

JOÃO PEDRO RICCI BOREJO
GABRIELA FONSECA FANUCCHI

DETERMINAÇÃO DE VIÉS DE DADOS
POLIÉDRICOS A PARTIR DE ANÁLISE
ESTATÍSTICA

São Paulo
2022

JOÃO PEDRO RICCI BOREJO
GABRIELA FONSECA FANUCCHI

**DETERMINAÇÃO DE VIÉS DE DADOS
POLIÉDRICOS A PARTIR DE ANÁLISE
ESTATÍSTICA**

Trabalho apresentado à Escola Politécnica
da Universidade de São Paulo para obtenção
do Título de Engenheiro Mecatrônico.

São Paulo
2022

JOÃO PEDRO RICCI BOREJO
GABRIELA FONSECA FANUCCHI

**DETERMINAÇÃO DE VIÉS DE DADOS
POLIÉDRICOS A PARTIR DE ANÁLISE
ESTATÍSTICA**

Trabalho apresentado à Escola Politécnica
da Universidade de São Paulo para obtenção
do Título de Engenheiro Mecatrônico.

Orientador:

Prof. Dr. Thiago C. Martins

São Paulo
2022

Dedicamos este trabalho às nossas famílias, em especial nossos pais, por terem nos fornecido a oportunidade, o apoio e o incentivo de buscarmos nossos sonhos e objetivos, mesmo frente circunstâncias adversas e com os mais variados obstáculos

AGRADECIMENTOS

Agradecimentos especiais ao professor Doutor Thiago Castro Martins, por toda a orientação e suporte ao longo dos últimos meses na concepção e criação do Projeto de Conclusão de Curso.

“Alea Jacta Est”

- Julius Caesar

RESUMO

O projeto tem como objetivos analisar a confiabilidade de dados poliédricos encontrados nos mercados mais comuns e utilizados em jogos de tabuleiros e para outros fins, por meio de uma ferramenta computacional que capture os lançamentos de dados poliédricos, em conjunto com um algoritmo que realize testes estatísticos de confiança da precisão do dado. Nesse sentido, o objetivo maior é o de fornecer uma ferramenta que audite a qualidade e precisão de um dado, testando a confiança dos produtos e confirmando se eles entregam aquilo que se propõe aos seus consumidores, consolidando demandas do dia a dia com os conteúdos observados no curso de engenharia.

Palavras-Chave – Dados, Estatística, Visão Computacional, Descritores, SIFT, Teste Chi-Quadrado, Teste de Kolmogorov-Smirnov.

ABSTRACT

The main purpose of this project is to analyze the reliability of poliedric dice sold in the core markets, used in board games and for other applications. Working tools include a computational program that captures and reads all the dice playing results and an algorithm that makes the statistics confidence tests to attribute a degree of precision to the dice. In that sense, the main goal is to provide a tool capable of auditing the quality and precision of a dice, testing the reliability of the product and confirming if they deliver what they promise to its consumers, attaching day to day demands with engineering knowledge.

Keywords – Dice, Statistics, Computer Vision, Descriptors, SIFT, Chi-Square Test, Kolmogorov-Smirnov Test.

LISTA DE FIGURAS

1	Kit de dados usado no projeto	11
2	Exemplo de protótipo amador	15
3	Exemplo de segundo protótipo amador	16
4	Distribuição Chi-Quadrado para diferentes graus de liberdade	20
5	Teste de Kolmogorov-Smirnov comparando duas amostras	21
6	Teste de Kolmogorov-Smirnov comparando a amostra obtida com os valores teóricos	22
7	Exemplo de imagem usada no projeto	30
8	Cadastro para o D6	31
9	Setup a partir de vista frontal	32
10	Setup a partir de vista superior	32
11	Teste livre com o modelo ORB	34
12	Teste livre com o Modelo SIFT	34
13	Teste realidade com o modelo ORB	35
14	Teste realidade com o Modelo SIFT	35
15	Comparação entre leituras	36
16	Distribuição normal para soma de resultados de um D6	38
17	Distribuição normal para os resultados obtidos	39
18	Tabela referência Chi-Quadrado	40
19	Distribuição uniforme esperada versus observada de um D6	41
20	Frequência acumulada esperada <i>vs</i> frequência acumulada observada de um D6	43

LISTA DE TABELAS

1	Tabela da distribuição Chi-Quadrado (χ^2)	19
2	Comparativo entre resultados teóricos e experimentais	22
3	Tabela de valores críticos para diferentes valores de α	24
4	Determinação aproximada do valor crítico de K	24
5	Função distribuição normal acumulada	27
6	Comparativo entre resultados teóricos e amostras	27
7	Matriz de confusão	37
8	Exemplo de distribuição normal perfeita - D6	38
9	Análise de chi-quadrado	39
10	Frequência esperada para o D6	41
11	Frequência observada para o D6	42
12	Diferença entre a frequência esperada acumulada e a frequência observada acumulada para o D6	42

SUMÁRIO

1	Introdução	11
1.1	Problema	11
1.2	Objetivo geral	12
1.3	Revisão da literatura	12
1.3.1	Reconhecimento de Imagem	12
1.3.2	Protótipos Mecânicos	15
1.3.3	Estatística aplicada aos requisitos do projeto	16
1.3.3.1	Teste chi-quadrado	18
1.3.3.2	Teste de aderência (Kolmogorov-Smirnov)	20
1.4	Requisitos	24
1.4.1	Requisitos teóricos	24
1.4.2	Determinação do número de lançamentos para cada tipo de dado	25
2	Materiais e Métodos	29
2.1	Elaboração da Base de Dados	29
2.1.1	Setup para realização dos testes	32
2.2	Uso das ferramentas de software Python e OpenCV	33
2.2.1	OpenCV	33
2.2.2	Descritores	33
3	Resultados e Discussão	37
3.1	Resultados do reconhecimento de imagem	37
3.2	Resultados dos testes estatísticos	37
3.2.1	Teste chi-quadrado	39

3.2.2	Teste de aderência (Kolmogorov-Smirnov)	40
3.2.3	Discussões e conclusão	43
4	Referências Bibliográficas	45
	Anexo A – Códigos Utilizados	47

1 INTRODUÇÃO

1.1 Problema

Os dados são pequenos poliedros gravados com determinadas instruções, ou, mais frequentemente, números em cada face. O dado mais clássico é o cubo (seis faces), mas no mundo do entretenimento moderno, em especial dos jogos de tabuleiro e outras modalidades, como os famosos RPGs, muitas vezes são utilizados dados de três, quatro, oito, doze ou até vinte faces.

Para artefatos como esses, existem tanto fabricantes em massa, quanto fabricantes de dados artesanais e sob encomenda. Em ambos os casos, é difícil auditar o processo de fabricação e testes desses poliedros, uma vez que, por questões concorrenciais, pouco se revela a respeito desses processos. Os vendedores de jogos muitas vezes confiam de maneira cega em seus fornecedores, por não haver um processo de teste regado e consistente. Quando os vendedores testam os produtos que adquirem para revender o fazem lançando os dados de forma manual e não se preocupando em registrar detalhadamente o número de eventos de cada face.

Ademais, ainda que o trabalho artesanal permita uma maior variedade e customização



Figura 1: Kit de dados usado no projeto

de produtos, além de dinamizar o mercado de jogos, torna-se mais desafiadora a padronização dos poliedros, por conta das peculiaridades e singularidades de cada produto fabricado. São muitos os fatores que podem desbalancear os poliedros, ainda que sutilmente, como por exemplo pinturas, artes e customizações que muitas vezes aparecem irregularmente pelas faces e podem causar uma distribuição não equivalente de massa.

Na Figura 1, é possível verificar que o conjunto de dados possui um processo de fabricação bastante específico e cada formato tem suas peculiaridades: diferentes números de chanfros, gravuras e cores, que podem causar um desequilíbrio de massa.

Nesse sentido, percebe-se a viabilidade e relevância de criar uma ferramenta para validar os resultados gerados lançamentos de cada dado, na linha de entender se ele apresenta (ou não) um viés, isto é, se os lados de cada fundo são equiprováveis ou não.

1.2 Objetivo geral

O projeto tem por objetivos propor e desenvolver uma ferramenta de lançamento de dados que capture os resultados, em conjunto com algoritmos de testes de confiança da precisão dos dados, de forma a fornecer uma ferramenta confiável da qualidade e precisão dos dados, assim garantindo produtos de confiança e que entreguem aquilo que se propõe.

1.3 Revisão da literatura

A literatura existente a respeito do estudo de dados enviesados se divide essencialmente em três partes: a literatura a respeito de inteligência artificial para tratar os dados, mais especificamente na área de visão computacional, e a literatura produzida com base em protótipos que foram feitos com o objetivo de fornecer uma solução automatizada que testasse o viés dos dados e a literatura mais clássica a respeito dos testes estatísticos a serem feitos com os dados. A seguir, vamos explorar os dois aspectos.

1.3.1 Reconhecimento de Imagem

A literatura existente na área de inteligência artificial e visão computacional utiliza principalmente *Deep Learning*, por meio de algoritmos de redes neurais. O método mais famoso é o das chamadas *Convolutional Neural Networks*, cujo racional aplicado é o treino do algoritmo com uma base de dados separada, de modo que o computador seja capaz de receber uma imagem como entrada, que passará por filtros atrelados a características da

imagem (as chamadas *features*), com diferentes pesos entre eles. Ao final desse processo, o algoritmo será capaz de gerar uma saída de modo a classificar a imagem recebida como entrada. Ainda, existem outros algoritmos mais sofisticados de *Deep Learning* que classificam de maneira semelhante, como a *Deep Belief Network (DBN)*, que é um modelo que trabalha a partir da geração de probabilidades com uma imagem de entrada.

Além dos métodos ora citados de reconhecimento de imagem, que são, por si só, mais sofisticados, existem outros métodos também amplamente explorados no meio acadêmico que são utilizados em aplicações similares. É o caso dos algoritmos ORB (*Oriented FAST and Rotated BRIEF*), que ganharam relevância a partir de 2011 por serem computacionalmente eficientes e terem uma performance notável em tempo real.

