exame.

Outro ponto relevante a ser considerado é a abordagem do tema "probabilidade" pelos professores. Muitas vezes, os problemas apresentados em sala de aula estão relacionados a espaços amostrais equiprováveis, o que pode levar os alunos a cometerem erros ao aplicar esses conceitos a situações do cotidiano em que essa igualdade de probabilidades não existe, como observado nos problemas de D'Alembert e dos filhos (Capítulo 4).

Este estudo, portanto, representa um complemento à obra de Mlodinow, abordando aspectos que apresentaram lacunas. Além de oferecer uma análise mais detalhada das resoluções de problemas, também busca explorar outras dimensões dos conceitos apresentados no livro, enriquecendo assim a compreensão do leitor sobre a probabilidade e o acaso. Dessa forma, espera-se que esta dissertação tenha um impacto significativo tanto para os docentes, auxiliando-os a abordar o tema de maneira mais eficaz, quanto para os alunos, estimulando o interesse e proporcionando uma compreensão mais abrangente e precisa dos fenômenos de natureza probabilística.

REFERÊNCIAS

- BASTOS, N. M. S. **Síntese e caracterização de novos polímeros orgânicos condutores derivados da polianilina**. 2004. 73f. Trabalho de conclusão de curso. Universidade Estadual do Norte Fluminense Darcy Ribeiro. Campo dos Goytacazes: UFF, 2004.
- BERRY, DONALD A. & CHAINSTAN. *Inference About Testosterone Abuse Among Athletes*. CHANCE: Vol. 17, No. 2, p. 5-8 , 2004.
- BONACCORSO, NORMA SUELI. **Aplicação do exame de DNA na elucidação de crimes**. 2005. 193f. Dissertação (Mestrado em Medicina Forense) – Universidade de São Paulo. São Paulo: USP, 2005.
- CRISTIANO, M. V. de M. B. **Sensibilidade e especificidade na curva de ROC: um estudo de caso**. 2017. Dissertação (Mestrado em Gestão de Sistemas de Informação Médica) – Universidade do Porto. Leiria, Portugal: UP, 2017.
- FELLER, W. **An Introduction to Probability Theory and Its Applications**. Vol. 1. Wiley. 3 edition, 1968.
- GIGERENZER, G. *Calculated Risks: How to Know When Numbers Deceive You*. 1ª Editions. New York: Simon & Schuster, 2002. 320p.
- KAHNEMAN, D. **Rápido e Devagar: Duas Formas de Pensar**. 1ª edição. Rio de Janeiro: Objetiva, 2012.
- KHANEMAN, D. et al. **Judgement under uncertainty: heuristic and biases**, 1ª edição, Cambridge: Cambridge University Press, 1982
- KAHNEMAN, D., & TVERSKY, A. *Availability: A heuristic for judging frequency and probability*. **Cognitive Psychology**, 1973, 5(2), 207–232. [https://doi.org/10.1016/0010-0285\(73\)90033-9](https://doi.org/10.1016/0010-0285(73)90033-9)
- KHOEHLER, JONATHAN J., *Error and exaggeration of DNA Evidence at Trial*, *Jurimetrics*, Vol. 34, p.21-39 , 1993.
- LIMA, BERNARDO NUNES B, et al, *Probabilidades no Esporte*, *Trim*, Vol. 5 , p.39-53, 2012
- LOTÉRIAS CAIXA. **Mega-sena**. 2023. Disponível em: <https://www.caixa.gov.br/loterias/comunicados-importantes/Paginas/default.aspx>. Acesso em: 15 out. 2023.
- MLODINOW, L. **O andar do bêbado. Como o acaso determina nossas vidas**. Rio de Janeiro: Zahar, 2009.
- MARQUES, J. P. **Beisebol. Todo Estudo**. 2023. Disponível em: <https://www.todoestudo.com.br/educacao-fisica/beisebol>. Acesso em: 10 nov. 2023.
- THOMPSON, R. B. hompson, R. B. **Mathematics for Business Decisions: Probability and Simulation**. Mathematical Association of America, 2003.

VON MISES, R., Probability, Statistics and Truth, 1928.

REVISTA LIFE, Math muscles is on the Race Against Ruth, Life, 18 ago 1961. Disponível em

https://books.google.com.br/books?id=oFQEAAAAMBAJ&pg=PA62&source=gbs_toc_r&redir_esc=y#v=onepage&q&f=false