

Whole-genome sequencing and analysis of *Apocynum cannabinum*

Guoqi Li (✉ guoqilee@163.com)

Ningxia University School of Ecology and Environment <https://orcid.org/0000-0002-6176-3621>

Lixiao Song

Ningxia University School of Ecology and Environment

Jinfeng Che

Ningxia University School of Ecology and Environment

Yanyun Chen

Ningxia University School of Ecology and Environment

Research Article

Keywords: *Apocynum cannabinum*, whole-genome sequencing, 10X genomics linked reads, Hi-C, speciation

Posted Date: April 7th, 2023

DOI: <https://doi.org/10.21203/rs.3.rs-2663915/v1>

License: © ⓘ This work is licensed under a Creative Commons Attribution 4.0 International License.
[Read Full License](#)

Abstract

Background

Apocynum cannabinum is an important plant resource from the Apocynaceae family. However, the lack of complete genome information has severely impeded research progress of molecular biology research in this plant. Whole-genome sequencing can provide an in-depth understanding of species growth, development, and evolutionary origin, and is the most effective method for scientifically exploring the ecological and economic value of a plant.

Methods and results

In this study, we employed Illumina HiSeq, single-molecule real-time sequencing, 10X genomics linked reads, and chromatin interaction (Hi-C), a new assembly technique, to successfully assemble the whole draft genome for *A. cannabinum* (260 Mb). The super-scaffold N50 genome size from the Hi-C assisted assembly was 21.16 Mb and was anchored to 11 chromosome, resulting in a high-quality reference genome at the chromosome level ($2n = 2x = 22$). We further annotated, analyzed, and predicted 22,793 protein-coding genes, of which the functions of 95.6% were already annotated, 92.3% contained conserved protein domains, and 78.7% were aligned to known metabolic pathways.

Conclusions

This high-quality *A. cannabinum* genome can be used to analyze growth and development and evaluate gene evolution at the genome level, as well as assist in the comparative genomics and genetic modification of other important medicinal plants in Apocynaceae. Comparative analysis of the gene families showed that *A. cannabinum* speciated around 35.8 (27.0–46.9) million years ago.

Introduction

Apocynum venetum L is a perennial herb plant or small shrub that is widely distributed in saline-alkali soils of arid semi-arid area. It has a broad ecological distribution in Northwestern deserts in China and is a multi-purpose plant with ecological, economic, and medicinal value [1, 2]. Because of its excellent ecological and economic value, *A. venetum* has a promising future of industrialization[3]. To increase *Apocynum's* genetic diversities and ecological and economic value, our laboratory, with the support of project from the State Forestry Administration of the People's Republic of China, introduced *Apocynum cannabinum* from the abroad in 2004 and successfully cultivated the species in 2006. Preliminary studies showed that *A. cannabinum* has better drought resistance, thicker and straighter stalks, larger leaves, and fewer bifurcations than *A. venetum* [4]. *A. cannabinum* is also known as dogbane or prairie dogbane and it is mostly produced in temperate and subtropical regions in North America [5]. *A. cannabinum* has strong adaptability and has high early succession rate [6, 7]. Besides traditional fibers, *A. cannabinum* also contains cardiac glycosides, resin, volatile oils, rubber, tannin, and starch [8]. Native American traditional medicine also uses *A. cannabinum* roots to treat heart disease and ascites [9].

Apocynum is a genus of Apocynaceae. There are 44–46 genera and 145–176 species of Apocynaceae in China. They are mainly distributed in provinces South of the Yangtze River and coastal islands, and a few in the North and Northwest [10]. Because Apocynaceae plants represent unique evolutionary groups and their important economic value, the systematic status of oleander family and the relationship between families and genera have been the research focus of phylogenetic scientist [11].