Os algoritmos ORB são essencialmente compostos por três etapas: (1) extração das *features*; (2) geração de descritores e (3) ligação ("*matching*") dos pontos. Na primeira etapa, a extração das *features* é feita a partir da análise dos *pixels* vizinhos a um determinado *pixel*. Caso sejam identificados muitos vizinhos significativamente diferentes - em aspectos como cor, luminosidade, formato, entre outros - do ponto analisado, o algoritmo classifica este ponto (*pixel*) como fronteira. A separação (ou identificação dos pontos diferentes) é feita com base em determinado *threshold*, que serve de parâmetro para os vizinhos. Por exemplo, caso o algoritmo esteja analisando um pixel p com uma luminosidade I_p e a luminosidade identificada em n (onde n é um número arbitrário dos vizinhos analisados) dos vizinhos de p analisados seja superior a $I_p + T$ ou inferior a $I_p - T$, onde T é o *threshold*, então p será considerado fronteiroço. Para concluir a extração das *features*, o algoritmo segue com a construção de pirâmides de escala das imagens e adiciona invariância de escala aos pontos. Para finalizar a primeira etapa, o algoritmo determina uma direção para o ponto. A direção é calculada a partir do vetor que liga o centro de massa (m_{00}) ao centróide (m_{10}, m_{01}) do ponto analisado. A direção do ponto é dada por:

$$\theta = \arctan\left(\frac{m_{01}}{m_{10}}\right)$$

Na segunda etapa, o algoritmo ORB gera descritores a partir do quanto calculado na etapa anterior, utilizando o método BRIEF. O BRIEF é um vetor binário descritor cuja função consiste de um número 0 ou 1:

$$\tau(p; x, y) = \begin{pmatrix} 1, p(x) < p(y); \\ 2, p(x) \geq p(y) \end{pmatrix}$$

Na equação acima, $p(x)$ é o valor que representa a luminosidade do ponto em x e $p(y)$ é o valor que representa a luminosidade do ponto em y .

Para reduzir o ruído, o algoritmo contém um filtro de Gauss. Seleciona-se o ponto p como central e - por exemplo, 256 - N pares de pontos vizinhos. A luminosidade de cada par de pontos é comparada utilizando a equação que determina a direção do ponto, ora apresentada. Dessa forma, o algoritmo constrói um vetor de N dimensões com N pares de pontos:

$$f_N(p) = \sum_{1 \leq i \leq N} 2^{i-1} \tau(p; x_i; y_i)$$

Para evitar perda de dados, a partir do descritor acima é gerado um descritor direcional:

$$g_{N(p,\theta)} = f_N(p) \in Q_\theta$$

Onde Q_θ é dado por:

$$Q_\theta = \begin{pmatrix} \cos\theta & \sin\theta \\ -\sin\theta & \cos\theta \end{pmatrix} \begin{pmatrix} x_1 \dots x_N \\ y_1 \dots y_N \end{pmatrix}$$

Por fim, a terceira e última etapa consiste em fazer as ligações (*matches*) dos pontos. Para tanto, o algoritmo mede a distância de determinado ponto aos descritores, e seleciona o que apresenta a menor distância.

Os algoritmos SIFT têm funcionamento muito semelhante ao algoritmo ORB. No entanto, seu grande diferencial é o fato de que, na etapa (2), os descritores gerados são invariantes à translação, rotação e dimensão das imagens. Possibilitando a geração de um maior número de *matches*.

O estudo dessa literatura foi essencial para a escolha dos métodos computacionais a serem usados, assim como o entendimento destes, possibilitando um embasamento teórico que levasse a uma programação mais fluída e melhor uso das ferramentas disponíveis.

A escolha de uso de descritores se deu, principalmente, pelo fato das propriedades de dados polidédricos. Cada fabricante, industrial ou artesanal, possui métodos diferentes de fabricação e desenho, o que acarreta em divergências muito grandes entre as faces de diversos dados, seja na forma como os números foram desenhados, ou em desenhos para

complementar o dado, esses fatores podem gerar uma confusão muito grande em um algoritmo baseado em machine learning, o que iria prejudicar a precisão do reconhecimento de imagem, prejudicando as análises do dado escolhido. No uso de descritores, essa imprecisão fica menor, como será mostrado mais a frente na seção de resultados, é possível obter números excelentes com esse método.

1.3.2 Protótipos Mecânicos

Quanto aos protótipos feitos com o objetivo de testar a acurácia dos poliedros, encontram-se alguns modelos feitos artesanalmente, em geral, com pouco método e sem qualquer tipo de propósito comercial ou industrial:

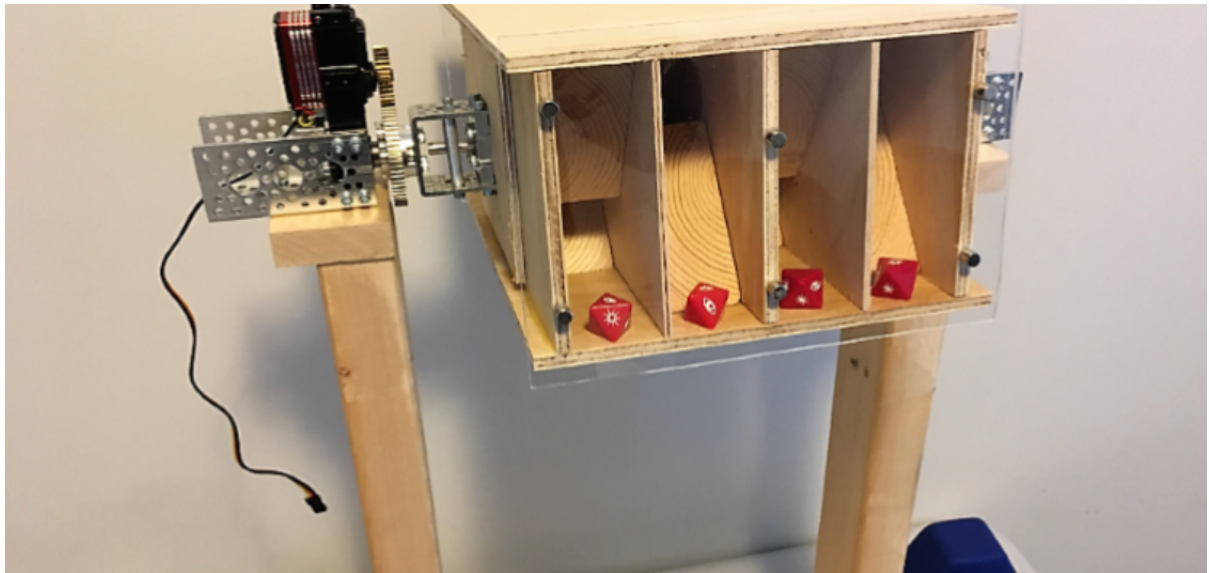


Figura 2: Exemplo de protótipo amador

O protótipo mostrado na Figura 2, além de coletar os dados (imagens) do resultado de cada lançamento, possui uma estrutura mecânica capaz de lançar o dado n vezes até obter o número de lançamentos desejado para realizar o teste de hipótese.

Mas, há outras formas de construir estas máquinas, desde algumas mais simples até formas mais sofisticadas, por exemplo, comparado ao projeto anterior, o projeto seguinte é mais sofisticado, cumprindo o mesmo papel mas de forma menos brusca:



Figura 3: Exemplo de segundo protótipo amador

Infelizmente não há uma literatura mais refinada quanto a projetos de lançamento de dados, afinal, a parte mecânica é consideravelmente trivial. No entanto, há uma vasta quantidade de projetos amadores que se apresentam como ferramentas interessantes de estudo para montagem de protótipos mecânico, sendo possível encontrar ferramentas que podem ser usadas para este.

1.3.3 Estatística aplicada aos requisitos do projeto

Um dos requisitos a ser estabelecido no projeto é o número de lançamentos de cada dado necessário para atribuir determinada significância ao resultado do teste de hipótese que será usado para determinar se os poliedros são enviesados ou não.

A revisão da literatura a respeito de estudos cujo objetivo é estimar uma proporção (podendo ser uma prevalência ou incidência), como no caso do projeto mostra que há

duas possibilidades para determinação do tamanho da amostra: caso em que se conhece o tamanho da população analisada e o caso em que se desconhece o tamanho da população analisada. Em ambos os cenários, três critérios principais são levados em consideração: (1) a proporção esperada para o desfecho de interesse; (2) o nível de confiança nos resultados e (3) o erro máximo de estimação, em pontos percentuais.

Para determinar o tamanho da amostra quando se conhece uma proporção, é utilizada a seguinte fórmula:

$$n = \frac{p(1-p)Z^2N}{\epsilon^2(N-1) + Z^2p(1-p)} \quad (1.1)$$

Onde:

n : tamanho da amostra;

p : proporção esperada;

Z : valor da distribuição normal para determinado nível de confiança;

N : tamanho da população;

ϵ : tamanho do intervalo de confiança (margem de erro).

Para determinar o tamanho da amostra quando se desconhece uma população, é utilizada a seguinte fórmula:

$$n = \frac{p(1-p)Z^2}{\epsilon^2} \quad (1.2)$$

Onde:

n : tamanho da amostra;

p : proporção esperada;

Z : valor da distribuição normal para determinado nível de confiança;

ϵ : tamanho do intervalo de confiança (margem de erro).

Nesse sentido, considerando que para fins do projeto se desconhece o tamanho da população, o mais conveniente é utilizar a segunda fórmula.

1.3.3.1 Teste chi-quadrado

Adicionalmente, para complementar a fundamentação teórica da análise das informações, podemos explorar o teste do qui-quadrado, também conhecido como teste dos mínimos quadrados. O teste do qui-quadrado de Pearson, também conhecido apenas como teste do qui-quadrado, é um teste não paramétrico, isto é, que não depende de parâmetros populacionais (como a média e a variância), cujo princípio básico é comparar proporções, ou seja, possíveis diferenças entre as frequências esperadas e observadas de determinada amostra.

Na prática, o teste é utilizado para verificar se a frequência de determinado acontecimento observada em uma amostra se desvia significativamente da frequência com que esse acontecimento é esperado.

Sua expressão matemática é dada por:

$$\chi^2 = \frac{1}{n} \sum_{i=0}^n (X_i - P(x_i))^2 \quad (1.3)$$

n : tamanho da amostra;

X_i : proporção observada;

$P(x_i)$: proporção esperada.

Com ela, é possível chegar na diferença entre os dados experimentais e os teóricos.

Uma vez obtido o valor do chi-quadrado para determinada amostra, ele é comparado ao valor teórico da tabela abaixo, que representa o gráfico monocaudal mostrado na figura subsequente. O valor teórico é determinado em função do número de graus de liberdade da amostra e da significância desejada à hipótese que se está testando.

