



hochschule mannheim

Comparison of compression tools for biological data and analysis of possible improvements

Gabriel Eichelkraut

Bachelor Thesis

for the acquisition of the academic degree Bachelor of Science (B.Sc.)

Course of Studies: Computer Science

Department of Computer Science

University of Applied Sciences Mannheim

01.12.22

Tutors

Prof. Elena Fimmel, Hochschule Mannheim

TBD

Eichelkraut, Gabriel:

Comparison of compression tools for biological data and analysis of possible improvements

/ Gabriel Eichelkraut. –

Bachelor Thesis, Mannheim: University of Applied Sciences Mannheim, 2022. 39 pages.

Eichelkraut, Gabriel:

Vergleich von Kompressionswerkzeugen für biologische Daten und Analyse von Verbesserungsmöglichkeiten

/ Gabriel Eichelkraut. –

Bachelor-Thesis, Mannheim: Hochschule Mannheim, 2022. 39 Seiten.

Erklärung

Hiermit erkläre ich, dass ich die vorliegende Arbeit selbstständig verfasst und keine anderen als die angegebenen Quellen und Hilfsmittel benutzt habe.

This work is licensed under a Creative Commons
“Attribution-ShareAlike 4.0 International” license.



Mannheim, 01.12.22

May Mustermaier

Gabriel Eichelkraut

Abstract

Comparison of compression tools for biological data and analysis of possible improvements

A variety of algorithms is used to compress sequenced DNA. New findings in the patterns of how the building blocks of DNA are distributed, might provide a chance to improve long used compression algorithms. The comparison of four widely used compression methods and a analysis on their implemented algorithms, leads to the conclusion that improvements are feasible. The closing discussion provides insights, in possible improvement approaches and which challenges they might involve.

Vergleich von Kompressionswerkzeugen für biologische Daten und Analyse von Verbesserungsmöglichkeiten

Verschiedene Algorithmen werden verwendet, um sequenzierte DNA zu speichern. Eine neue Entdeckung darüber wie die Bausteine der DNA angeordnet sind, bietet die Möglichkeit vorhandene kompressions methoden zum speichern von sequenzierten DNA zu verbessern.

Diese Arbeit vergleicht vier weit verbreitete Kompressionsmethoden und analysiert deren Verwendung von Algorithmen. Durch die Ergebnisse lässt sich der Schluss ziehen, dass Verbesserungen in der Implementation von arithmetischer codierung möglich sind. Die abschließende Diskussion betrachtet mögliche vorgehensweisen zur Vebesserung und welche Aufgaben diese mit sich ziehen könnten.

Contents

1. Introduction	1
2. The Structure of the Human Genome and how its Digital Form is Compressed	3
2.1. Structure of Human DNA	3
2.2. File Formats used to Store DNA	5
2.2.1. FASTA and FASTq	7
2.2.2. Sequence Alignment Map	9
2.3. Compression approaches	10
2.3.1. Dictionary coding	11
2.3.2. Shannons Entropy	13
2.3.3. Arithmetic coding	14
2.3.4. Huffman encoding	17
2.4. Implementations in Relevant Tools	22
2.4.1. Genome Compressor (GeCo)	22
2.4.2. Samtools	23
3. Environment and Procedure to Determine the State of The Art Efficiency and Compressionratio of Relevant Tools	27
3.1. Server specifications and test environment	28
3.2. Operating System and Additionally Installed Packages	29
3.3. Selection, Receivment, and Preperation of Testdata	29
4. Results and Discussion	32
4.1. Interpretation of Results	33
4.2. View on Possible Improvements	36
List of Abbreviations	v
List of Tables	vi
List of Figures	vii
Listings	viii

Contents

Bibliography	ix
Index	xiv
A. Erster Anhang: Lange Tabelle	xvi

Chapter 1

Introduction

Understanding the biological code of living things, is a always developing task which is important for multiple aspects of our life. The results of research in this area provides knowledge that helps development in the medical sector, agriculture and more **wang-22**, [1], [2]. Getting insights into this biological code is possible through storing and studying information, embedded in genomes [3]. Since life is complex, there is a lot of information, which requires a lot of memory [4], [5].

With compression algorithms and their implementation in tools, the problem of storing information got smaller. Compressed data requires less space and therefore less time to be transported over networks [6]. This advantage is scalable, and since genetic information needs a lot of storage, even in a compressed state, improvements are welcomed [7]. Since this field is, compared to others, like computer theory which created the foundation for compression algorithms, relatively new, there is much to discover and new findings are not unusual [6]. From some of these findings, new tools can be developed. In general they focus on increasing at least one of two factors: the speed at which data is compressed and the compression ratio, meaning the difference between uncompressed and compressed data [4], [6], [7].

New discoveries in the universal rules of stochastic organization of genomes might provide a base for new algorithms and therefore new tools or an improvement of existing ones for genome compression [8]. The aim of this work is to analyze the current state of the art for compression tools for biological data and implemented, probabilistic algorithms. Further this work will determine if there is room for improvement.

The discussion will include a superficial analysis of how and where this new

approach could get implemented and what problems possibly need to be taken care of in the process.

To reach a common ground, the first pages will give the reader a quick overview on the structure of human DNA. There will also be an fundamental explanation for some basic terms, used in biology and computer science. The first step into the theory of genome compression will be taken, by describing differences in common file formats, used to store genome information. From there, a section which is relevant for understanding compression will follow. It will analyze differences between compression approaches, go over some history of coding theory and lead to a deeper look into the fundamentals of state of the art compression algorithms. The chapter will end with a few pages about implementations of compression algorithms in tools relevant.

In order to measure a improvement, a baseline must be set. Therefore the efficiency and effecitivity of suiteable state of the art tools will be meassured. To be as precise as possible, the middle part of this work focuses on setting up an environment, picking input data, installing and executing tools and finaly meassuring and documenting the results.

The results of this compared with the understanding of how the tools work, will show if there is the need of a improvement and on what factor it should focus. The end of this work will be used to discuss the properties of a a possible improvement, how feasability could be determined and which problems such a project woudl need overcome.

Chapter 2

The Structure of the Human Genome and how its Digital Form is Compressed

2.1. Structure of Human DNA

To strengthen the understanding of how and where biological information is stored, this section starts with a quick and general rundown on the structure of any living organism.

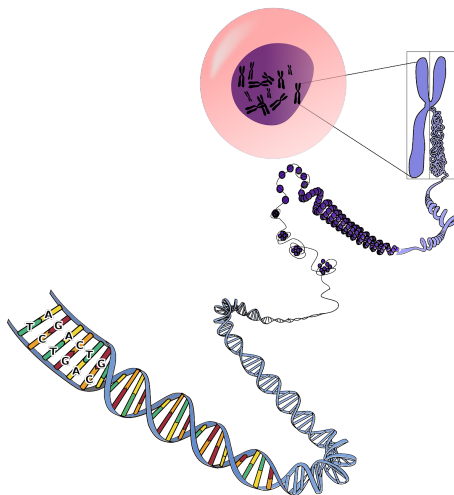


Figure 2.1.: A superficial representation of the physical positioning of genomes. Showing a double helix (bottom), a chromosome (upper right) and a cell (upper center).

All living organisms, like plants and animals, are made of cells. To get a rough impression, a human body can consist out of several trillion cells. A cell in itself, is the smallest living organism. Most cells consists of a outer section and a core which is a called nucleus. In 2.1 the nucleus is illustrated as a purple, circle like scheme,

2. The Structure of the Human Genome and how its Digital Form is Compressed

inside a lighter circle. The nucleus contains chromosomes. Those chromosomes contain genetic information, about its organism, in form of Deoxyribonucleic Acid (DNA) [9].

DNA is often seen in the form of a double helix, like shown in 2.2. A double helix consists, as the name suggests, of two single helix [3].

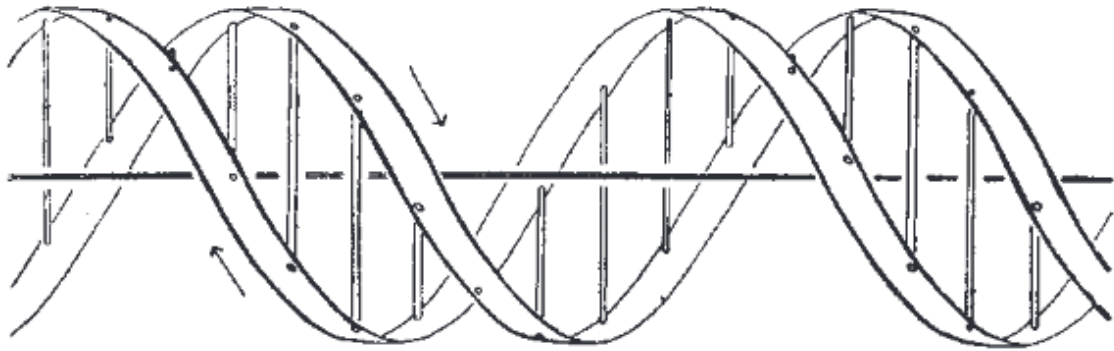


Figure 2.2.: A purely diagrammatic figure of the components DNA is made of. The smaller, inner rods symbolize nucleotide links and the outer ribbons the phosphate-sugar chains [3].

Each of them consists of two main components: the sugar phosphate backbone, which is not relevant for this work and the bases. The sugar phosphate backbones are illustrated as flat stripes, circulating around the horizontal line in 2.2. Pairs of bases are symbolized as vertical bars between the sugar phosphates. The arrangement of Bases represents the information, stored in the DNA. What is here described as base is an organic molecule, which is also called nucleotide [3].

For this work, nucleotides are the most important parts of the DNA. A nucleotide can occur in one of four forms: it can be either adenine, thymine, guanine or cytosine. Each of them got a Counterpart with which a bond can be established: adenine can bond with thymine, guanine can bond with cytosine.

From the perspective of a computer scientist: The content of one helix must be stored, to persist the full information. In more practical terms: The nucleotides of only one (entire) helix needs to be stored physically, to save the information of the whole DNA. The other half can be determined by “inverting” the stored one. An example would show the counterpart for e.g.: adenine, guanine, adenine chain which would be a chain of thymine, cytosine, thymine. For the sake of simplicity, one does not write out the full name of each nucleotide, but only its initial. So the example would change to AGA in one Helix, TCT in the other.

This representation is commonly used to store DNA digitally. Depending on the sequencing procedure and other factors, more information is stored and therefore

more characters are required but for now 'A', 'C', 'G' and 'T' should be the only concern.

2.2. File Formats used to Store DNA

As described in previous chapters DNA can be represented by a string with the buildingblocks A,T,G and C. Using a common file format for saving text would be impractical because the ammount of characters or symbols in the used alphabet, defines how many bits are used to store each single symbol.

The American Standard Code for Information Interchange (ASCII) [10] table is a character set, registered in 1975 and to this day still in use to encode texts digitally. For the purpose of communication bigger character sets replaced ASCII. It is still used in situations where storage is short.

The buildingblocks of DNA require a minimum of four letters, so at least two bits are needed. Storing a single A with ASCII encoding, requires 8 bit (excluding magic bytes and the bytes used to mark the End of File (EOF)) Since there are at least 2^8 or 128 displayable symbols with ASCII encoding, this leaves a great overhead of unused combination.

In most tools, more than four symbols are used. This is due to the complexity in sequencing DNA. It is not 100% preceice, so additional symbols are used to mark nucelotides that could not or could only partly get determined. Further a so called quality score is used to indicate the certainty, for each single nucleotide, that it got sequenced correctly [5], [11].