ν	0.1%	0.5%	1.0%	2.5%	5.0%	10.0%	12.5%	20.0%	25.0%	33.3%	50.0%
1	0.000	0.000	0.000	0.001	0.004	0.016	0.025	0.064	0.102	0.186	0.455
2	0.002	0.010	0.020	0.051	0.103	0.211	0.267	0.446	0.575	0.811	1.386
3	0.024	0.072	0.115	0.216	0.352	0.584	0.692	1.005	1.213	1.568	2.366
4	0.091	0.207	0.297	0.484	0.711	1.064	1.219	1.649	1.923	2.378	3.357
5	0.210	0.412	0.554	0.831	1.145	1.610	1.808	2.343	2.675	3.216	4.351
6	0.381	0.676	0.872	1.237	1.635	2.204	2.441	3.070	3.455	4.074	5.348
7	0.598	0.989	1.239	1.690	2.167	2.833	3.106	3.822	4.255	4.945	6.346
8	0.857	1.344	1.646	2.180	2.733	3.490	3.797	4.594	5.071	5.826	7.344
9	1.152	1.735	2.088	2.700	3.325	4.168	4.507	5.380	5.899	6.716	8.343
10	1.479	2.156	2.558	3.247	3.940	4.865	5.234	6.179	6.737	7.612	9.342
11	1.834	2.603	3.053	3.816	4.575	5.578	5.975	6.989	7.584	8.514	10.341
12	2.214	3.074	3.571	4.404	5.226	6.304	6.729	7.807	8.438	9.420	11.340
13	2.617	3.565	4.107	5.009	5.892	7.042	7.493	8.634	9.299	10.331	12.340
14	3.041	4.075	4.660	5.629	6.571	7.790	8.266	9.467	10.165	11.245	13.339
15	3.483	4.601	5.229	6.262	7.261	8.547	9.048	10.307	11.037	12.163	14.339
16	3.942	5.142	5.812	6.908	7.962	9.312	9.837	11.152	11.912	13.083	15.338
17	4.416	5.697	6.408	7.564	8.672	10.085	10.633	12.002	12.792	14.006	16.338
18	4.905	6.265	7.015	8.231	9.390	10.865	11.435	12.857	13.675	14.931	17.338
19	5.407	6.844	7.633	8.907	10.117	11.651	12.242	13.716	14.562	15.859	18.338
20	5.921	7.434	8.260	9.591	10.851	12.443	13.055	14.578	15.452	16.788	19.337
21	6.447	8.034	8.897	10.283	11.591	13.240	13.873	15.445	16.344	17.720	20.337
22	6.983	8.643	9.542	10.982	12.338	14.041	14.695	16.314	17.240	18.653	21.337
23	7.529	9.260	10.196	11.689	13.091	14.848	15.521	17.187	18.137	19.587	22.337
24	8.085	9.886	10.856	12.401	13.848	15.659	16.351	18.062	19.037	20.523	23.337
25	8.649	10.520	11.524	13.120	14.611	16.473	17.184	18.940	19.939	21.461	24.337
26	9.222	11.160	12.198	13.844	15.379	17.292	18.021	19.820	20.843	22.399	25.336
27	9.803	11.808	12.879	14.573	16.151	18.114	18.861	20.703	21.749	23.339	26.336
28	10.391	12.461	13.565	15.308	16.928	18.939	19.704	21.588	22.657	24.280	27.336
29	10.986	13.121	14.256	16.047	17.708	19.768	20.550	22.475	23.567	25.222	28.336
30	11.588	13.787	14.953	16.791	18.493	20.599	21.399	23.364	24.478	26.165	29.336
40	17.916	20.707	22.164	24.433	26.509	29.051	30.008	32.345	33.660	35.643	39.335
50	24.674	27.991	29.707	32.357	34.764	37.689	38.785	41.449	42.942	45.184	49.335

Tabela 1: Tabela da distribuição Chi-Quadrado (χ^2)

Chi-Quadrado para diferentes graus de liberdade

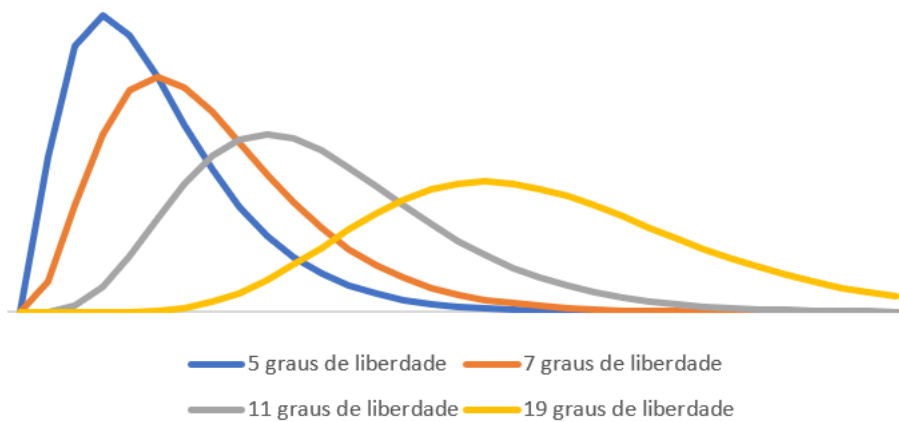


Figura 4: Distribuição Chi-Quadrado para diferentes graus de liberdade

No teste qui quadrado, assim como quando analisamos uma distribuição normal, quanto maior a amostra de dados utilizados no experimento, mais este se aproxima da teoria, o que implica em um longo processo de testagem do projeto de máquina pensado, de forma a ter os valores de amostragem mais confiáveis e precisos possíveis, para assim julgar os dados da melhor forma possível.

1.3.3.2 Teste de aderência (Kolmogorov-Smirnov)

O teste de Kolmogorov-Smirnov - também conhecido como teste de aderência, tem por objetivo principal comparar a frequência acumulada de determinada amostra com a frequência acumulada esperada para aqueles lançamentos, isto é, calcular a diferença entre uma distribuição conhecida e uma distribuição obtida empiricamente.

Diferentemente de outros testes estatísticos, como o *t de Student* e testes que utilizam tão somente a média, mediana ou o desvio padrão da amostra, o teste de Kolmogorov-Smirnov identifica a maior diferença ao longo de toda a distribuição de frequências de uma determinada amostra. Ainda, é possível utilizar o teste de Kolmogorov-Smirnov para comparar a frequência de duas amostras.

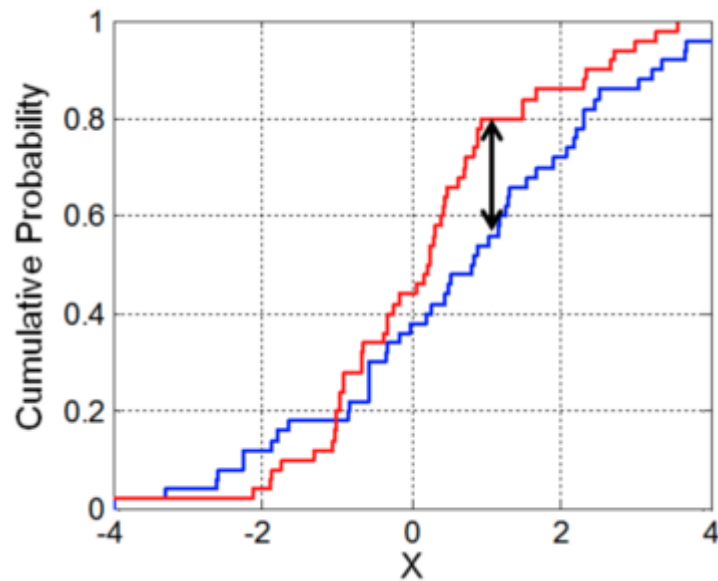


Figura 5: Teste de Kolmogorov-Smirnov comparando duas amostras

O K_{teste} é dado pela fórmula abaixo, que representa a maior diferença entre a frequência acumulada obtida e a frequência acumulada esperada (ou obtida novamente, em caso de comparação entre duas amostras).

$$K_{teste} = MAX(|A_i - O_i|) \quad (1.4)$$

i : tamanho da amostra;

A_i : frequência relativa acumulada teórica;

O_i : frequência relativa acumulada observada.

Para fins de ilustração, a seguir compara-se uma amostra arbitrária de frequência de números de 1 a 10 com a frequência esperada para uma distribuição normal teórica de média 4,5 e desvio padrão 1,5. A tabela ilustra os resultados obtidos, sendo que a quinta coluna é obtida de forma teórica e a quarta coluna é obtida a partir da amostra apresentada na segunda coluna.

Números	Frequência	Frequência Relativa	Oi	Ai	Ai-Oi
0	1	0.0164	0.0164	0.0013	0.0150
1	3	0.0492	0.0656	0.0098	0.0558
2	4	0.0656	0.1311	0.0478	0.0834
3	8	0.1311	0.2623	0.1587	0.1036
4	24	0.3934	0.6557	0.3694	0.2836
5	12	0.1967	0.8525	0.6306	0.2219
6	5	0.0820	0.9344	0.8413	0.0931
7	3	0.0492	0.9836	0.9522	0.0314
8	0	0.0000	0.9836	0.9902	0.0066
9	1	0.0164	1.0000	0.9987	0.0013
10	0	0.0000	1.0000	0.9999	0.0001

Tabela 2: Comparativo entre resultados teóricos e experimentais

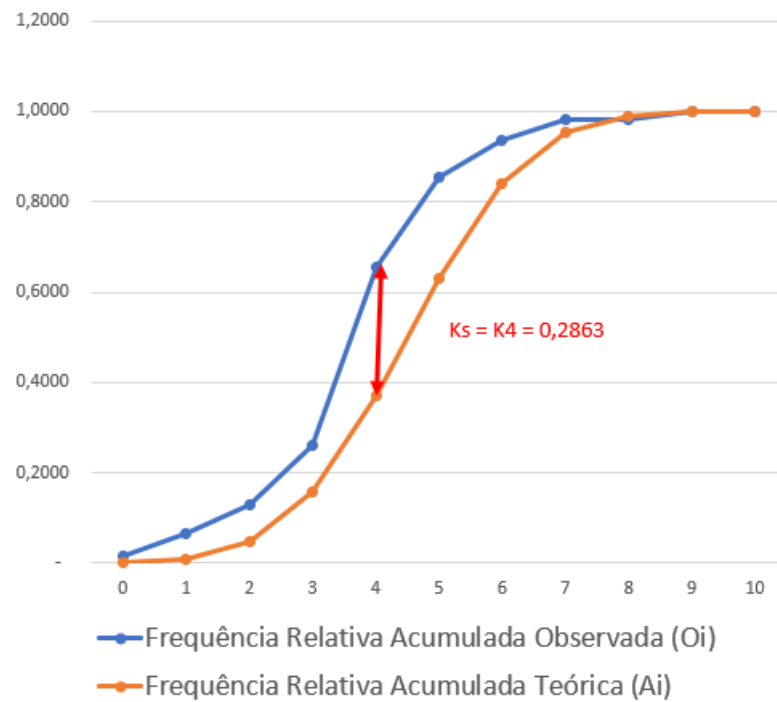


Figura 6: Teste de Kolmogorov-Smirnov comparando a amostra obtida com os valores teóricos

Para o caso exemplificado, o K_{teste} é de 0.2863 para o número 4, que representa a maior distância entre a curva das frequências acumuladas teóricas e a curva das frequências acumuladas da amostra.

Para saber se a frequência da amostra está ou não aderente à frequência esperada,

compara-se K_{teste} com o respectivo valor crítico. O valor crítico de K pode ser obtido por meio da tabela mostrada abaixo.

Significância(α)	0.1	0.05	0.025	0.01	0.005
1	0.9000	0.9500	0.9750	0.9900	0.9950
2	0.6838	0.7764	0.8419	0.9000	0.9293
3	0.5648	0.6360	0.7076	0.7846	0.8290
4	0.4927	0.5652	0.6239	0.6889	0.7342
5	0.4470	0.5094	0.5633	0.6272	0.6685
6	0.4104	0.4680	0.5193	0.5774	0.6166
7	0.3815	0.4361	0.4834	0.5384	0.5758
8	0.3583	0.4096	0.4543	0.5065	0.5418
9	0.3391	0.3875	0.4300	0.4796	0.5133
10	0.3226	0.3687	0.4092	0.4566	0.4889
11	0.3083	0.3524	0.3912	0.4367	0.4677
12	0.2958	0.3382	0.3754	0.4192	0.4490
13	0.2847	0.3255	0.3614	0.4036	0.4325
14	0.2748	0.3142	0.3489	0.3897	0.4176
15	0.2659	0.3040	0.3376	0.3771	0.4042
16	0.2578	0.2947	0.3273	0.3657	0.3920
17	0.2504	0.2863	0.3180	0.3553	0.3809
18	0.2436	0.2785	0.3094	0.3457	0.3706
19	0.2373	0.2714	0.3014	0.3369	0.3612
20	0.2316	0.2647	0.2941	0.3287	0.3524
21	0.2262	0.2586	0.2872	0.3210	0.3443
22	0.2212	0.2528	0.2809	0.3139	0.3367
23	0.2165	0.2475	0.2749	0.3073	0.3295
24	0.2120	0.2424	0.2693	0.3010	0.3229
25	0.2079	0.2377	0.2640	0.2952	0.3166
26	0.2040	0.2332	0.2591	0.2896	0.3106
27	0.2003	0.2290	0.2544	0.2844	0.3050
28	0.1968	0.2250	0.2499	0.2794	0.2997
29	0.1935	0.2212	0.2457	0.2747	0.2947
30	0.1903	0.2176	0.2417	0.2702	0.2899
31	0.1873	0.2141	0.2379	0.2660	0.2853
32	0.1844	0.2108	0.2342	0.2619	0.2809

Significância(α)	0.1	0.05	0.025	0.01	0.005
33	0.1817	0.2077	0.2308	0.2580	0.2768
34	0.1791	0.2047	0.2274	0.2543	0.2728
35	0.1766	0.2018	0.2242	0.2507	0.2690
36	0.1742	0.1991	0.2212	0.2473	0.2653
37	0.1719	0.1965	0.2183	0.2440	0.2618
38	0.1697	0.1939	0.2154	0.2409	0.2584
39	0.1675	0.1915	0.2127	0.2379	0.2552
40	0.1655	0.1891	0.2101	0.2349	0.2521

Tabela 3: Tabela de valores críticos para diferentes valores de α

No caso de valores maiores que 40, é possível determinar o valor crítico por meio de expressões em função do tamanho da amostra e que variam de acordo com o nível de significância desejado, conforme tabela mostrada abaixo.

Significância (α)	0.20	0.10	0.05	0.02	0.01
Fórmula	$\frac{1.07}{\sqrt{n}}$	$\frac{1.22}{\sqrt{n}}$	$\frac{1.36}{\sqrt{n}}$	$\frac{1.52}{\sqrt{n}}$	$\frac{1.63}{\sqrt{n}}$

Tabela 4: Determinação aproximada do valor crítico de K

1.4 Requisitos

1.4.1 Requisitos teóricos

Para iniciar o projeto, deve-se definir quais são seus requisitos, isto é, a quais critérios ele deve atender de forma a deixar o seu escopo e objetivos claros.

O objetivo principal do projeto é suprir uma carência de testagem de dados artesanais no mundo de RPG, de forma a possibilitar a averiguação da qualidade do produto, assim como sua precisão. Para isso, a ideia principal é definir a probabilidade de, ao se testar o dado, ele apresentar algum tipo de viés em seus resultados. Com isso em mente, foi pensado o teste de hipótese para testagem do dado e definição de seus parâmetros.

O teste de hipótese permite decidir se uma afirmação sobre um parâmetro populacional é verdadeira, por meio de dados amostrais. Para exemplificar o funcionamento, a seguir tem-se uma exemplificação de resultados para 30 lançamentos de dados:

- Média: 3,458

- Mediana: 3
- Desvio padrão: 1,641

Nesse caso, assume-se que nossa hipótese é que, para um dado de 6 lados não viciado, a média de resultados de seus lançamentos deveria ser 3,5, considerando-se infinitas rolagens, e deseja-se testar essa hipótese com 2% de significância. Assim, podemos usar o teste *t-student* (utilizado quando não se conhece a variância da população, apenas da amostra) para analisar o problema:

$$t = \frac{\bar{x} - u}{s/\sqrt{n}} \quad (1.5)$$

Substituindo, temos:

$$t = \frac{3,458 - 3,5}{1,641^2/\sqrt{30}} = -0,00467 \quad (1.6)$$

Para uma significância de 2% (1% na tabela de *t-student*, já que a distribuição é bicaudal), resulta-se em $z = -2,33$, indicando que o dado não é viciado com uma confiança de 98%.

Nesse caso, o dado usado provavelmente passou por um processos de testes e teve uma fabricação de boa qualidade, uma vez que era um dado feito de forma industrial, diferente de dados artesanais, feitos a mão e sem processos de testagem, assim, testá-los de forma extensa será essencial para nossa máquina, de forma a deixar o teste mais preciso e diminuir os possíveis erros de média e desvio padrão amostrais.

Com isso, é possível definir qual a probabilidade de um dado viciado passar no teste, sendo que o nível de significância é escolhido de forma arbitrária pelos usuários, mas, quanto maior a significância, maior a chance do dado ser pouco viciado.

1.4.2 Determinação do número de lançamentos para cada tipo de dado

Para determinar o número de lançamentos a ser feito para cada tipo de dado no âmbito dos testes estatísticos realizados neste projeto, foi utilizada a fórmula 1.1 apresentada na seção 1.3.3. da presente monografia, de modo que tratam-se de populações desconhecidas.

Para testar a acurácia de cada dado, foi arbitrado um intervalo de confiança de 95%. Nesse sentido, para realizar os testes utilizando a formula citada, considerando uma dis-

tribuição normal bicaudal, os valores de Z e ϵ , de acordo com a tabela abaixo, foram determinados em 1.96 e 0.05.

x	0.00	0.01	0.02	0.03	0.04	0.05	0.06	0.07	0.08	0.09
0.0	0.5000	0.5040	0.5080	0.5120	0.5160	0.5199	0.5239	0.5279	0.5319	0.5359
0.1	0.5398	0.5438	0.5478	0.5517	0.5557	0.5596	0.5636	0.5675	0.5714	0.5753
0.2	0.5793	0.5832	0.5871	0.5910	0.5948	0.5987	0.6026	0.6064	0.6103	0.6141
0.3	0.6179	0.6217	0.6255	0.6293	0.6331	0.6368	0.6406	0.6443	0.6480	0.6517
0.4	0.6554	0.6591	0.6628	0.6664	0.6700	0.6736	0.6772	0.6808	0.6844	0.6879
0.5	0.6915	0.6950	0.6985	0.7019	0.7054	0.7088	0.7123	0.7157	0.7190	0.7224
0.6	0.7257	0.7291	0.7324	0.7357	0.7389	0.7422	0.7454	0.7486	0.7517	0.7549
0.7	0.7580	0.7611	0.7642	0.7673	0.7703	0.7734	0.7764	0.7794	0.7823	0.7852
0.8	0.7881	0.7910	0.7939	0.7967	0.7995	0.8023	0.8051	0.8078	0.8106	0.8133
0.9	0.8159	0.8186	0.8212	0.8238	0.8264	0.8289	0.8315	0.8340	0.8365	0.8389
1.0	0.8413	0.8438	0.8461	0.8485	0.8508	0.8531	0.8554	0.8577	0.8599	0.8621
1.1	0.8643	0.8665	0.8686	0.8708	0.8729	0.8749	0.8770	0.8790	0.8810	0.8830
1.2	0.8849	0.8869	0.8888	0.8907	0.8925	0.8944	0.8962	0.8980	0.8997	0.9015
1.3	0.9032	0.9049	0.9066	0.9082	0.9099	0.9115	0.9131	0.9147	0.9162	0.9177
1.4	0.9192	0.9207	0.9222	0.9236	0.9251	0.9265	0.9279	0.9292	0.9306	0.9319
1.5	0.9332	0.9345	0.9357	0.9370	0.9382	0.9394	0.9406	0.9418	0.9429	0.9441
1.6	0.9452	0.9463	0.9474	0.9484	0.9495	0.9505	0.9515	0.9525	0.9535	0.9545
1.7	0.9554	0.9564	0.9573	0.9582	0.9591	0.9599	0.9608	0.9616	0.9625	0.9633
1.8	0.9641	0.9649	0.9656	0.9664	0.9671	0.9678	0.9686	0.9693	0.9699	0.9706
1.9	0.9713	0.9719	0.9726	0.9732	0.9738	0.9744	0.9750	0.9756	0.9761	0.9767
2.0	0.9772	0.9778	0.9783	0.9788	0.9793	0.9798	0.9803	0.9808	0.9812	0.9817
2.1	0.9821	0.9826	0.9830	0.9834	0.9838	0.9842	0.9846	0.9850	0.9854	0.9857
2.2	0.9861	0.9864	0.9868	0.9871	0.9875	0.9878	0.9881	0.9884	0.9887	0.9890
2.3	0.9893	0.9896	0.9898	0.9901	0.9904	0.9906	0.9909	0.9911	0.9913	0.9916
2.4	0.9918	0.9920	0.9922	0.9925	0.9927	0.9929	0.9931	0.9932	0.9934	0.9936
2.5	0.9938	0.9940	0.9941	0.9943	0.9945	0.9946	0.9948	0.9949	0.9951	0.9952

x	0.00	0.01	0.02	0.03	0.04	0.05	0.06	0.07	0.08	0.09
2.6	0.9953	0.9955	0.9956	0.9957	0.9959	0.9960	0.9961	0.9962	0.9963	0.9964
2.7	0.9965	0.9966	0.9967	0.9968	0.9969	0.9970	0.9971	0.9972	0.9973	0.9974
2.8	0.9974	0.9975	0.9976	0.9977	0.9977	0.9978	0.9979	0.9979	0.9980	0.9981
2.9	0.9981	0.9982	0.9982	0.9983	0.9984	0.9984	0.9985	0.9985	0.9986	0.9986
3.0	0.9987	0.9987	0.9987	0.9988	0.9988	0.9989	0.9989	0.9989	0.9990	0.9990
3.1	0.9990	0.9991	0.9991	0.9991	0.9992	0.9992	0.9992	0.9992	0.9993	0.9993
3.2	0.9993	0.9993	0.9994	0.9994	0.9994	0.9994	0.9994	0.9995	0.9995	0.9995
3.3	0.9995	0.9995	0.9995	0.9996	0.9996	0.9996	0.9996	0.9996	0.9996	0.9997
3.4	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9997	0.9998
3.5	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998	0.9998
3.6	0.9998	0.9998	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999
3.7	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999
3.8	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999	0.9999
3.9	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000	1.0000

Tabela 5: Função distribuição normal acumulada

Em seguida, substituindo os valores da fórmula para dados com 6, 8, 12 e 20 faces, foram obtidos os resultados apresentados na tabela a seguir.