More common everyday-usage text encodings like unicode require 16 bits per letter. So settling with ASCII has improvement capabilities but is, on the other side, more efficient than using bulkier alternatives like unicode.

Formats for storing uncompressed genomic data, can be sorted into several categories. Three noticable ones would be [5]:

- Sequenced reads
- Aligned data
- Sequence variation

2. The Structure of the Human Genome and how its Digital Form is Compressed

The categories are listed on their complexity, considering their usecase and data structure, in ascending order. Starting with sequence variation, also called haplotype describes formats storing graph based structures that focus on analysing variations in different genomes [12], [13]. Sequenced reads focus on storing continuous protein chains from a sequenced genome [5]. Aligned data is somewhat similar to sequenced reads with the difference that instead of a whole chain of genomes, overlapping subsequences are stored. This could be described as a rawer form of sequenced reads. This way aligned data stores additional information on how certain a specific part of a genome is read correctly [5], [13]. The focus of this work lays on compression of sequenced data but not on the likelihood of how accurate the data might be. Therefore, only formats that are able to store sequenced reads will be worked with. Note that some aligned data formats are also able to store aligned reads, since latter is just a less informative representation of first [5], [13].

Several people and groups have developed different file formats to store genomes. Unfortunately, the only standard for storing genomic data is fairly new [14], [15]. Therefore, formats and tools implementing this standard are mostly still in development. In order to not go beyond scope, this work will focus only on file formats that fulfill following criteria:

- The format has reputation. This can be indicated through:
 - A scientific paper, that proved its superiority to other relevant tools.
 - A broad usage of the format determined by its use on ftp servers, which focus on supporting scientific research.
- The format should not specialize on only one type of DNA or target a specific technology.
- The format stores nucleotide sequences and does not necessarily include International Union of Pure and Applied Chemistry (IUPAC) codes besides A, C, G and T [16].
- The format is open source. Otherwise, improvements can not be tested, without buying the software and/or requesting permission to disassemble and reverse engineer the software or parts of it.

2. The Structure of the Human Genome and how its Digital Form is Compressed

Information on available formats were gathered through various Internet platforms [17]–[19] and scientific papers [5], [11], [13]. Some common file formats found:

- File Format for Storing Genomic Data (FASTA)/Multi-FASTA
- File Format Based on FASTA (FASTq) [11]
- Sequence Alignment Map (SAM)/Binary Alignment Map (BAM) [13], [20]
- Compressed Reference-oriented Alignment Map (CRAM) [13], [20]

Since methods to store this kind of Data are still in development, there are many more file formats. From the selection listed above, FASTA and FASTq seem to have established the reputation of a inoficial standard for sequenced reads [5], [13], [21]–[23].

Considering the first criteria, by searching through anonymously accesable FTP servers, only two formats are used commonly: FASTA or its extension FASTq and the BAM Format [24]–[26].

2.2.1. FASTA and FASTq

The rather simple FASTA format, are widely used when it comes to storing sequenced reads, without a quality store [5], [13]. Since it is a uncompressed format, FASTA files are often transmitted compressed with an external tool like gzip [25], [26].

```
Header section
[
  >2 dna:chromosome chromosome:GRCh38:2:1:242193529:1 REF
  CCTAACCCCTCACCCTCACCCTCGACCCCGACCCCGACCCCGACCCCGACCCCGAAC
  CCGACCCCGACCCCGACCCAAACCCCTAACCCCTAAAACCCCTAACCCCTAGCCCTAGCCCTAG
  CCCTAGCCCTAACCCCTAACCCCTAACCCCTAAGCCGAAGCCTAACTCGTGTCTGACTTTG
  ...
Sequence section
```

The diagram illustrates the structure of a FASTA file. It is divided into two main parts: the 'Header section' and the 'Sequence section'. The header section is enclosed in square brackets and contains a single line starting with a greater-than sign (>) followed by a sequence identifier: '2 dna:chromosome chromosome:GRCh38:2:1:242193529:1 REF'. The sequence section is also enclosed in square brackets and contains multiple lines of nucleotide sequences (A, C, G, T) separated by spaces. The first line of the sequence section is 'CCTAACCCCTCACCCTCACCCTCGACCCCGACCCCGACCCCGACCCCGACCCCGAAC', followed by 'CCGACCCCGACCCCGACCCAAACCCCTAACCCCTAAAACCCCTAACCCCTAGCCCTAGCCCTAG', 'CCCTAGCCCTAACCCCTAACCCCTAACCCCTAAGCCGAAGCCTAACTCGTGTCTGACTTTG', and finally '...'. Brackets on the left side of the text indicate the extent of each section.

Figure 2.3.: Edited example of a FASTA file. The original was received from the ensemble server [26].

2. The Structure of the Human Genome and how its Digital Form is Compressed

The format consists of two repeated sections. The first section consists of one line and stores metadata about the sequenced genome and the file itself. This line, also called header, contains a comment section starting with > followed by a custom text [4], [11]. The comment section is usually used to store information about the sequenced genome and sometimes metadata about the file itself like its size in bytes. The other section contains the sequenced genome whereas each nucleotide is represented by character A, C, G or T. There are more nucleotide characters that store additional information and some characters for representing amino acids, but in order to not go beyond scope, only A, C, G, and T will be paid attention to [16]. The second section can have multiple lines of sequences. A similar format is the Multi-FASTA file format, it consists of concatenated FASTA files.[5].

In addition to its predecessor, FASTq files contain a quality score. The file content consists of four sections, whereby no section is stored in more than one line. All four lines contain information about one sequence. The exact structure of FASTq is formatted in this order [11]:

- Line 1: Sequence identifier aka. Title, starting with an @ and an optional description.
- Line 2: The sequence consisting of nucleotides, symbolized by A, T, G and C.
- Line 3: A '+' that functions as a separator or delimiter. Optionally followed by content of Line 1.
- Line 4: quality line(s). consisting of letters and special characters in the ASCII scope.

The quality scores have no fixed format. To name a few, there is the sanger format, the solexa format introduced by Solexa Inc., the Illumina and the QUAL format which is generated by the PHRED software.

The quality value shows the estimated probability of error in the sequencing process [11].

In 2.3 the described structure is illustrated. The sequence and the delimiter section were altered, to illustrate the structure of this format better. In the header section, SRR002906.1 is the sequence identifier, the following text is a description. In the delimiter line, the header section without the leading @ could be written again. The last line shows the header for the second sequence.

2. The Structure of the Human Genome and how its Digital Form is Compressed

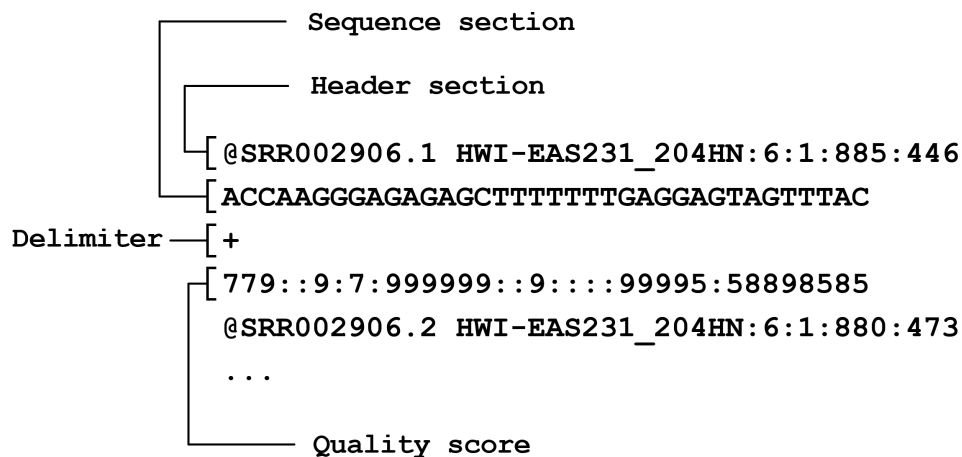


Figure 2.4.: Altered example of a FASTq file. The original was received from ncbi server [25].

2.2.2. Sequence Alignment Map

SAM often seen in its compressed, binary representation BAM with the file extension .bam, is part of the SAMtools package, a utility tool for processing SAM/BAM and CRAM files. The SAM/BAM file is a text based format delimited by the whitespace character called tabulation or **tab** for short [13]. It uses 7-bit US-ASCII, to be precise Charset ANSI X3.4-1968 [27]. The structure is more complex than the one in FASTq and described best, accompanied by an example:

```

Coor      12345678901234  5678901234567890123456789012345
ref       AGCATGTTAGATAA**GATAGCTGTGCTAGTAGGCAGTCAGCGCCAT

+r001/1   TTAGATAAAGGATA*CTG
+r002     aaaAGATAA*GGATA
+r003     gcctaAGCTAA
+r004           ATAGCT.....TCAGC
-r003           ttagctTAGGC
-r001/2           CAGCGGCAT

```

Figure 2.5.: SAM file structure example [20].

Compared to FASTA SAM and further compression forms, store more information. As displayed in 2.5 this is done by adding, identifier for Reads e.g. **+r003**, aligning subsequences and writing additional symbols like dots e.g. **ATAGCT.....** in the

split alignment +r004 [5]. A full description of the information stored in SAM files would be of little value to this work, therefore further information on is left out but can be found in [13] or at [20].

Samtools provide the feature to convert a FASTA file into SAM format. Since there is no way to calculate mentioned, additional information from the information stored in FASTA, the converted files only store two lines. The first one stores metadata about the file and the second stores the nucleotide sequence in just one line.

2.3. Compression approaches

The process of compressing data serves the goal to generate an output that is smaller than its input [28].

In many cases, like in gene compressing, the compression is ideally lossless. This means it is possible with any compressed data, to receive the full information that were available in the origin data, by decompressing it. Lossy compression on the other hand, might excludes parts of data in the compression process, in order to increase the compression rate. The excluded parts are typically not necessary to transmit the origin information. This works with certain audio and pictures files or with network protocols like Universal Datagram Protocol (UDP) which are used to transmit video/audio streams live [29], [30].

For storing DNA a lossless compression is needed. To be preceive a lossy compression is not possible, because there is no unnecessary data. Every nucleotide and its exact position is needed for the sequence to be complete and usefull.

Before going on, the difference between information and data should be emphasized.

Data contains information. In digital data clear, physical limitations delimit what and how much of something can be stored. A bit can only store 0 or 1, eleven bit can store up to 2^{11} combinations of bit and a 1 GB drive can store no more than 1 GB data. Information on the other hand, is limited by the way how it is stored. What exactly defines informations, depends on multiple factors. The context in which information is transmitted and the source and destination of the information. This can be in form of a signal, transfered from one entity to another or information that is persisted so it can be obtained at a later point in time.

For the scope of this work, information will be seen as the type and position of nucleotides, sequenced from DNA. To get even more preceise, it is a chain of char-

acters from a alphabet of A, C, G, and T, since this is the *de facto* standard for digital persistence of DNA [14]. The boundaries of information, when it comes to storing capabilities, can be illustrated by using the example mentioned above. A drive with the capacity of 1 GB could contain a book in form of images, where the content of each page is stored in a single image. Another, more resourceful way would be storing just the text of every page in UTF-16 [31]. The information, the text would provide to a potential reader would not differ. Changing the text encoding to ASCII and/or using compression techniques would reduce the required space even more, without losing any information.