Dado	p	n	μ	σ	Número de Lançamentos
D6	0.1667	214	3.5000	1.8708	1284
D8	0.1250	169	4.5000	2.4495	1352
D12	0.0833	118	6.5000	3.6056	1416
D20	0.0500	73	10.5000	5.9161	1460

Tabela 6: Comparativo entre resultados teóricos e amostras

n : tamanho da amostra;

p : proporção esperada;

Z : valor da distribuição normal para determinado nível de confiança;

ϵ : tamanho do intervalo de confiança (margem de erro);

μ : média;

σ : variância.

A última coluna mostra o resultado da multiplicação do valor de n obtido para cada dado diferente, de modo que obtém-se o número mínimo de lançamentos necessário para realizar os testes estatísticos de acurácia para o respectivo dado com 5% de confiança.

2 MATERIAIS E MÉTODOS

2.1 Elaboração da Base de Dados

A elaboração da base de dados foi realizada usando uma WebCam HD da Logitech C920s, presa a um monitor de forma fixa para gerar estabilidade, folhas de papel sulfite para geração de um fundo branco, e um conjunto de dados poliédricos.

Primeiramente a câmera foi presa ao suporte, de forma que estivesse estável e todas as fotos estivessem em um padrão. Em seguida, foi usada uma folha de papel sulfite para determinar onde o dado deveria ser lançado, de forma que ele sempre ficasse o mais centralizado possível, e houvesse a máxima definição da câmera. Em seguida, foram realizados os lançamentos, de acordo com o número necessário calculado. Com isso, foram tiradas as fotos dos lançamentos até chegar nos valores conhecidos.

Em seguida, foi utilizado um script para renomear todos os arquivos, de forma que fosse fácil checar a quantidade e procurar possíveis fotos de baixa qualidade com maior facilidade. Além disso, outro script realizava um corte nas imagens, focando apenas na região da folha sulfite, de forma que possíveis imprecisões fossem retiradas, melhorando o reconhecimento de imagem. A seguir, na figura 7, temos um exemplo de uma das imagens de lançamentos que foram usadas:



Figura 7: Exemplo de imagem usada no projeto

Além dessa base de dados, há a necessidade de realizar um cadastro de cada face, de forma que haja algo para se comparar com cada imagem para encontrar o lado no qual o dado caiu. Para aumentar a precisão, diminuindo o impacto de reflexos de iluminação, foram usados 4 fotos de cada face do dado. Dessa forma, foi possível criar cadastros de alta qualidade, a seguir, temos o exemplo dos cadastros do dado D6:

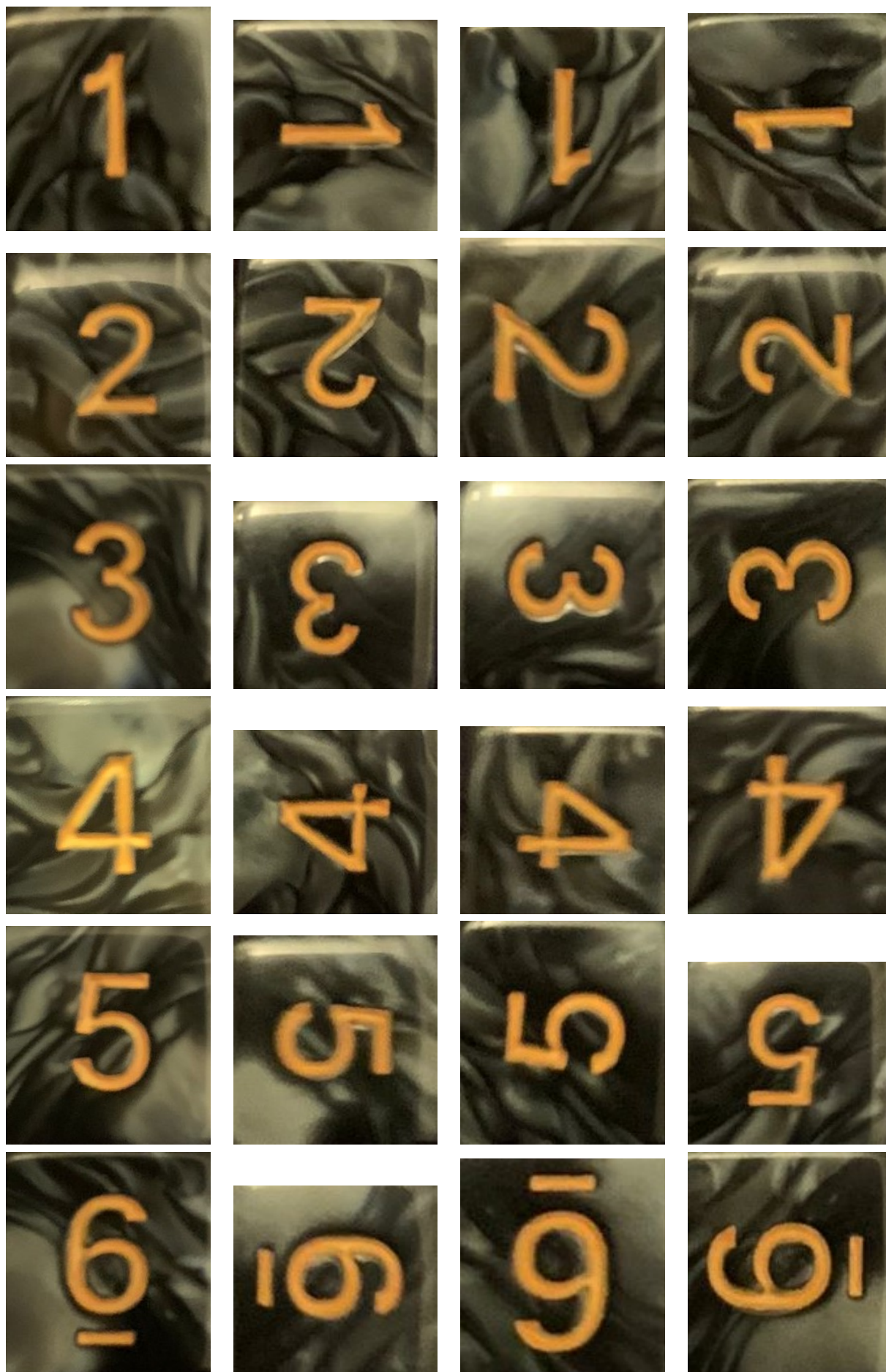


Figura 8: Cadastro para o D6

2.1.1 Setup para realização dos testes

Para realização dos testes, foi montado um *setup* com a câmera acoplado ao computador, buscando a melhor iluminação possível. Ainda, buscou-se colocar um fundo branco atrás dos dados, que eram escuros, para realizar a captura das imagens com um contraste adequado ao reconhecimento de imagem. Por fim, a iluminação indireta porém suficiente para visualização foi essencial para obtenção de boas imagens, de fácil visualização e sem reflexo de luz.



Figura 9: Setup a partir de vista frontal



Figura 10: Setup a partir de vista superior

2.2 Uso das ferramentas de software Python e OpenCV

Para realizar o tratamento da base de dados e ler as informações coletadas com a imagem da câmera, foi utilizado o Python e sua biblioteca OpenCV. Ao longo do processo, verificou-se a necessidade de incluir outras ferramentas de modo a refinar a precisão dos dados coletados, de modo que foi se aprimorando a base de cadastro por meio de ajustes na iluminação e foco das imagens capturadas. Mais adiante, serão explicados o funcionamento de cada ferramenta utilizada.

2.2.1 OpenCV

A biblioteca OpenCV é uma biblioteca *open source* com diversas ferramentas de visão computacional, dentre elas funções que permitem o fácil uso dos descritores por meio dos métodos SIFT, SURF, e ORB, conforme detalhado na seção a seguir. Com ela, além de ocorrer o processamento e leitura das imagens coletadas, ocorre também a identificação dos números, pelo processo descrito em seguida.

2.2.2 Descritores

Os descritores, conhecidos como *descriptors* na língua inglesa, são uma poderosa ferramenta computacional, constituídos de estruturas de informações que descrevem dados computacionais. Para este projeto, os descritores visuais são aplicáveis na construção da solução. Os descritores visuais buscam explicar matematicamente uma imagem, marcando nela os pontos de destaque como bordas, sombras, e pontos com características interessantes para análise dos dados.

Existem dois tipos de descritores que são mais amplamente utilizados: o SIFT e ORB. Os chamados SIFT (*Scale-invariant feature transform*) são conhecidos como os descritores mais poderosos que existem, pois são capazes de realizar aproximações de imagens mesmo que com diversas diferenças entre elas, como ângulos e iluminação, e até novos elementos, obtendo uma precisão espetacular. Por outro lado, os ORBs (*Oriented FAST and Rotated BRIEF*), são mais simples e triviais em seu uso e não possuem toda a sofisticação dos SIFT, o que os torna menos eficiente na comparação de imagens.

Foram montados 2 códigos, um utilizando o método SIFT e outro o método ORB, de forma a termos uma medida de comparação entre esses métodos:

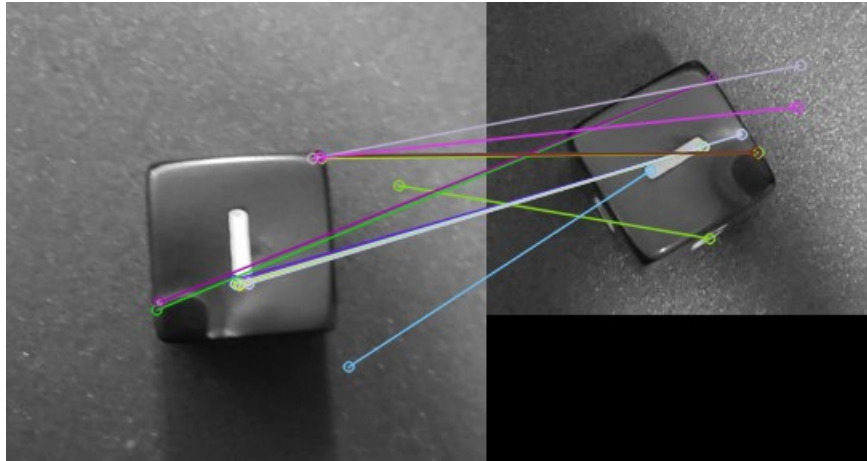


Figura 11: Teste livre com o modelo ORB

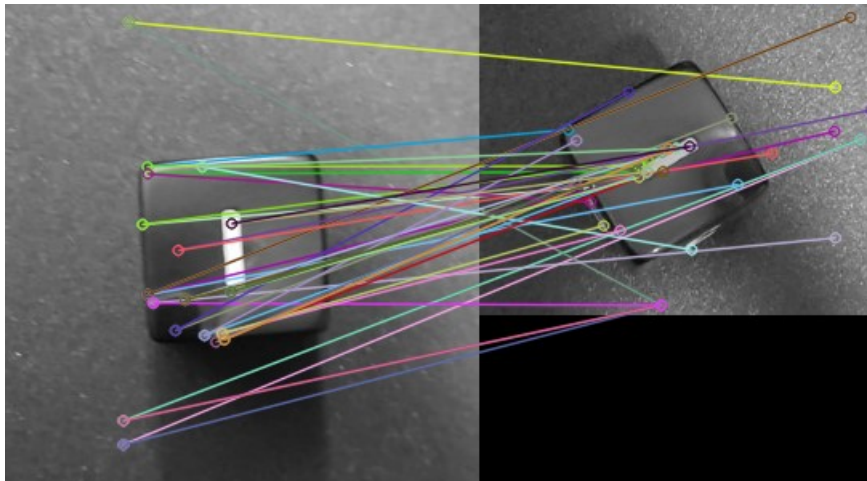


Figura 12: Teste livre com o Modelo SIFT

Ainda que não seja 100% preciso, o modelo SIFT claramente apresenta um índice de acertos consideravelmente maior que o modelo ORB quando as condições de ambiente são descontroladas e estamos comparando fotos com elementos diversos.

No entanto, um teste interessante era simular o que seria feito no projeto proposto. O planejamento foi montado de forma que, deveria ser escolhido um dado, e então os lados dele seriam "cadastrados", isto é, haveria uma foto para cada lado, e essas fotos seriam comparadas com os lançamentos. A foto que apresentasse maior semelhança em relação aos descritores (existem alguns métodos para essa identificação, que serão explorados a seguir), seria o valor do lançamento do dado. Com isso, ao simular algo parecido com a futura realidade do projeto, foram obtidos os seguintes resultados:

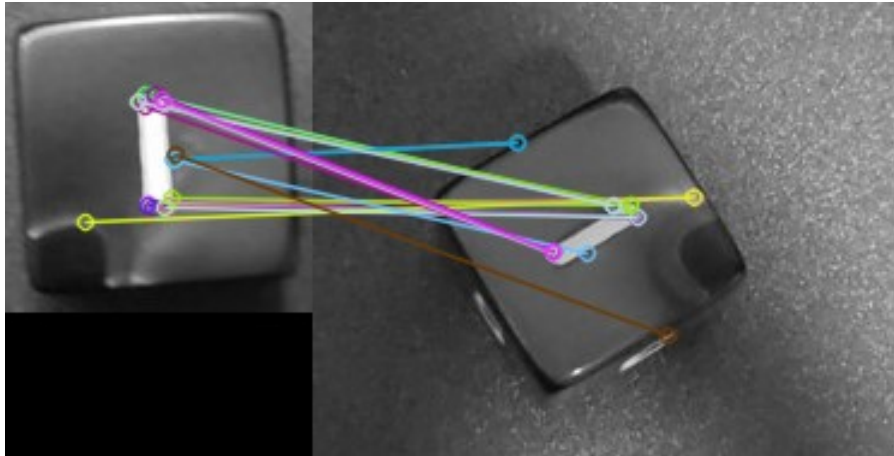


Figura 13: Teste realidade com o modelo ORB

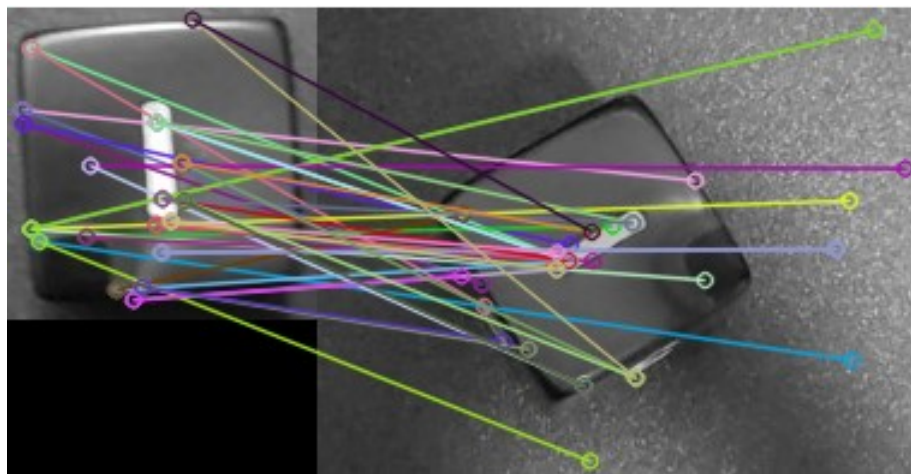


Figura 14: Teste realidade com o Modelo SIFT

A ideia principal é definir uma comparação de distância entre os *matches* nas duas imagens. Caso a comparação entre os dois esteja dentro de uma margem de erro definida pelo usuário ela será considerada boa, e podem-se contar o número de bons *matches* em um par de imagens. A imagem correspondente ao número que tiver maior quantidade de *matches* revelará o número encontrado. O principal método utilizado será este, a despeito de outros métodos também disponíveis para a identificação das imagens. A série de testes apresentada a seguir promoveu a identificação de pontos importantes para o projeto:

- Necessidade de um material liso e de poucos detalhes para a superfície do equipamento, de forma a ter o menor impacto possível na análise a partir dos descritores;
- Uso de uma iluminação consistente e de boa qualidade. Em alguns testes o ângulo e a potência da luz na imagem capturada geraram ruído nos resultados;

- Escolha de uma câmera de alta qualidade, uma vez que quanto maior a qualidade das fotos percebeu-se uma melhora nos *matches* dos descritores, tanto em qualidade quanto em precisão.

Com isso, foi montado um ambiente controlado para a realização dos lançamentos para a montagem da base de dados, no qual a luz era ideal, assim como a câmera usada, que possuía uma resolução altíssima, e o fundo tinha a menor quantidade de detalhes para evitar que houvesse interferência nos resultados encontrados.

A seguir, há exemplos de leitura de uma imagem diretamente pelo método SIFT, que se apresentou mais eficaz, com os 6 lados do D6 comparados com um lançamento de dado de valor 5, no qual percebe-se como houve uma diferença com o ajuste de ambiente e com um refino do código:

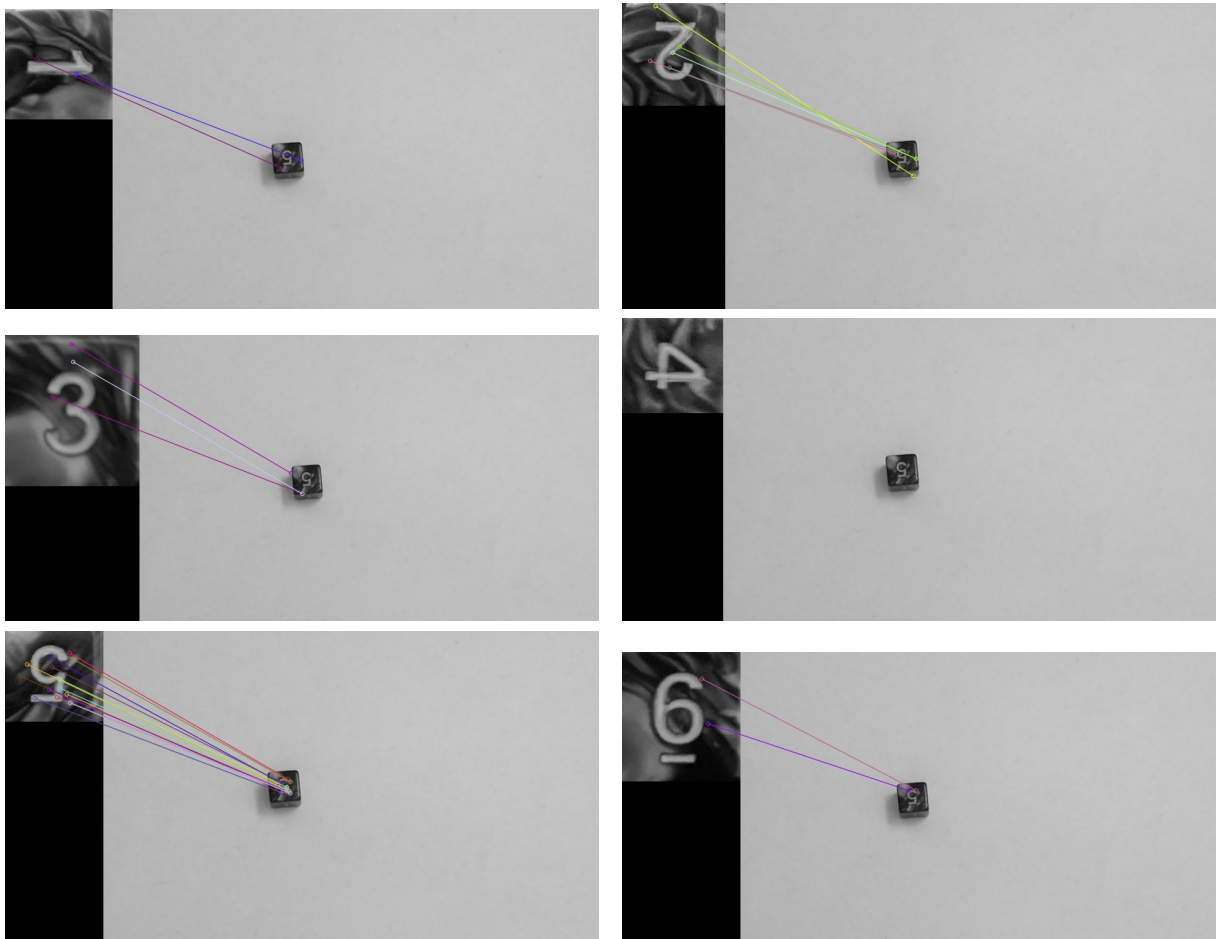


Figura 15: Comparação entre leituras

Com as imagens anteriores, é possível observar a potência do algoritmo SIFT, assim como sua precisão, pela diferença de matches entre as diferentes faces, como na face 4, que não teve nenhum match considerado satisfatório pela regra de Lowe's.

3 RESULTADOS E DISCUSSÃO

3.1 Resultados do reconhecimento de imagem

Para se aprofundar mais no estudo de descritores, foram realizados alguns testes com lançamentos manuais de dados, foram 200 lançamentos de um D6, realizados de forma manual, que foram tabelados e fotografados ao serem lançados, e, em seguida, analisados por um *script* que se usava de descritores para dizer qual o número do dado. A seguir temos a matriz de confusão dos resultados obtidos, nas colunas temos os resultados esperados, enquanto nas linhas os resultados obtidos:

	1	2	3	4	5	6
1	34	0	1	0	0	4
2	0	33	1	0	0	0
3	1	0	31	0	0	1
4	0	0	0	32	0	0
5	0	0	0	0	37	0
6	0	0	0	0	1	24

Tabela 7: Matriz de confusão

Dos 200 testes realizados, foi obtida uma precisão de 95,5%, isto é, 191 testes tiveram o resultado correto, enquanto 9 apresentaram erros. Essa precisão é extremamente alta, gerando uma segurança alta em relação aos resultados obtidos, além de mostrar que o algoritmo é extremamente confiável no caso de um dado D6, gerando um ruído pequeno nos resultados das análises estatísticas, tornando-as confiáveis.