For DNA a lossless compression is needed. To be precise a lossy compression is not possible, because there is no unnecessary data. Every nucleotide and its position is needed for the sequenced DNA to be complete. For lossless compression two major approaches are known: the dictionary coding and the entropy coding. Methods from both fields, that acquired reputation, are described in detail below [4], [7], [32], [33].

2.3.1. Dictionary coding

Dictionary coding, as the name suggest, uses a dictionary to eliminate redundant occurrences of strings. Strings are a chain of characters representing a full word or just a part of it. For a better understanding this should be illustrated by a short example: Looking at the string 'stationary' it might be smart to store 'station' and 'ary' as separate dictionary entries. Which way is more efficient depends on the text that should get compressed. The dictionary should only store strings that occur in the input data. Also storing a dictionary in addition to the (compressed) input data, would be a waste of resources. Therefore the dictionary is part of the text. Each first occurrence is left uncompressed. Each occurrence of a string, after the first one, points either to its first occurrence or to the last replacement of its occurrence.

2.6 illustrates how this process is executed. The bar on top of the figure, which extends over the full width, symbolizes any text. The squares inside the text are repeating occurrences of text segments. In the dictionary coding process, the square annotated as `first occ.` is added to the dictionary. `second` and `third occ.` get replaced by a structure `<pos, len>` consisting of a pointer to the position of the first occurrence `pos` and the length of that occurrence `len`. The bar at the bottom of

2. The Structure of the Human Genome and how its Digital Form is Compressed

the figure shows how the compressed text for this example would be structured. The dotted lines would only consist of two bytes, storing position and length, pointing to first occ.. Decompressing this text would only require to parse the text from left to right and replace every $\langle \text{pos}, \text{len} \rangle$ with the already parsed word from the dictionary. This means jumping back to the parsed position stored in the replacement, reading for as long as the length dictates, copying the read section, jumping back and pasting the section.

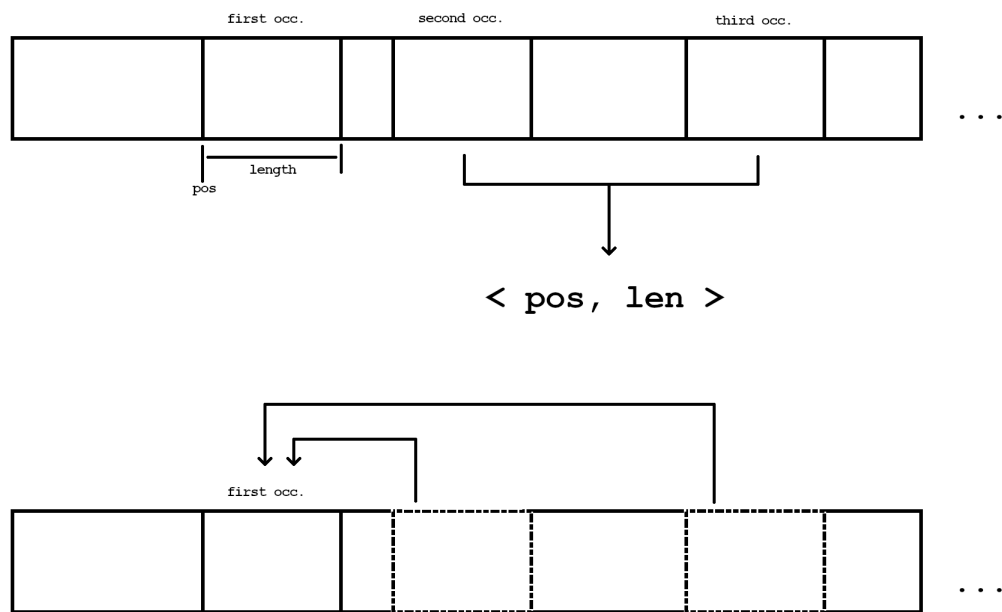


Figure 2.6.: Schematic sketch, illustrating the replacement of multiple occurrences done in dictionary coding.

The LZ Family

The computer scientist Abraham Lempel and the electrical engineer Jacob Ziv created multiple algorithms that are based on dictionary coding. They can be recognized by the substring LZ in its name, like LZ77 and LZ78 which are short for Lempel Ziv 1977 and 1978 [34]. The number at the end indicates when the algorithm was published. Today LZ78 is widely used in unix compression solutions like gzip and bzip2. Those tools are also used in compressing DNA.

2. The Structure of the Human Genome and how its Digital Form is Compressed

LZ77 basically works, by removing all repetition of a string or substring and replacing them with information where to find the first occurrence and how long it is. Lempel and Ziv described restricted the pointer in a range to integers. Today a pointer, length pair is typically stored in two bytes. One bit is reserved to indicate that the next 15 bit are a position, length pair. More than 8 bit are available to store the pointer and the rest is reserved for storing the length. Exact amounts depend on the implementation [34], [35].

Unfortunately, implementations like the ones out of LZ Family, do not use probabilities to compress and are therefore not in the main scope for this work. To strengthen the understanding of compression algorithms this section will remain. Also it will be useful for the explanation of a hybrid coding method, which will get described later in this chapter.

2.3.2. Shannons Entropy

The founder of information theory Claude Elwood Shannon described entropy and published it in 1948 [6]. In this work he focused on transmitting information. His theorem is applicable to almost any form of communication signal. His findings are not only useful for forms of information transmission.

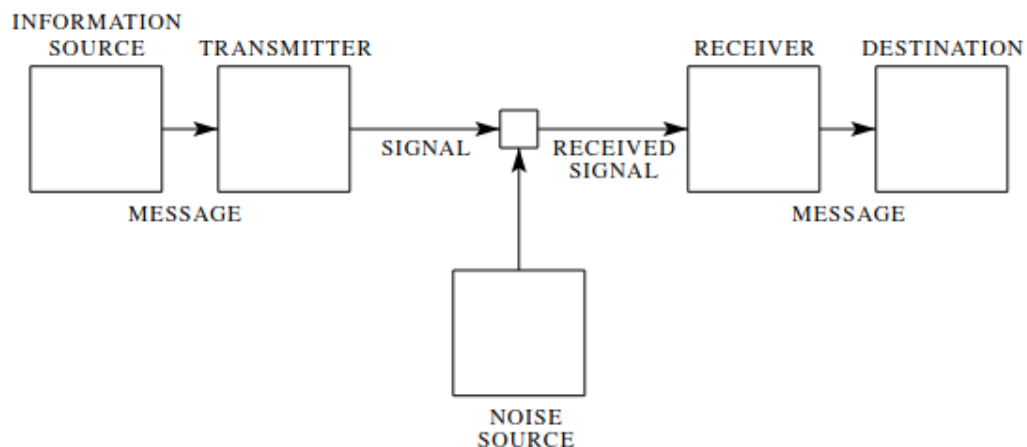


Figure 2.7.: Schematic diagram of a general communication system by Shannons definition. [6]

Altering 2.7 would show how this can be applied to other technology like compression. The Information source and destination are left unchanged, one has to keep in

2. The Structure of the Human Genome and how its Digital Form is Compressed

mind, that it is possible that both are represented by the same physical actor. Transmitter and receiver would be changed to compression/encoding and decompression/decoding. Inbetween those two, there is no signal but instead any period of time [6].

Shannons Entropy provides a formula to determine the 'uncertainty of a probability distribution' in a finite field.

$$H(X) := \sum_{x \in X, \text{prob}(x) \neq 0} \text{prob}(x) \cdot \log_2\left(\frac{1}{\text{prob}(x)}\right) \equiv - \sum_{x \in X, \text{prob}(x) \neq 0} \text{prob}(x) \cdot \log_2(\text{prob}(x)). \quad (2.1)$$

He defined entropy as shown in figure (2.1). Let X be a finite probability space. Then $x \in X$ are possible final states of an probability experiment over X . Every state that actually occurs, while executing the experiment generates information which is measured in *Bits* with the part of the equation displayed in 2.2 [6], [36]:

$$\log_2\left(\frac{1}{\text{prob}(x)}\right) \equiv -\log_2(\text{prob}(x)). \quad (2.2)$$

2.3.3. Arithmetic coding

This coding method is an approach to solve the problem of wasting memory due to the overhead which is created by encoding certain lengths of alphabets in binary [7], [37]. For example: Encoding a three-letter alphabet requires at least two bit per letter. Since there are four possible combinations with two bit, one combination is not used, so the full potential is not exhausted. Looking at it from another perspective and thinking a step further: Less storage would be required, if there would be a possibility to encode more than one letter in two bit.

Dr. Jorma Rissanen described arithmetic coding in a publication in 1976 [37]. This works goal was to define an algorithm that requires no blocking. Meaning the input text could be encoded as one instead of splitting it and encoding the smaller texts or single symbols. He stated that the coding speed of arithmetic coding is comparable to that of conventional coding methods [37].

2. The Structure of the Human Genome and how its Digital Form is Compressed

Before getting into the arithmetic coding algorithm, the following section will go over some details on how digital fractions are handled by computers. This knowledge will be helpful in understanding how arithmetic coding works.

In computers, arithmetic operations on floating point numbers are processed with integer representations of given floating point number [38]. The number $0.4 \cdot 10^{-1}$ would be represented by $4 \cdot 10^{-1}$.

A interval would be represented by natural numbers between 0 and 100 and $\dots \cdot 10^{-x}$. x starts with the value 2 and grows as the integers grow in length, meaning only if a uneven number is divided. For example: Dividing a uneven number like $5 \cdot 10^{-1}$ by two, will result in $25 \cdot 10^{-2}$. On the other hand, subdividing $4 \cdot 10^y$ by two, with any negative real number as y would not result in a greater x the length required to display the result will match the length required to display the input number [7], [39].

Binary fractions are limited in from of representing decimal fractions. This is due to the fact that every other digit, adds zero or half of the value before. In other terms: $b \cdot 2^{-n}$ determines the value of $b \in 0, 1$ at position n behind the decimal point.

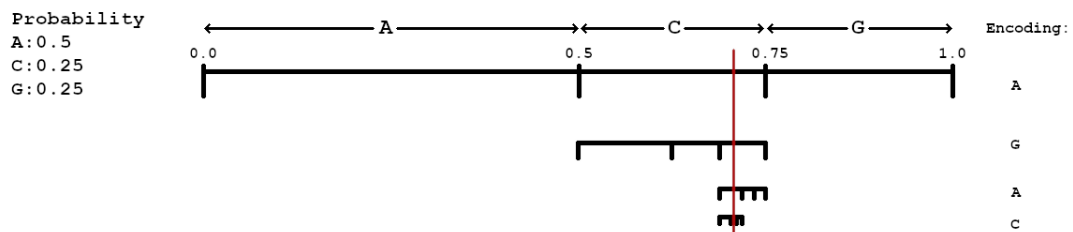


Figure 2.8.: Illustrative example of arithmetic coding.

The encoding of the input text, or a sequence is possible by projecting it on a binary encoded fraction between 0 and 1. To get there, each character in the alphabet is represented by an interval between two fractions, in the space between 0.0 and 1.0. In 2.8 this space is illustrated by the line in the upper center, with a scaling from 0.0 on the left, to 1.0 on the right side. The interval for each symbol is determined by its distribution, in the input text (interval start) and the the start of the next character (interval end). The sum of all intervals will result in one [7].