3.2 Resultados dos testes estatísticos

Para realizar os testes estatísticos que nos permitem avaliar a precisão do dado, foram realizados 1483 lançamentos para um D6, número superior ao calculado na seção 1.4.2. e

mostrado na tabela 2 (1284 para o D6), equivalente ao número de lançamentos n mínimo para esse dado.

Para testar a confiabilidade do dado, foram montados algoritmos que realizam os testes de Chi-quadrado e Kolmogorov-Smirnov (aderência) com os resultados obtidos na leitura dos 1483 lançamentos, realizada pelo algoritmo de visão computacional.

Há alguns gráficos que devem ser observados, começando por uma comparação dos resultados obtidos com os resultados esperados:

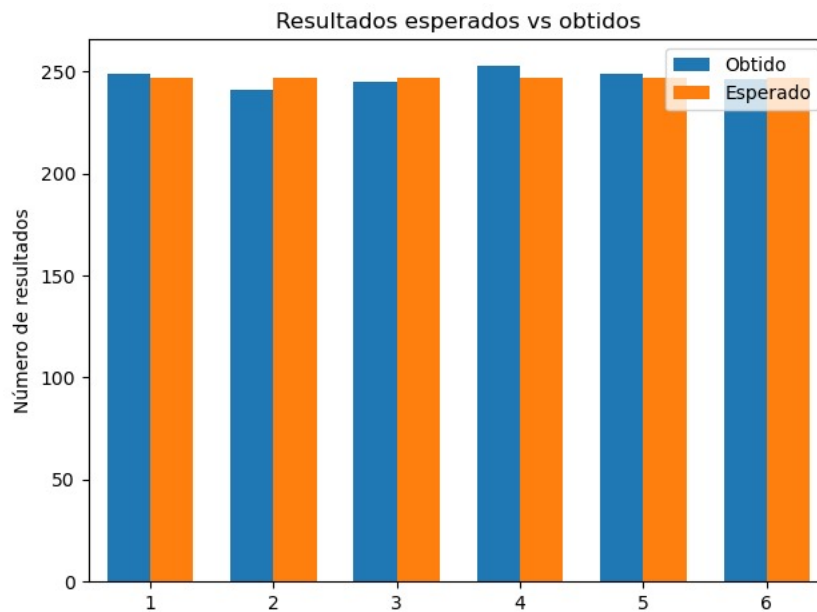


Figura 16: Distribuição normal para soma de resultados de um D6

Pelo gráfico apresentado, é possível observar que os resultados apresentam divergências mínimas em relação ao esperado.

Além disso, uma análise interessante é realizar uma soma de resultados do dado, de 2 em 2, isto é, o 1º lançamento com o 2º lançamento, o 3º lançamento com o 4º lançamento, e assim por diante, de forma que uma curva normal seja formada, uma vez que ao somar 2 lançamentos de dado, forma-se uma distribuição normal, seguindo as seguintes probabilidades:

2	3	4	5	6	7	8	9	10	11	12
2,8%	5,6%	8,3%	11,1%	13,9%	16,7%	13,9%	11,1%	8,3%	5,6%	2,8%

Tabela 8: Exemplo de distribuição normal perfeita - D6

Para o teste realizado, foi montada uma curva normal, que representou o seguinte

resultado:

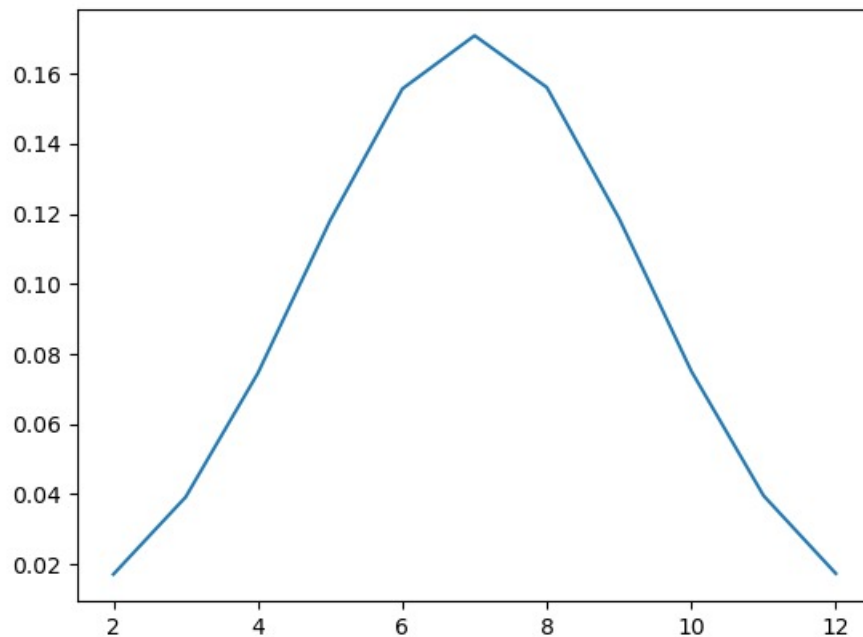


Figura 17: Distribuição normal para os resultados obtidos

Observa-se que a curva normal comporta-se de forma lógica, de acordo com o esperado, sem divergências que indiquem problemas ou incoerências nas amostras.

3.2.1 Teste chi-quadrado

O teste de chi-quadrado é essencial para realizar uma comparação entre o resultado esperado e o resultado obtido, de forma a validar se o dado escolhido é viciado, dentro de um intervalo de confiança. Assim, gera-se a seguinte tabela:

Valor	Observado	Esperado	$(\text{Obs}-\text{Esp})^2 / \text{Esp}$
1	249	248	0,0136
2	241	248	0,1539
3	245	248	0,0190
4	253	248	0,1377
5	249	248	0,0136
6	246	248	0,006

Tabela 9: Análise de chi-quadrado

Com isso, obtem-se um valor de X^2 de 0,343.

Agora, analisando uma tabela de distribuição de chi-quadrado, como a seguinte:

Degrees of Freedom	Chi-Square (χ^2) Distribution									
	Area to the Right of Critical Value									
	0.995	0.99	0.975	0.95	0.90	0.10	0.05	0.025	0.01	0.005
1	—	—	0.001	0.004	0.016	2.706	3.841	5.024	6.635	7.879
2	0.010	0.020	0.051	0.103	0.211	4.605	5.991	7.378	9.210	10.597
3	0.072	0.115	0.216	0.352	0.584	6.251	7.815	9.348	11.345	12.838
4	0.207	0.297	0.484	0.711	1.064	7.779	9.488	11.143	13.277	14.860
5	0.412	0.554	0.831	1.145	1.610	9.236	11.071	12.833	15.086	16.750
6	0.676	0.872	1.237	1.635	2.204	10.645	12.592	14.449	16.812	18.548
7	0.989	1.239	1.690	2.167	2.833	12.017	14.067	16.013	18.475	20.278
8	1.344	1.646	2.180	2.733	3.490	13.362	15.507	17.535	20.090	21.955
9	1.735	2.088	2.700	3.325	4.168	14.684	16.919	19.023	21.666	23.589
10	2.156	2.558	3.247	3.940	4.865	15.987	18.307	20.483	23.209	25.188
11	2.603	3.053	3.816	4.575	5.578	17.275	19.675	21.920	24.725	26.757
12	3.074	3.571	4.404	5.226	6.304	18.549	21.026	23.337	26.217	28.299
13	3.565	4.107	5.009	5.892	7.042	19.812	22.362	24.736	27.688	29.819
14	4.075	4.660	5.629	6.571	7.790	21.064	23.685	26.119	29.141	31.319
15	4.601	5.229	6.262	7.261	8.547	22.307	24.996	27.488	30.578	32.801
16	5.142	5.812	6.908	7.962	9.312	23.542	26.296	28.845	32.000	34.267
17	5.697	6.408	7.564	8.672	10.085	24.769	27.587	30.191	33.409	35.718
18	6.265	7.015	8.231	9.390	10.865	25.989	28.869	31.526	34.805	37.156
19	6.844	7.633	8.907	10.117	11.651	27.204	30.144	32.852	36.191	38.582
20	7.434	8.260	9.591	10.851	12.443	28.412	31.410	34.170	37.566	39.997
21	8.034	8.897	10.283	11.591	13.240	29.615	32.671	35.479	38.932	41.401
22	8.643	9.542	10.982	12.338	14.042	30.813	33.924	36.781	40.289	42.796
23	9.260	10.196	11.689	13.091	14.848	32.007	35.172	38.076	41.638	44.181
24	9.886	10.856	12.401	13.848	15.659	33.196	36.415	39.364	42.980	45.559
25	10.520	11.524	13.120	14.611	16.473	34.382	37.652	40.646	44.314	46.928
26	11.160	12.198	13.844	15.379	17.292	35.563	38.885	41.923	45.642	48.290
27	11.808	12.879	14.573	16.151	18.114	36.741	40.113	43.194	46.963	49.645
28	12.461	13.565	15.308	16.928	18.939	37.916	41.337	44.461	48.278	50.993
29	13.121	14.257	16.047	17.708	19.768	39.087	42.557	45.722	49.588	52.336
30	13.787	14.954	16.791	18.493	20.599	40.256	43.773	46.979	50.892	53.672
40	20.707	22.164	24.433	26.509	29.051	51.805	55.758	59.342	63.691	66.766
50	27.991	29.707	32.357	34.764	37.689	63.167	67.505	71.420	76.154	79.490
60	35.534	37.485	40.482	43.188	46.459	74.397	79.082	83.298	88.379	91.952
70	43.275	45.442	48.758	51.739	55.329	85.527	90.531	95.023	100.425	104.215
80	51.172	53.540	57.153	60.391	64.278	96.578	101.879	106.629	112.329	116.321
90	59.196	61.754	65.647	69.126	73.291	107.565	113.145	118.136	124.116	128.299
100	67.328	70.065	74.222	77.929	82.358	118.498	124.342	129.561	135.807	140.169

Figura 18: Tabela referência Chi-Quadrado

No caso do D6 temos um total de 5 graus de liberdade, assim, deve-se usar a linha 5. Assumindo que gostaria-se que um dado tivesse pelo menos 95% de precisão, o valor obtido deveria ser inferior a 1,145, e, como o valor obtido foi de 0,343, pelo teste de Chi-Quadrado ele tem um viés menor que 0,05%, indicando que o dado é extremamente confiável, assim, tem seu uso aprovado.

3.2.2 Teste de aderência (Kolmogorov-Smirnov)

Para os lançamentos feitos com o dado de seis faces, foi realizado o teste de Kolmogorov-Smirnov. Para dados não viciados, a frequência esperada é uniforme e, portanto, a frequência esperada para cada face é de 0.1667.