In order, to remain in a presentable range, the example in 2.8 uses an alphabet of only three characters: A, C and G. For the sequence AGAC a probability distribution as shown in the upper left corner and listed in 2.3.3 was calculated. The inter-

2. The Structure of the Human Genome and how its Digital Form is Compressed

vals resulting from this probabilities, are visualized by the three sections marked by outwards pointing arrows at the top. The interval for A extends from 0.0 until the start of C at 0.5, which extends to the start of G at 0.75 and so on.

Table 2.1.: Probabilities for A, C and G as shown in example 2.8

Symbol	Probability	Interval
A	$\frac{2}{4} = 0.11$	$x \in \mathbb{Q} 0.0 \leq x < 0.5$
C	$\frac{1}{4} = 0.71$	$x \in \mathbb{Q} 0.5 \leq x < 0.75$
G	$\frac{1}{4} = 0.13$	$x \in \mathbb{Q} 0.75 \leq x < 1.0$

In the encoding process, the first symbol read from the sequence determines a interval, its symbol is associated with. Every following symbol determines a subinterval, which is determined by subdividing the previous interval into sections proportional to the probabilities from 2.3.3. Starting with A, the most left interval in 2.8 is subdivided into intervals visualized below. Leaving a available space of $[0.0, 0.5)$. From there the interval, representing G is subdivided, and so on until the last symbol C is processed. This leaves a interval of $[0.40625, 0.421275)$.

To store the encoding result in as few bits as possible, only a single number, between upper and lower end of the last interval will be stored. To encode in binary, the binary floating point representation of any number inside the interval, for the last character is calculated.

For this example, the number 0.41484375 in decimal, or 0.0110101 in binary, would be calculated.

To summarize the encoding process in short [7], [39]:

- The interval representing the first character is noted.
- Its interval is split into smaller intervals, with the ratios of the initial intervals between 0.0 and 1.0.
- The interval representing the second character is choosen.
- This process is repeated, until a interval for the last character is determined.
- A binary floating point number is determined wich lays in between the interval that represents the represents the last symbol.

2. The Structure of the Human Genome and how its Digital Form is Compressed

For the decoding process to work, the EOF symbol must be present as the last symbol in the text. The compressed file will store the probabilities of each alphabet symbol as well as the floatingpoint number. The decoding process executes in a similar procedure as the encoding. The stored probabilities determine intervals. Those will get subdivided, by using the encoded floating point as guidance, until the EOF symbol is found. By noting in which interval the floating point is found, for every new subdivision, and projecting the probabilities associated with the intervals onto the alphabet, the original text can be read [7], [37], [39].

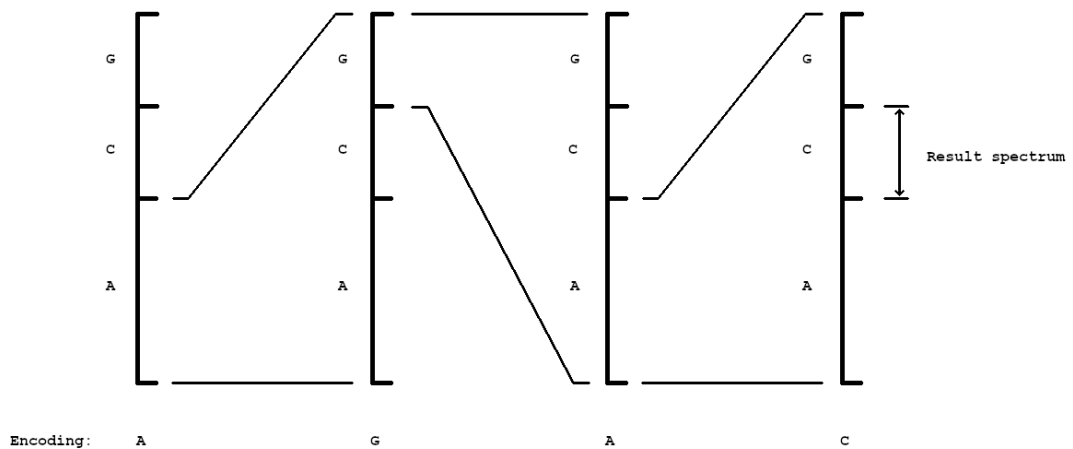


Figure 2.9.: Illustrative rescaling in arithmetic coding process.

The described coding is only feasible on machines with infinite precision [39]. As soon as finite precision comes into play, the algorithm must be extended, so that a certain length in the resulting number will not be exceeded. Since digital datatypes are limited in their capacity, like unsigned 64-bit integers which can store up to $2^{64} - 1$ bit or any number between 0 and 18,446,744,073,709,551,615. That might seem like a great amount at first, but considering a unfavorable alphabet, that extends the results length by one on each symbol that is read, only texts with the length of 63 can be encoded (62 if EOF is excluded) [7].

2.3.4. Huffman encoding

D. A. Huffman's work focused on finding a method to encode messages with a minimum of redundancy. He referenced a coding procedure developed by Shannon and

2. The Structure of the Human Genome and how its Digital Form is Compressed

Fano and named after its developers, which worked similar. The Shannon-Fano coding is not used today, due to the superiority in both efficiency and effectivity, in comparison to Huffman. Even though his work was released in 1952, the method he developed is in use today. Not only tools for genome compression but in compression tools with a more general usage [40].

Compression with the Huffman algorithm also provides a solution to the problem, described at the beginning of 2.3.2, on waste through unused bit, for certain alphabet lengths. Huffman did not save more than one symbol in one bit, like it is done in arithmetic coding, but he decreased the number of bit used per symbol in a message. This is possible by setting individual bit lengths for symbols, used in the text that should get compressed [41]. As with other codings, a set of symbols must be defined. For any text constructed with symbols from mentioned alphabet, a binary tree is constructed, which will determine how the symbols will be encoded. As in arithmetic coding, the probability of a letter is calculated for given text. The binary tree will be constructed after following guidelines [4]:

- Every symbol of the alphabet is one leaf.
- The right branch from every knot is marked as a 1, the left one is marked as a 0.
- Every symbol got a weight, the weight is defined by the frequency the symbol occurs in the input text. This might be a fraction between 0 and 1 or an integer. In this scenario it will be described as the first.
- The less weight a leaf has, the higher the probability is, that this node is read next in the symbol sequence.
- The leaf with the lowest probability is most left and the one with the highest probability is most right in the tree.

A often mentioned difference between Shannon-Fano and Huffman coding, is that first is working top down while the latter is working bottom up. This means the tree starts with the lowest weights. The nodes that are not leafs have no value ascribed to them. They only need their weight, which is defined by the weights of their individual child nodes [4], [33].

2. The Structure of the Human Genome and how its Digital Form is Compressed

Given $K(W, L)$ as a node structure, with the weight or probability as W_i and codeword length as L_i for the node K_i . Then will L_{av} be the average length for L in a finite chain of symbols, with a distribution that is mapped onto W [41].

$$L_{av} = \sum_{i=0}^{n-1} w_i \cdot l_i \quad (2.3)$$

The equation (2.3) describes the path, to the desired state, for the tree. The upper bound n is assigned the length of the input text. The tuple in any node K consists of a weight w_i , that also references a symbol, and the length of a codeword l_i . This codeword will later encode a single symbol from the alphabet. Working with digital codewords, an element in L contains a sequence of zeros and ones. Since there in this coding method, there is no fixed length for codewords, the premise of prefix free code must be adhered to. This means there can be no codeword that match the sequence of any prefix of another codeword. To illustrate this: 0, 10, 11 would be a set of valid codewords but adding a codeword like 01 or 00 would make the set invalid because of the prefix 0, which is already a single codeword.

With all important elements described: the sum that results from this equation is the average length a symbol in the encoded input text will require to be stored [33], [41].

For this example a four letter alphabet, containing A, C, G and T will be used. For this alphabet, the binary representation encoded in ASCII is listed in the second column of 2.3.4. The average length for any symbol encoded in ASCII is eight, while only using four of the available 2^8 symbols, a overhead of 252 unused bit combinations. For this example it is more vivid, using a imaginary encoding format, without overhead. It would result in a average codeword length of two, because four symbols need a minimum of 2^2 bit.

Table 2.2.: ASCII Codes and probabilities for A, C, G and T

Symbol	ASCII Code	Probability	Occurences
A	0100 0001	$\frac{11}{100} = 0.11$	11
C	0100 0011	$\frac{71}{100} = 0.71$	71
G	0101 0100	$\frac{13}{100} = 0.13$	13
T	0000 1010	$\frac{5}{100} = 0.05$	5

The exact input text is not relevant, since only the resulting probabilities are needed. To make this example more illustrative, possible occurences are listed in the most

2. The Structure of the Human Genome and how its Digital Form is Compressed

right column of 2.3.4. The probability for each symbol is calculated by dividing the message length by the times the symbol occurred. This and the resulting probabilities on a scale between 0.0 and 1.0, for this example are shown in 2.3.4 [41].

Creating a tree will be done bottom up. In the first step, for each symbol from the alphabet, a node without any connection is formed .

<A>, <T>, <C>, <G>

Starting with the two lowest weighted symbols, a node is added to connect both. With the added, blank node the count of available nodes got down by one. The new node weights as much as the sum of weights of its child nodes so the probability of 0.16 is assigned to <A,T>.

<A, T>, <C>, <G>

From there on, the two leafs will only get rearranged through the rearrangement of their temporary root node. Now the two lowest weights are paired as described, until there are only two subtrees or nodes left which can be combined by a root.

<G, <A, T> >, <C>

The <G, <A, T> > has a probability of 0.29. Adding the last, highest weighted node C results in a root node with the probability of 1.0. For a better understanding of this example, and to help further explanations, the resulting tree is illustrated in 2.10.

2. The Structure of the Human Genome and how its Digital Form is Compressed

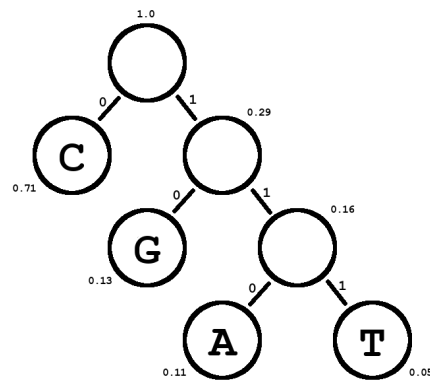


Figure 2.10.: Final version of the Huffman tree for described example.

As illustrated in 2.10 the left branches are assigned with 0 and right branches with 1, following a path until a leaf is reached reveals the encoding for this particular leaf. With a corresponding tree, created from with the weights, the binary sequences to encode the alphabet can be seen in the second column of 2.3.4.

Table 2.3.: Huffman codes for A, C, G and T

Symbol	Huffman Code	Occurences
A	100	11
C	0	71
G	11	13
T	101	5

Since high weighted and therefore often occurring leafs are positioned to the left, short paths lead to them and so only few bit are needed to encode them. Following the tree on the other side, the symbols occur more rarely, paths get longer and so do the codeword. Applying (2.3) to this example, results in 1.45 bit per encoded symbol. In this example the text would require over one bit less storage for every second symbol [41].

Leaving the theory and entering the practice, brings some details that lessen this improvement by a bit. A few bytes are added through the need of storing the information contained in the tree. Also, like described in 2.2 most formats, used for persisting DNA, store more than just nucleotides and therefore require more characters [11], [13].

2.4. Implementations in Relevant Tools

This section should give the reader a overview, how a small variety of compression tools implement described compression algorithms. It is written with the goal to compensate a problem that occurs in scientific papers, and sometimes in technical specifications for programs. They often lack information on the implementation, in a satisfying dimension [13], [20], [21].

The information on the following pages was received through static code analysis. Meaning the comprehension of a programs behaviour or its interactions due to the analysis of its source code. This is possible because the analysed tools are openly published and licenced under GNU General Public License (GPL) v3 [21] and Massachusetts Institute of Technology (MIT)/Expat [20], which permits the free use for scientific purposes [42], [43].

2.4.1. GeCo

This tool has three development stages. the first GeCo released in 2016 GeCo. This tool happens to have the smallest codebase, with only eleven C files. The two following extensions GeCo2, released in 2020 and the latest version GeCo3 have bigger codebases [44]. They also provide features like the usage of a neural network, which are of no help for this work. Since the file, providing arithmetic coding functionality, do not differ between all three versions, the first release was analyzed. The header files, that this tool includes in `geco.c`, can be split into three categories: basic operations, custom operations and compression algorithms. The basic operations include header files for general purpose functions, that can be found in almost any c++ Project. The provided functionality includes operations for text-output on the command line interface, memory management, random number generation and several calculations on up to real numbers.

Custom operations happens to include general purpose functions too, with the difference that they were written, altered or extended by GeCos developer. The last category consists of several C Files, containing implementations of two arithmetic coding implementations: **first** `bitio.c` and `arith.c`, **second** `arith_aux.c`.

The first two were developed by John Carpinelli, Wayne Salamonsen, Lang Stuver and Radford Neal (is only mentioned in the latter). Comparing the two files, `bitio.c` has less code, shorter comments and much more not functioning code sec-

tions. Overall the conclusion would be likely that `arith.c` is some kind of official release, whereas `bitio.c` serves as an experimental file for the developers to create proof of concepts. The described files adapt code from Armando J. Pinho licensed by University of Aveiro DETI/IEETA written in 1999.

The second implementation was also licensed by University of Aveiro DETI/IEETA, but no author is mentioned. From interpreting the function names and considering the length of function bodies `arith_aux.c` could serve as a wrapper for basic functions that are often used in arithmetic coding.

Since original versions of the files licensed by University of Aveiro could not be found, there is no way to determine if the files comply with their originals or if changes have been made. This should be considered while following the static analysis.

Following function calls in all three files led to the conclusion that the most important function is defined as `arithmetic_encode` in `arith.c`. In this function the actual arithmetic encoding is executed. This function has no redirects to other files, only one function call `ENCODE_RENORMALISE` the remaining code consists of arithmetic operations only [44].

Following function calls into the compressor section of `geco.c`, to find the call of `arith.c` no sign of multithreading could be identified. This fact leaves additional optimization possibilities and will be discussed at the end of this work.

2.4.2. Samtools

BAM

Compression in this format is done by an implementation called BGZF, which is a block compression on top of a widely used algorithm called DEFLATE.

DEFLATE The DEFLATE compression algorithm combines LZ77 and Huffman coding. It is used in well known tools like gzip. Data is split into blocks. Each block stores a header consisting of three bits. A single block can be stored in one of three forms. Each of which is represented by an identifier that is stored with the last two bits in the header.

- 00 No compression.

2. The Structure of the Human Genome and how its Digital Form is Compressed

- 01 Compressed with a fixed set of Huffman codes.
- 10 Compressed with dynamic Huffman codes.

The last combination 11 is reserved to mark a faulty block. The third, leading bit is set to flag the last data block [35]. As described in 2.3.1 a compression with LZ77 results in literals, a length for each literal and pointers that are represented by the distance between pointer and the literal it points to. The LZ77 algorithm is executed before the Huffman algorithm. Further compression steps differ from the already described algorithm and will extend to the end of this section.

Besides header bit and a data block, two Huffman code trees are stored. One encodes literals and lengths and the other distances. They happen to be in a compact form. This is achieved by a addition of two rules on top of the rules described in 2.3.3: Codes of identical lengths are ordered lexicographically, directed by the characters they represent. And the simple rule: shorter codes precede longer codes. To illustrate this with an example: For a text consisting out of C and G, following codes would be set, for an encoding of two bit per character: C: 00, G: 01. With another character A in the alphabet, which would occur more often than the other two characters, the codes would change to a representation like this:

Symbol	Huffman code
A	0
C	10
G	11

Since A precedes C and G, it is represented with a 0. To maintain prefix-free codes, the two remaining codes are not allowed to contain a leading 0. C precedes G lexicographically, therefore the (in a numerical sense) smaller code is set to represent C. With this simple rule, the alphabet can be compressed too. Instead of storing codes itself, only the code length is stored [35]. This might seem unnecessary when looking at a single compressed bulk of data, but when compressing blocks of data, a smaller alphabet can make a relevant difference.

BGZF extends this by creating a series of blocks. Each can not extend a limit of 64 Kilobyte. Each block contains a standard gzip file header, followed by compressed data.

2. The Structure of the Human Genome and how its Digital Form is Compressed

CRAM

The improvement of BAM [23] called CRAM, also features a block structure [20]. The whole file can be separated into four sections, stored in ascending order: File definition, a CRAM Header Container, multiple Data Container and a final CRAM EOF Container.

The complete structure is displayed in 2.11. The following paragraph will give a brief description to the high level view of a CRAM file, illustrated as the most upper bar. Followed by a closer look at the data container, which components are listed in the bar, at the center of 2.11. The most in depth explanation will be given to the bottom bar, which shows the structure of so called slices.

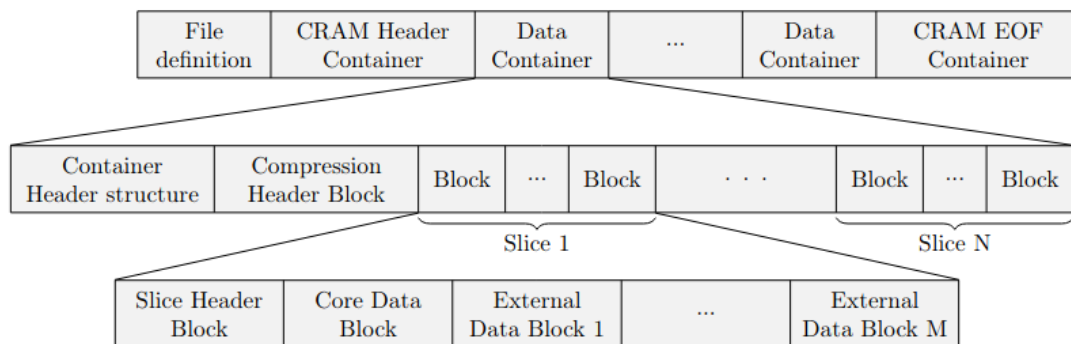


Figure 2.11.: CRAM file format structure [20].

The File definition, illustrated on the left side of the first bar in 2.11, consists of 26 uncompressed bytes, storing formatting information and a identifier. The CRAM header contains meta information about Data Containers and is optionally compressed with gzip. This container can also contain a uncompressed zero-padded section, reserved for SAM header information [20]. This saves time, in case the compressed file is altered and its compression need to be updated. The last container in a CRAM file serves as a indicator that the EOF is reached. Since in addition information about the file and its structure is stored, a maximum of 38 uncompressed bytes can be reached.

A Data Container can be split into three sections. From this sections the one storing the actual sequence consists of blocks itself, displayed in 2.11 as the bottom row.

- Container Header.

2. The Structure of the Human Genome and how its Digital Form is Compressed

- Compression Header.
- A variable amount of Slices.
 - Slice Header.
 - Core Data Block.
 - A variable amount of External Data Blocks.

The Container Header stores information on how to decompress the data stored in the following block sections. The Compression Header contains information about what kind of data is stored and some encoding information for SAM specific flags [20]. The actual data is stored in the Data Blocks. Those consist of encoded bit streams. According to the Samtools specification, the encoding can be one of the following: External, Huffman and two other methods which happen to be either a form of Huffman coding or a shortened binary representation of integers [20]. The External option allows to use gzip, bzip2 which is a form of multiple coding methods including run length encoding and Huffman, an encoding from the LZ family called LZMA or a combination of arithmetic and Huffman coding called rANS [13].

Chapter 3

Environment and Procedure to Determine the State of The Art Efficiency and Compressionratio of Relevant Tools

Since improvements must be measured, defining a baseline which would need to be beaten beforehand is necessary. Others have dealt with this task several times with common algorithms and tools, and published their results. But since the test case, that need to be build for this work, is rather uncommon in its compilation, the available data are not very useful. Therefore, new test data must be created.

The goal of this is, to determine a baseline for efficiency and effectivity of state of the art tools, used to compress DNA. This baseline is set by two important factors:

- Efficiency: **Duration** the Process had run for
- Effectivity: The difference in **Size** between input and compressed data

As a third point, the compliance that files were compressed losslessly should be verified. This is done by comparing the source file to a copy that got compressed and than decompressed again. If one of the two processes should operate lossy, a difference between the source file and the copy a difference in size should be recognizable.

3.1. Server specifications and test environment

To be able to recreate this in the future, relevant specifications and the commands that revealed this information are listed in this section.

Reading from `/proc/cpuinfo` reveals processor specifications. Since most of the information displayed in the seven entries is redundant, only the last entry is shown. Below are relevant specifications listed:

```
cat /proc/cpuinfo
```

- available logical processors: 0 - 7
- vendor: GenuineIntel
- cpu family: 6
- model nr, name: 58, Intel(R) Core(TM) i7-3770 CPU @ 3.40GHz
- microcode: 0x15
- MHz: 2280.874
- cache size: 8192 KB
- cpu cores: 4
- fpu and fpu exception: yes
- address sizes: 36 bits physical, 48 bits virtual

Full CPU specification can be found in 4.2.

The installed Random Access Memory (RAM) was offering a total of 16GB with four 4GB instances. For this paper relevant specifications are listed below: Command used to list

```
dmidecode --type 17
```

- Total/Data Width: 64 bits
- Size: 4GB
- Type: DDR3
- Type Detail: Synchronous
- Speed/Configured Memory Speed: 1600 Megatransfers/s

3.2. Operating System and Additionally Installed Packages

To leave the testing environment in a consistent state, not project specific processes running in the background, should be avoided. Due to following circumstances, a current Linux distribution was chosen as a suitable operating system:

- factors that interfere with a consistent efficiency value should be avoided
- packages, support and user experience should be present to an reasonable ammount

Some background processes will run while the compression analysis is done. This is owed to the demand of an increasingly complex operating system to execute complex programs. Considering that different tools will be exeuted in this environment, minimizing the background processes would require building a custom operating system or configuring an existing one to fit this specific use case. The boundary set by the time limitation for this work rejects named alternatives. Choosing **Debian GNU/Linux** version **11** features enough packages to run every tool without spending to much time on the setup.

The graphical user interface and most other optional packages were omitted. The only additional package added in the installation process is the ssh server package. Further a list of packages required by the compression tools were installed. At last, some additional packages were installed for the purpose of simplifying work processes and increasing the safety of the environment.

- installation process: ssh-server
- tool requirements: git, libhts-dev, autoconf, automake, cmake, make, gcc, perl, zlib1g-dev, libbz2-dev, liblzma-dev, libcurl4-gnutls-dev, libssl-dev, libncurses5-dev, libomp-dev
- additional packages: ufw, rsync, screen, sudo

A complete list of installed packages as well as individual versions can be found in the appendix.