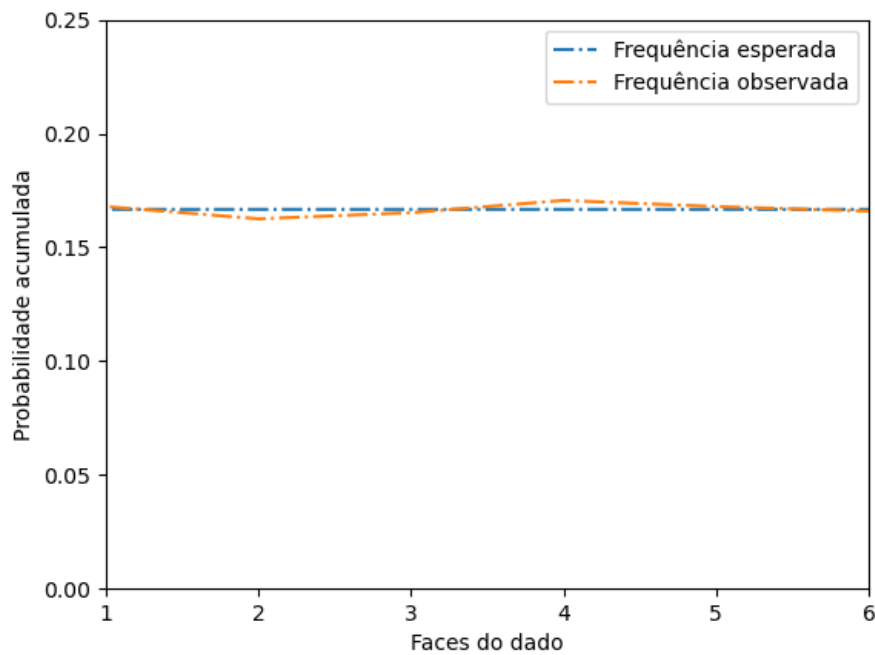


Figura 19: Distribuição uniforme esperada versus observada de um D6

Dessa forma, a tabela que descreve a frequência esperada e a frequência esperada acumulada para um D6 está exposta a seguir:

Face	Frequência Esperada	Frequência Esperada Acumulada
1	0.1667	0.1667
2	0.1667	0.3333
3	0.1667	0.5000
4	0.1667	0.6666
5	0.1667	0.8333
6	0.1667	1.0000

Tabela 10: Frequência esperada para o D6

Nos dados colhidos com o uso de reconhecimento de imagem, a frequência obtida para cada face dos dados é mostrada na tabela a seguir.

Face	Frequência Observada	Frequência Observada Acumulada
1	0.1679	0.1679
2	0.1625	0.3304
3	0.1652	0.4956
4	0.1706	0.6662
5	0.1679	0.8341
6	0.1658	1.0000

Tabela 11: Frequência observada para o D6

Com a combinação dos resultados expostos nas tabelas acima, é possível obter os valores das diferenças da frequência acumulada esperada para a frequência acumulada observada, conforme tabela abaixo. O maior valor - que se consagra o resultado do teste de Kolmogorov-Smirnov (teste de aderência) é de 0.004383, para a terceira face do dado.

	K
1	0.0012
2	0.0029
3	0.0043
4	0.0004
5	0.0007
6	0.0000

Tabela 12: Diferença entre a frequência esperada acumulada e a frequência observada acumulada para o D6

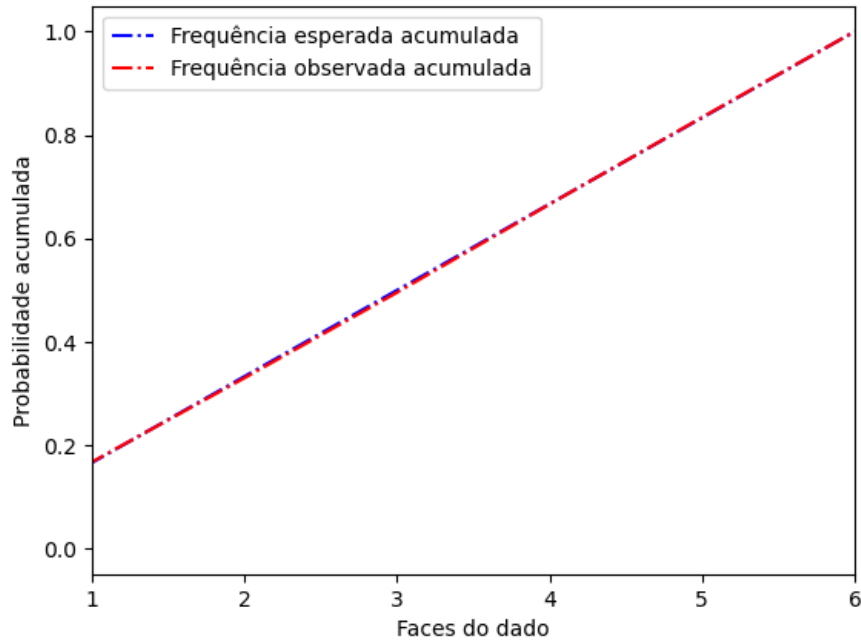


Figura 20: Frequência acumulada esperada *vs* frequência acumulada observada de um D6

O valor do K crítico (K_c) para um teste de 5% de significância com n (tamanho da amostra) igual a 1483 pode ser obtido por meio da equação:

$$K_c = \frac{1.36}{\sqrt{n}}$$

O resultado do valor crítico para esse teste é de 0.035316. Considerando o resultado obtido de um K máximo de 0.004383, entende-se que o dado não é enviesado porque o K máximo obtido é inferior ao valor crítico.

3.2.3 Discussões e conclusão

Por fim, após os testes realizados, conclui-se que o dado escolhido não apresenta um viés relevante, com uma confiança de maior de 95%. Além disso, também é possível perceber a relevância da realização de mais de um teste para a confirmação de hipóteses desenvolvidas em projetos de engenharia.

Também observa-se a potência de ferramentas de visão computacional, com enfoque na biblioteca OpenCV, expandindo os horizontes de engenharia e projetos que são possíveis serem construídos usando esta biblioteca em conjunto com outras ferramentas

computacionais. Um grande enfoque também fica nos descritores, principal ferramenta utilizada no projeto, que abriu caminhos não pensados anteriormente, e também gera uma alternativa viável e interessante ao reconhecimento de imagem a partir de aprendizado de máquinas, apresentando resultados extremamente satisfatórios e competitivos com esses métodos.

Quanto às contribuições de do projeto, há a ferramenta de leitura computacional para estudo de lançamento de dados, que fornece uma forma simples, rápida e viável para usuários testarem os seus dados poliédricos. Além disso, há a criação de bases de dados de lançamentos que podem ser usadas para projetos futuros, para o treinamento de Machine Learning por exemplo, de forma a se explorar novas possibilidades para a execução das tarefas apresentadas anteriormente, e, como engenheiros, buscando sempre novas soluções para os problemas encontrados. Para trabalhos futuros, sugere-se o aprimoramento da precisão da ferramenta de visão computacional desenvolvida para que sua aplicabilidade seja expandida a outros tipos de dado. Ainda, sugere-se a criação de uma ferramenta mecânica de lançamento automático dos dados.

4 REFERÊNCIAS BIBLIOGRÁFICAS

MALONEY, Dan. Automated Dice Tester Uses Machine Vision to Ensure a Fair Game. Hackaday, 2019. Disponível em: <https://hackaday.com/2019/06/20/automated-dice-tester-uses-machine-vision-to-ensure-a-fair-game/>. Acesso em: 10 de março de 2022.

LEYSHON, Tom. Modelling the probability distributions of dice. Towards data science, 2021. Disponível em: <https://towardsdatascience.com/modelling-the-probability-distributions-of-dice-b6ecf87b24ea>. Acesso em: 15 de março de 2022.

WU, Meyin. CHEN, Li. Image recognition based on deep learning. College of Computer Science and Technology, Wuhan University of Science and Technology. Key Laboratory of Intelligent Information Processing and Real-time Industrial System.

WIMSATT, Hunter. PANZADE, Aarohl. SHANKAR, Kaaustaub. CAMPBELL, Warren. Using Machine Learning to Interpret Dice Rolls. Western Kentucky University, School of Engineering and Applied Sciences, 2021.

SAHA, Sumit. A Comprehensive Guide to Convolutional Neural Networks - the ELI5 way. Towards data science, 2018. Disponível em: <https://towardsdatascience.com/a-comprehensive-guide-to-convolutional-neural-networks-the-eli5-way-3bd2b1164a53>. Acesso em: 18 de abril de 2022.

Kit 7 Dados Rpg D4 D6 D8 D10 D10 D12 D20 Dungeons Dragons. Mercado Livre, 2022. Disponível em: <https://produto.mercadolivre.com.br/MLB-695737269-kit-7-dados->

rpg-d4-d6-d8-d10-d10-d12-d20-dungeons-dragons-JM. Acesso em: 18 de abril de 2022.

LUO, Chuan. Overview of Image Matching Based on ORB Algorithm. Department of Mechanical Engineering, Shandong University, 2019.

AGRANONIK, Marilyn. HIRAKATA, Vânia Naomi. Cálculo de tamanho de amostra: proporções. Revista HCPA, 2011.

ADÁMEK, Roman. Dice Rolling Machine. YouTube, 2016. Disponível em: <https://www.youtube.com/watch?v=W9zJ0b91SQ0>. Acesso em: 16 de maio de 2022.

ARSENAULT, Marc-Olivier. Kolmogorov-Smirnov test. A needed tool in your data science toolbox. Towards data science, 2017. Disponível em: <https://towardsdatascience.com/kolmogorov-smirnov-test-84c92fb4158d>. Acesso em: 10 de novembro de 2022.

FERNANDEZ, Javier. Choose the appropriate normality test Shapiro-Wilk test, Kolmogorov-Smirnov test, and D'Agostino-Pearson's K^2 test. Towards data science, 2022. Disponível em: <https://towardsdatascience.com/choose-the-appropriate-normality-test-d53146ca1f1c>. Acesso em: 10 de novembro de 2022.

RIBEIRO, Evandro Marcos Saidel. Teste Kolmogorov-Smirnov. YouTube, 2020. Disponível em: <https://www.youtube.com/watch?v=JJZfeHpbMgo>. Acesso em: 10 de novembro de 2022.

JESUS, Samuel Rebouças de. ROCHA, Walber Conceição de Jesus. BITTENCOURT, João Carlos Nunes. Análise de Desempenho de Detectores e Descritores de Características Utilizando a Plataforma Computacional Raspberry Pi. Centro de Ciências Exatas e Tecnológicas. Universidade Federal do Recôncavo da Bahia. Brasil, 2019.

ANEXO A – CÓDIGOS UTILIZADOS

Para acessar as imagens e o código do projeto, Clique [AQUI](#)