3.3. Selection, Receivement, and Preperation of Testdata

Following criteria is required for test data to be appropriate:

3. Environment and Procedure to Determine the State of The Art Efficiency and Compressionratio of Relevant Tools

- The test file is in a format that all or at least most of the tools can work with, meaning FASTA or FASTq files.
- The file is publicly available and free to use (for research).

A second, bigger set of testfiles were required. This would verify the test results are not limited to small files. A minimum of one gigabyte of average filesize were set as a boundary. This corresponds to over five times the size of the first set.

Since there are multiple open File Transfere Protocol (FTP) servers which distribute a variety of files, finding a suitable first set is rather easy. The ensembl database featured defined criteria, so the first available set called:

`Homo_sapiens.GRCh38.dna.chromosome`

were chosen [26]. This sample includes 20 chromosomes, whereby considering the filenames, one chromosome is contained in each single file. After retrieving and unpacking the files, write privileges on them was withdrawn. So no tool could alter any file contents, without sufficient permission. Finding a second, bigger set happened to be more complicated. FTP offers no fast, reliable way to sort files according to their size, regardless of their position. Since available servers [24]–[26] offer several thousand files, stored in variating, deep directory structures, mapping filesize, filetype and file path takes too much time and resources for the scope of this work. This problematic combined with a easily triggered overflow in the samtools library, resulted in a set of several, manually searched and tested FASTq files. Compared to the first set, there is a noticable lack of quantity, but the filesizes happen to be of a fortunate distribution. With pairs of two files in the ranges of 0.6, 1.1, 1.2 and one file with a size of 1.3 gigabyte, effects on scaling sizes should be clearly visible.

Following tools and parameters where used in this process:

```
\$ wget http://ftp.ensembl.org/pub/release-107/fastq/homo_sapiens/dna/
    Homo_sapiens.GRCh38.dna.chromosome.{2,3,4,5,6,7,8,9,10}.fa.gz
\$ gzip -d ./*
\$ chmod -w ./*
```

The chosen tools are able to handle the FASTA format. However Samtools must convert FASTA files into their SAM format bevor the file can be compressed. The compression will firstly lead to an output with BAM format, from there it can be compressed further into a CRAM file. For CRAM compression, the time needed for each step, from converting to two compressions, is summed up and displayed as one. For the compression time into the BAM format, just the conversion and the

3. Environment and Procedure to Determine the State of The Art Efficiency and Compressionratio of Relevant Tools

single compression time is summed up. The conversion from FASTA to SAM is not displayed in the results. This is due to the fact that this is no compression process, and therefor has no value to this work.

Even though SAM files are not compressed, there can be a small but noticeable difference in size between the files in each format. Since FASTA should store less information, by leaving out quality scores, this observation was counterintuitive. Comparing the first few lines showed two things: the header line were altered and newlines were removed. The alteration of the header line would result in just a few more bytes. To verify, no information was lost while converting, both files were temporary stripped from metadata and formatting, so the raw data of both files can be compared. Using `diff` showed no differences between the stored characters in each file.

Chapter 4

Results and Discussion

The tables A and A contain raw measurement values for the two goals, described in 3. The table A lists how long each compression procedure took, in milliseconds. A contains file sizes in bytes. In these tables, as well as in the other ones associated with tests in the scope of this work, the a name scheme is used, to improve readability. The filenames were replaced by File followed by two numbers separated by a point. For the first test set, the number prefix 1. was used, the second set is marked with a 2.. For example, the fourth file of each test, in tables are named like this File 1.4 and File 2.4. The name of the associated source file for the first set is:

Homo_sapiens.GRCh38.dna.chromosome.4.fa

Since the source files of the second set are not named as consistent as in the first one, a third column in 4 was added, which is mapping table ID. and source file name.

The files contained in each test set, as well as their size can be found in the tables 4 and 4. The first test set contained a total of 2.8 GB unevenly spread over 21 files, while the second test set contained 7 GB in total, with a quantity of seven files.

Table 4.1.: Files contained in the First Test Set and their Sizes in **MB!**

ID.	Size in MB!
File 1.1	241.38
File 1.2	234.823
File 1.3	192.261
File 1.4	184.426
File 1.5	176.014
File 1.6	165.608

4. Results and Discussion

File 1.7	154.497
File 1.8	140.722
File 1.9	134.183
File 1.10	129.726
File 1.11	130.976
File 1.12	129.22
File 1.13	110.884
File 1.14	103.786
File 1.15	98.888
File 1.16	87.589
File 1.17	80.724
File 1.18	77.927
File 1.19	56.834
File 1.20	62.483
File 1.21	45.289

Table 4.2.: Files contained in the Second Test Set, their Sizes in **MB!** and Source File Names

ID.	Size in MB!	Source File Name
File 2.1	1188.976	SRR002905.recal.fastq
File 2.2	1203.314	SRR002906.recal.fastq
File 2.3	627.467	SRR002815.recal.fastq
File 2.4	676.0	SRR002816.recal.fastq
File 2.5	1066.431	SRR002817.recal.fastq
File 2.6	1071.095	SRR002818.recal.fastq
File 2.7	1240.564	SRR002819.recal.fastq

4.1. Interpretation of Results

The units milliseconds and bytes store a high precision. Unfortunately they are harder to read and compare, solely by the readers eyes. Therefore the data was altered. Sizes in 4.1 are displayed in percentage, in relation to the respective source file. Meaning the compression with GeCo on:

`Homo_sapiens.GRCh38.dna.chromosome.11.fa`

resulted in a compressed file which were only 17.6% as big. Runtimes in 4.1 were converted into seconds and have been rounded to two decimal places. Also a line was added to the bottom of each table, showing the average percentage or runtime for each process.

Table 4.3.: File sizes in different compression formats in **percent**

ID.	GeCo %	Samtools BAM%	Samtools CRAM %
File 1.1	18.32	24.51	22.03

4. Results and Discussion

File 1.2	20.28	26.56	23.57
File 1.3	20.4	26.58	23.66
File 1.4	20.3	26.61	23.56
File 1.5	20.12	26.46	23.65
File 1.6	20.36	26.61	23.6
File 1.7	19.64	26.15	23.71
File 1.8	20.4	26.5	23.67
File 1.9	17.01	23.25	20.94
File 1.10	20.15	26.36	23.7
File 1.11	19.96	26.14	23.69
File 1.12	20.1	26.26	23.74
File 1.13	17.8	22.76	20.27
File 1.14	17.16	22.31	20.11
File 1.15	16.21	21.69	19.76
File 1.16	17.43	23.48	21.66
File 1.17	18.76	25.16	23.84
File 1.18	20.0	25.31	23.63
File 1.19	17.6	24.53	23.91
File 1.20	19.96	25.6	23.67
File 1.21	16.64	22.06	20.44
Total	18.98	24.99	22.71

Overall, Samtools BAM resulted in 71.76% size reduction, the CRAM method improved this by roughly 2.5%. GeCo provided the greatest reduction with 78.53%. This gap of about 4% comes with a comparatively great sacrifice in time.

Table 4.4.: Compression duration in seconds

ID.	GeCo	Samtools BAM	Samtools CRAM
File 1.1	23.5	3.786	16.926
File 1.2	24.65	3.784	17.043
File 1.3	2.016	3.123	13.999
File 1.4	19.408	3.011	13.445
File 1.5	18.387	2.862	12.802
File 1.6	17.364	2.685	12.015
File 1.7	15.999	2.503	11.198
File 1.8	14.828	2.286	10.244
File 1.9	12.304	2.078	9.21
File 1.10	13.493	2.127	9.461
File 1.11	13.629	2.132	9.508
File 1.12	13.493	2.115	9.456
File 1.13	99.902	1.695	7.533
File 1.14	92.475	1.592	7.011
File 1.15	85.255	1.507	6.598
File 1.16	82.765	1.39	6.089
File 1.17	82.081	1.306	5.791
File 1.18	79.842	1.277	5.603

4. Results and Discussion

File 1.19	58.605	0.96	4.106
File 1.20	64.588	1.026	4.507
File 1.21	41.198	0.721	3.096
Total	42.57	2.09	9.32

As 4.1 is showing, the average compression duration for GeCo is at 42.57s. That is a little over 33s, or 78% longer than the average runtime of samtools for compressing into the CRAM format.

Since CRAM requires a file in BAM format, the third row is calculated by adding the time needed to compress into BAM with the time needed to compress into CRAM. While SAM format is required for compressing a FASTA into BAM and further into CRAM, in itself it does not features no compression. However, the conversion from SAM to FASTA can result in a decrease in size. At first this might be contra intuitive since, as described in 2.2.1 SAM stores more information than FASTA. This can be explained by comparing the sequence storing mechanism. A FASTA sequence section can be spread over multiple lines whereas SAM files store a sequence in just one line, converting can result in a SAM file that is smaller than the original FASTA file. Before interpreting this data further, a quick view into development processes: GeCo stopped development in the year 2016 while Samtools is being developed since 2015, to this day, with over 70 people contributing.

Reviewing 4.1 one will notice, that GeCo reached a runtime over 60 seconds on every run. Instead of displaying the runtime solely in seconds, a leading number followed by an m indicates how many minutes each run took.

Table 4.5.: File sizes in different compression formats in **percent**

ID.	GeCo%	Samtools BAM%	Samtools CRAM%
File 2.1	1.00	6.28	5.38
File 2.2	0.98	6.41	5.52
File 2.3	1.21	8.09	7.17
File 2.4	1.20	7.70	6.85
File 2.5	1.08	7.58	6.72
File 2.6	1.09	7.85	6.93
File 2.7	0.96	5.83	4.63
Total	1.07	7.11	6.17

Table 4.6.: Compression duration in seconds

ID.	GeCo	Samtools BAM	Samtools CRAM
File 2.1	1m58.427	16.248	23.016
File 2.2	1m57.905	15.770	22.892
File 2.3	1m09.725	07.732	12.858
File 2.4	1m13.694	08.291	13.649
File 2.5	1m51.001	14.754	23.713
File 2.6	1m51.315	15.142	24.358
File 2.7	2m02.065	16.379	23.484
Total	1m43.447	13.474	20.567

In both tables 4.1 and 4.1 the already identified pattern can be observed. Looking at the compression ratio in 4.1 a maximum compression of 99.04% was reached with GeCo. In this set of test files, file seven were the one with the greatest size (1.3 Gigabyte). Closely folled by file one and two (1.2 Gigabyte).

4.2. View on Possible Improvements

So far, this work went over formats for storing genomes, methods to compress files (in mentioned formats) and through tests where implementations of named algorithms compress several files and analyzed the results. The test results show that GeCo provides a better compression ratio than Samtools and takes more time to run through. So in this testrun, implementations of arithmetic coding resulted in a better compression ratio than Samtools BAM with the mix of huffman coding and LZ77, or Samtools custom compression format CRAM. Comparing results in [5], supports this statement. This study used FASTA/Multi-FASTA files from 71MB to 166MB and found that GeCo had a variating compression ratio from 12.34 to 91.68 times smaller than the input reference and also resulted in long runtimes up to over 600 minutes [5]. Since this study focused on another goal than this work and therefore used different test variables and environments, the results can not be compared. But what can be taken from this, is that arithmetic coding, at least in GeCo is in need of a runtime improvement.

The actual mathematical proove of such an improvemnt and its implementation can not be covered because it would to beyond scope. But in order to set up a foundation for this task, the rest of this work will consist of considerations and problem analysis, which should be thought about and dealt with to develop a improvement.

S.V. Petoukhov described his findings about the distribution of nucleotides [8]. With the probability of one nucleotide, in a sequence of sufficient length, information about the direct neighbours is revealed. For example, with the probability of C, the probabilities for sets (n-plets) of any nucleotide N, including C can be determined without counting them [8].

Considering this and the measured results, an improvement in the arithmetic coding process and therefore in GeCos efficiency, would be a good start to equalize the great gap in the compression duration. Combined with a tool that is developed with today's standards, there is a possibility that even greater improvements could be archived.

How would a theoretical improvement approach look like? As described in 2.3.2, entropy coding requires to determine the probabilities of each symbol in the alphabet. The simplest way to do that, is done by parsing the whole sequence from start to end and increasing a counter for each nucleotide that got parsed. With new findings discovered by Petoukhov in consideration, the goal would be to create an entropy coding implementation that beats current implementation in the time needed to determine probabilities. A possible approach would be that the probability of one nucleotide can be used to determine the probability of other nucleotides, by a calculation rather than the process of counting each one. This approach throws a few questions that need to be answered in order to plan a implementation [8]:

- How many probabilities are needed to calculate the others?
- Is there space for improvement in the parsing/counting process?
- How can the variation between probabilities be determined?

The question for how many probabilities are needed, needs to be answered, to start working on any kind of implementation. This question will only get answered by theoretical prove. It could happen in form of a mathematical equation, which proves that counting all occurrences of one nucleotide reveals can be used to determine all probabilities. Since this task is time and resource consuming and there is more to discuss, finding an answer will be postponed to another work.

The Second point must be asked, because the improvement in counting only one nucleotide in comparison to counting three, would be too little to be called relevant. Especially if multithreading is an option. Since in the static code analysis in 2.4

revealed no multithreading, the analysis for improvements when splitting the workload onto several threads should be considered, before working on an improvement based on Petoukhov's findings. This is relevant, because some improvements, like the one described above, will lose efficiency if only subsections of a genome are processed. A tool like OpenMC for multithreading C programs would possibly supply the required functionality to develop a proof of concept [8], [45]. But how could an improvement look like, not considering possible difficulties multithreading would bring? To answer this, first a mechanism to determine a possible improvement must be determined. To compare parts of a program and their complexity, the Big-O notation is used. Unfortunately this is only covering loops and conditions as a whole. Therefore a more detailed view on operations must be created: Considering a single threaded loop with the purpose to count every nucleotide in a sequence, the process of counting can be split into several operations, defined by this pseudocode.

```
while (sequence not end)
do
    next_nucleotide = read_next_nucleotide(sequence)
    for (element in alphabet_probabilities)
    do
        if (element equals next_nucleotide)
            element = element + 1
        fi
    done
done
```

This loop will iterate over a whole sequence, counting each nucleotide. In line three, an inner loop can be found which iterates over the alphabet, to determine which symbol should be increased. Considering the findings, described above, the inner loop can be left out, because there is no need to compare the read nucleotide against more than one symbol. The Big-O notation for this code, with any sequence with the length of n , would be decreased from $O(n^2)$ to $O(n \cdot 1)$ or simply $O(n)$ [46]. Which is clearly an improvement in complexity and therefore also in runtime. The runtime for calculations of the other symbols' probabilities must be considered as well and compared against the nested loop to be certain, that the overall runtime was improved.

Getting back to the question how multithreading would impact improvements: A implementation like the one described above, could also work with multithreading. Since the ratio of the difference between $O(n^2)$ and $O(n)$ does not differ with the reduction of n . Multiple threads, processing parts of a sequence with the length of n , would also benefit, because any fraction of n^2 will always be greater than the corresponding fraction of n . This results can either be summed up for global probabilities or get used individually on each associated subsequence. Either way, the presented improvement approach should be applicable to both parsing methods.

This leaves a list of problems, which needs to be regarded in the approach of developing an improvement. If there is space for improvement in the parsing/counting process, what problems need to be addressed:

- reducing one process by adding additional code must be estimated and set into relation.
- for a tool that does not feature multithreading, how would multithreading affect the improvement result?

An important question that needs to be answered would be: If Petoukhov's findings show that, through similarities in the distribution of each nucleotide, one can lead to the approximation of the other three. Entropy codings work with probabilities, how does that affect the coding mechanism? With an equal probability for each nucleotide, entropy coding can not be treated as a whole. This is due to the fact, that Huffman coding makes use of differing probabilities. An equal distribution means every character will be encoded in the same length which would make the encoding process unnecessary. Arithmetic coding on the other hand is able to handle equal probabilities. The fact that there are obviously chains of repeating nucleotides in genomes. For example File 2.2, which contains this subsequence is found at line 90:

```
AAAAAAAAAAAAAAAAAAAAAAAAATAAATATTTTATTT
```

Without determining probabilities, one can see that the amount of As outnumbers Ts and neither C nor G are present. With the whole 1.2 gigabytes, the distribution will align more, but by cutting out a subsection, of relevant size, with unequal distributions will have an impact on the probabilities of the whole sequence. If a greater sequence would lead to a more equal distribution, this knowledge could be used to help determining distributions on subsequences of one with equally distributed probabilities.

List of Abbreviations

ASCII American Standard Code for Information Interchange
BAM Binary Alignment Map
CRAM Compressed Reference-oriented Alignment Map
DNA Deoxyribonucleic Acid
EOF End of File
FASTA File Format for Storing Genomic Data
FASTq File Format Based on FASTA
FTP File Transfere Protocol
GB Gigabyte
GeCo Genome Compressor
GPL GNU General Public License
IUPAC International Union of Pure and Applied Chemistry
LZ77 Lempel Ziv 1977
MIT Massachusetts Institute of Technology
RAM Random Access Memory
SAM Sequence Alignment Map
UDP Universal Datagram Protocol
UTF Unicode Transformation Format

List of Tables

2.1. Probability distribution for arithmetic encoding example shown in 2.8	16
2.2. Huffman example table no. 1 (pre encoding)	19
2.3. Huffman example table no. 2 (post encoding)	21
4.1. First Test Set Files and their Sizes in MB	32
4.2. Second Test Set Files and their Sizes in MB	33
4.3. Compression Effectivity for first test set	33
4.4. Compression Efficiency for first test set	34
4.5. Compression Effectivity for second test set	35
4.6. Compression Effectivity for greater files	36
A.1. Compression efficiency in milliseconds	xvi
A.2. Compression effectivity in byte	xvi

List of Figures

2.1.	A superficial representation of the physical positioning of genomes. Showing a double helix (bottom), a chromosome (upper right) and a cell (upper center).	3
2.2.	A purely diagrammatic figure of the components DNA is made of. The smaller, inner rods symbolize nucleotide links and the outer ribbons the phosphate-sugar chains [3].	4
2.3.	Edited example of a FASTA file. The original was received from the ensemble server [26].	7
2.4.	Altered example of a FASTq file. The original was received from ncbi server [25].	9
2.5.	SAM file structure example [20].	9
2.6.	Schematic sketch, illustrating the replacement of multiple occurrences done in dictionary coding.	12
2.7.	Schematic diagram of a general communication system by Shannons definition. [6]	13
2.8.	Illustrative example of arithmetic coding.	15
2.9.	Illustrative rescaling in arithmetic coding process.	17
2.10.	Final version of the Huffman tree for described example.	21
2.11.	CRAM file format structure [20].	25

Listings

Bibliography

- [1] Philomin Juliana, Ravi Prakash Singh, Jesse Poland, *et al.*, “Elucidating the genetics of grain yield and stress-resilience in bread wheat using a large-scale genome-wide association mapping study with 55,568 lines”, *Scientific Reports*, vol. 11, no. 1, Mar. 2021. DOI: 10.1038/s41598-021-84308-4.
- [2] Arno G. Motulsky, “Impact of genetic manipulation on society and medicine”, *Science*, vol. 219, no. 4581, pp. 135–140, Jan. 1983. DOI: 10.1126/science.6336852.
- [3] J. D. WATSON and F. H. C. CRICK, “Molecular structure of nucleic acids: A structure for deoxyribose nucleic acid”, *Nature*, vol. 171, no. 4356, pp. 737–738, Apr. 1953. DOI: 10.1038/171737a0.
- [4] Anas Al-Okaily, Badar Almarri, Sultan Al Yami, and Chun-Hsi Huang, “Toward a better compression for DNA sequences using huffman encoding”, *Journal of Computational Biology*, vol. 24, no. 4, pp. 280–288, Apr. 1, 2017. DOI: 10.1089/cmb.2016.0151.
- [5] Morteza Hosseini, Diogo Pratas, and Armando Pinho, “A survey on data compression methods for biological sequences”, *Information*, vol. 7, no. 4, p. 56, Oct. 2016. DOI: 10.3390/info7040056.
- [6] C. E. Shannon, “A mathematical theory of communication”, *Bell System Technical Journal*, vol. 27, no. 3, pp. 379–423, Jul. 1948. DOI: 10.1002/j.1538-7305.1948.tb01338.x.
- [7] Alistair Moffat, Radford M. Neal, and Ian H. Witten, “Arithmetic coding revisited”, *ACM Transactions on Information Systems*, vol. 16, no. 3, pp. 256–294, Jul. 1998. DOI: 10.1145/290159.290162.
- [8] Sergey V. Petoukhov, “Tensor rules in the stochastic organization of genomes and genetic stochastic resonance in algebraic biology”, Oct. 2021. DOI: 10.20944/preprints202110.0093.v1.

- [9] Eva Bianconi, Allison Piovesan, Federica Facchin, *et al.*, “An estimation of the number of cells in the human body”, *Annals of Human Biology*, vol. 40, no. 6, pp. 463–471, Jul. 2013. DOI: 10.3109/03014460.2013.807878.
- [10] ISO/IEC JTC 1/SC 2 Coded character sets, “Information technology — 8-bit single-byte coded graphic character sets — part 1: Latin alphabet no. 1”, International Organization for Standardization, Geneva, CH, Standard, Apr. 1998.
- [11] Peter J. A. Cock, Christopher J. Fields, Naohisa Goto, Michael L. Heuer, and Peter M. Rice, “The sanger FASTQ file format for sequences with quality scores, and the solexa/illumina FASTQ variants”, *Nucleic Acids Research*, vol. 38, no. 6, pp. 1767–1771, Dec. 2009. DOI: 10.1093/nar/gkp1137.
- [12] Wai Yee Low, Rick Tearle, Ruijie Liu, *et al.*, “Haplotype-resolved genomes provide insights into structural variation and gene content in angus and brahman cattle”, *Nature Communications*, vol. 11, no. 1, Apr. 2020. DOI: 10.1038/s41467-020-15848-y.
- [13] Petr Danecek, James K Bonfield, Jennifer Liddle, *et al.*, “Twelve years of SAMtools and BCFtools”, *GigaScience*, vol. 10, no. 2, Jan. 2021. DOI: 10.1093/gigascience/giab008.
- [14] ISO Central Secretary, “Mpge-g”, en, International Organization for Standardization, Standard ISO/IEC 23092-1:2020, Oct. 2020. [Online]. Available: <https://www.iso.org/standard/23092.html>.
- [15] Claudio Albert, Tom Paridaens, Jan Voges, *et al.*, “An introduction to MPEG-g, the new ISO standard for genomic information representation”, Sep. 2018. DOI: 10.1101/426353.
- [16] Andrew D. Johnson, “An extended IUPAC nomenclature code for polymorphic nucleic acids”, *Bioinformatics*, vol. 26, no. 10, pp. 1386–1389, Mar. 2010. DOI: 10.1093/bioinformatics/btq098.
- [17] Paul Flicek. “Ensembl project”. (Oct. 24, 2022), [Online]. Available: <http://www.ensembl.org/>.
- [18] Santa Cruz UCSC - University of California. “Ucsc genome browser”. (Oct. 28, 2022), [Online]. Available: <https://genome.ucsc.edu/> (visited on 10/28/2022).
- [19] “Global alliance for genomics and health”. (Oct. 10, 2022), [Online]. Available: <https://github.com/samtools/hts-specs..>

- [20] The SAM/BAM Format Specification Working Group. “Sequence alignment/map format specification”. version 44b4167. (Aug. 22, 2022), [Online]. Available: <https://github.com/samtools/hts-specs> (visited on 09/12/2022).
- [21] Diogo Pratas, Armando J. Pinho, and Paulo J. S. G. Ferreira, “Efficient compression of genomic sequences”, in *2016 Data Compression Conference (DCC)*, IEEE, Mar. 2016. DOI: 10.1109/DCC.2016.60.
- [22] Kelvin V. Kredens, Juliano V. Martins, Osmar B. Dordal, *et al.*, “Vertical lossless genomic data compression tools for assembled genomes: A systematic literature review”, *PLOS ONE*, vol. 15, no. 5, Rashid Mehmood, Ed., e0232942, May 2020. DOI: 10.1371/journal.pone.0232942.
- [23] Markus Hsi-Yang Fritz, Rasko Leinonen, Guy Cochrane, and Ewan Birney, “Efficient storage of high throughput DNA sequencing data using reference-based compression”, *Genome Research*, vol. 21, no. 5, pp. 734–740, Jan. 2011. DOI: 10.1101/gr.114819.110.
- [24] “Igsr: The international genome sample resource”. (Nov. 10, 2022), [Online]. Available: <https://ftp.1000genomes.ebi.ac.uk>.
- [25] “Ncbi national center for biotechnology information”. (Nov. 1, 2022), [Online]. Available: <https://ftp.ncbi.nlm.nih.gov/genomes/>.
- [26] “Ensembl rapid release”. (Oct. 15, 2022), [Online]. Available: <https://ftp.ensembl.org>.
- [27] K. Simonsen and, “Character mnemonics and character sets”, RFC 1345, Jun. 1992.
- [28] Colin McIntosh, *Cambridge International Dictionary of English*. Cambridge University Press, 2013, p. 1856.
- [29] J. Postel, “User datagram protocol”, RFC Editor, Tech. Rep. 768, Aug. 28, 1980, 3 pp. DOI: 10.17487/RFC0768. [Online]. Available: <https://www.rfc-editor.org/info/rfc768>.
- [30] Manish RajShivare, Yogendra P. S. Maravi, and Sanjeev Sharma, “Analysis of header compression techniques for networks: A review”, *International Journal of Computer Applications*, vol. 80, no. 5, pp. 13–20, Oct. 2013. DOI: 10.5120/13856-1701.
- [31] ISO, “Iso/iec 10646:2020 utf”, Tech. Rep.

- [32] Kashfia Sailunaz, Mohammed Rokibul Alam Kotwal, and Mohammad Nurul Huda, “Data compression considering text files”, *International Journal of Computer Applications*, vol. 90, no. 11, pp. 27–32, Mar. 2014. DOI: 10.5120/15765-4456.
- [33] Alistair Moffat, “Huffman coding”, *ACM Computing Surveys*, vol. 52, no. 4, pp. 1–35, Jul. 2020. DOI: 10.1145/3342555.
- [34] J. Ziv and A. Lempel, “A universal algorithm for sequential data compression”, *IEEE Transactions on Information Theory*, vol. 23, no. 3, pp. 337–343, 1977. DOI: 10.1109/TIT.1977.1055714.
- [35] L Peter Deutsch, “DEFLATE compressed data format specification version 1.3”, Tech. Rep., May 1996. DOI: 10.17487/rfc1951. [Online]. Available: <https://www.rfc-editor.org/rfc/rfc1951>.
- [36] Hans Delfs and Helmut Knebl, *Introduction to Cryptography, Principles and Applications (Information Security and Cryptography)*. Springer, 2007, p. 368.
- [37] J. J. Rissanen, “Generalized kraft inequality and arithmetic coding”, *IBM Journal of Research and Development*, vol. 20, no. 3, pp. 198–203, May 1976. DOI: 10.1147/rd.203.0198.
- [38] “Ieee standard for floating-point arithmetic”, *IEEE Std 754-2019 (Revision of IEEE 754-2008)*, pp. 1–84, 2019. DOI: 10.1109/IEEESTD.2019.8766229.
- [39] Ian H. Witten, Radford M. Neal, and John G. Cleary, “Arithmetic coding for data compression”, *Communications of the ACM*, vol. 30, no. 6, pp. 520–540, Jun. 1987. DOI: 10.1145/214762.214771. [Online]. Available: <https://doi.org/10.1145/214762.214771>.
- [40] L. Peter Deutsch, Jean-Loup Gailly, Mark Adler, L. Peter Deutsch, and Glenn Randers-Pehrson, “Gzip file format specification version 4.3”, RFC 1952, May 1996.
- [41] David A. Huffman, “A method for the construction of minimum-redundancy codes”, *Proceedings of the Institute of Radio Engineers*, vol. 40, no. 9, pp. 1098–1101, Sep. 1952.
- [42] “Gnu public license”. (), [Online]. Available: <http://www.gnu.org/licenses/gpl-3.0.html>.
- [43] “Mit license”. (), [Online]. Available: <https://spdx.org/licenses/MIT.html>.

- [44] Cobilab. “Repositories for the three versions of geco”. (Nov. 19, 2022), [Online]. Available: <https://github.com/cobilab>.
- [45] Michael J. Quinn, *Parallel Programming in C with MPI and OpenMP*. McGraw-Hill Education Group, 2003.
- [46] Firdous Mala and Rouf Ali, “The big-o of mathematics and computer science”, vol. 6, pp. 1–3, Jan. 2022. DOI: 10.26855/jamc.2022.03.001.

CPU specification. Due to redundance, the information is limited to the last core, beginning at: processor : 7

```
cat /proc/cpuinfo
```

```
processor : 0
```

```
...
```

```
processor : 7
```

```
vendor_id : GenuineIntel
```

```
cpu family : 6
```

```
model : 58
```

```
model name : Intel(R) Core(TM) i7-3770 CPU @ 3.40GHz
```

```
stepping : 9
```

```
microcode : 0x15
```

```
cpu MHz : 2412.891
```

```
cache size : 8192 KB
```

```
physical id : 0
```

```
siblings : 8
```

```
core id : 3
```

```
cpu cores : 4
```

```
apicid : 7
```

```
initial apicid : 7
```

```
fpu : yes
```

```
fpu_exception : yes
```

```
cpuid level : 13
```

```
wp : yes
```

```
flags : fpu vme de pse tsc msr pae mce cx8 apic sep mtrr pge mca cmov pat pse36
```

```
clflush dts acpi mmx fxsr sse sse2 ss ht tm pbe syscall nx rdtscp lm constant_tsc
```

```
arch_perfmon pebs bts rep_good nopl xtopology nonstop_tsc cpuid aperfmperf pni
```

```
pclmulqdq dtes64 monitor ds_cpl vmx smx est tm2 ssse3 cx16 xtpr pdcm pcid
```

```
sse4_1 sse4_2 x2apic popcnt tsc_deadline_timer aes xsave avx f16c rdrand lahf_lm
```

```
cpuid_fault epb pti tpr_shadow vnmi flexpriority ept vpid fsgsbase smep erms xsaveopt
```

```
dtherm ida arat pln pts
```

```
vmx flags : vnmi preemption_timer invvpid ept_x_only flexpriority tsc_offset vtptr
```

```
mtf vapid ept vpid unrestricted_guest
```

```
bugs : cpu_meltdown spectre_v1 spectre_v2 spec_store_bypass l1tf mds swapgs
```

```
itlb_multihit srbds mmio_unknown
```

```
bogomips : 6784.56
```

```
clflush size : 64
```

```
cache_alignment : 64
```

```
address sizes : 36 bits physical, 48 bits virtual
```

```
power management:
```

manually installed packages: autoconf

automake

bzip2

cmake

gcc

git

htop

libbz2-dev

libcurl4-gnutls-dev

libhts-dev

libhtscodecs2

liblzma-dev

libncurses5-dev

libomp-dev

libssl-dev

zlib1g-dev

openssh-client

perl

rsync

screen

sudo

ufw

vim

wget

Appendix A

Erster Anhang: Lange Tabelle

Table A.1.: Compression duration of various tools, measured in milliseconds

ID.	GeCo	Samtools BAM	Samtools CRAM
File 1.1	235005	3786	16926
File 1.2	246503	3784	17043
File 1.3	20169	3123	13999
File 1.4	194081	3011	13445
File 1.5	183878	2862	12802
File 1.6	173646	2685	12015
File 1.7	159999	2503	11198
File 1.8	148288	2286	10244
File 1.9	12304	2078	9210
File 1.10	134937	2127	9461
File 1.11	136299	2132	9508
File 1.12	134932	2115	9456
File 1.13	999022	1695	7533
File 1.14	924753	1592	7011
File 1.15	852555	1507	6598
File 1.16	827651	1390	6089
File 1.17	820814	1306	5791
File 1.18	798429	1277	5603
File 1.19	586058	960	4106
File 1.20	645884	1026	4507
File 1.21	411984	721	3096
File 2.1	58427	16248	23016
File 2.2	57905	15770	22892
File 2.3	9725	7732	12858
File 2.4	13694	8291	13649
File 2.5	51001	14754	23713
File 2.6	51315	15142	24358
File 2.7	2065	16379	23484

Table A.2.: File sizes for different formats in byte

ID.	Uncompressed Source File	GeCo	Samtools BAM	Samtools CRAM
File 2.1	253105752	46364770	62048289	55769827
File 2.2	246230144	49938168	65391181	58026123
File 2.3	201600541	41117340	53586949	47707954
File 2.4	193384854	39248276	51457814	45564837

A. Erster Anhang: Lange Tabelle

File 2.5	184563953	37133480	48838053	43655371
File 2.6	173652802	35355184	46216304	40980906
File 2.7	162001796	31813760	42371043	38417108
File 2.8	147557670	30104816	39107538	34926945
File 2.9	140701352	23932541	32708272	29459829
File 2.10	136027438	27411806	35855955	32238052
File 2.11	137338124	27408185	35894133	32529673
File 2.12	135496623	27231126	35580843	32166751
File 2.13	116270459	20696778	26467775	23568321
File 2.14	108827838	18676723	24284901	21887811
File 2.15	103691101	16804782	22486646	20493276
File 2.16	91844042	16005173	21568790	19895937
File 2.17	84645123	15877526	21294270	20177456
File 2.18	81712897	16344067	20684650	19310998
File 2.19	59594634	10488207	14616042	14251243
File 2.20	65518294	13074402	16769658	15510100
File 2.21	47488540	7900773	10477999	9708258
File 2.1	1246731616	12414797	78260121	67130756
File 2.2	1261766002	12363734	80895953	69649632
File 2.3	657946854	7966180	53201724	47175349
File 2.4	708837816	8499132	54569686	48521201
File 2.5	1118234394	12088239	84764250	75118457
File 2.6	1123124224	12265535	88147227	77826446
File 2.7	1300825946	12450651	75860986	60239362