With the development and maturity of sequencing technology, an increasing number of plant genomics studies had been initiated to elaborate the molecular mechanism of plant growth, development, evolution and origin at the genome level [12]. The second-generation and third-generation sequencing, containing 10X genomics linked reads and Hi-C, made many complex plant genomes sequencing possible [13]. The technology of 10X genomics linked reads is essentially to introduce the barcode sequence into the long sequence segment, and then used them to construct more accurate super-scaffolds to explore more detailed genetic information [14]. Hi-C sequencing method divides second-, third- generation or optical map-assisted assembled draft genome sequences into chromosome sets to determine the order and orientation of the various sequences on the chromosomes [15]. Therefore, the combination of 10X genomics linked reads and Hi-C three-dimensional conformation capture technology greatly increases the accuracy of whole-genome sequencing and assembly quality. Although there are a lot of whole sequencing of economic plants and crop plants have been finished, the sequencing data about *A. cannabinum* are insufficient relatively, which hinder the further study of *Apocynum*, especially in pharmacology, molecular biology and genetic evolution. To evaluate the genome composition, component analysis, metabolic pathways, phylogenetic relationships, and speciation time of *A. cannabinum*, we carried out whole-genome sequencing of *A. cannabinum*, which will be helpful to elucidate its stress resistance mechanisms at the molecular level ultimately. Our findings provide theoretical support for the subsequent screening of superior genes, variety improvement, genetic engineering, and the scientific utilization of *A. cannabinum*.

Materials And Methods

Experimental materials

The experimental material was cultivated in pots after the belowground roots had been collected from the *A. cannabinum* experimental base in Pingluo county, Shizuishan city, Ningxia Province in end-March 2021. In mid-April 2021, healthy plants showing good growth were selected, and young leaves and stems at the apex were collected, snap-frozen in liquid nitrogen, and stored in a – 80°C freezer for subsequent experiments.

Experimental methods

Sample extraction and measurements

Whole-genome sequencing requires high DNA integrity and purity, and thus the modified cetyltrimethylammonium bromide (CTAB) method was used to extract genomic DNA from *A. cannabinum*

[16]. Agarose gel electrophoresis and NanoDrop 2000 were used to measure the integrity and concentration of the template, respectively, with the requirement of $A_{260}/A_{280} \geq 1.80$.

Library construction and sequencing

Qualified DNA samples were fragmented using a Covaris ultrasonicator. Following this, magnetic beads were used for the enrichment and purification of large DNA fragments. This was followed by end terminal repair, addition of A-tails, addition of sequencing adapters, purification, and PCR amplification. We constructed two libraries of PacBio, 10X genomics linked reads, Hi-C and second-generation small fragment libraries, respectively, and sequenced by Illumina Hiseq and PacBio RSII sequencing platforms.

Genome assembly

PacBio genome assembly

First, self-correction of the PacBio data was performed whereby the reads were aligned with each other. Self-correction was conducted according to the insertion and deletion of bases and the probability of sequencing errors to obtain pre-assembly reads. Following this, data assembly was carried out. As third-generation data reads are long (mean read length: 10–15 kb, longest read > 40 kb), the Overlap-Layout-Consensus [17] method was used, i.e., the overlapping relationship of the reads was used for splicing to obtain a consensus sequence. Finally, the software Pilon [18] was used for another round of correction using the second-generation data of the assembly results from the previous step in order to increase the accuracy of the results, ultimately obtaining high-quality consensus sequences.

10X Genomics-assisted genome assembly

In order to obtain high-quality genome assembly sequences, 10X Genomics-assisted assembly was also used in addition to the third-generation PacBio data assembly, mainly using fragscaff (<https://sourceforge.net/projects/fragscaff/>) software. Linked reads obtained from the 10X genomics linked reads library sequencing were used for alignment with the consensus sequence obtained from the third-generation assembly results. Linked reads were assembled on the original basis to form scaffolding [19].

Hi-C-assisted genome assembly

First, quality control and adapter removal were performed on the generated data to obtain high-quality clean reads. Following this, data were aligned to the assembled genome sequence, and PCR repeats were removed. The interaction data obtained after noise correction were used to construct a chromosome interaction matrix. Based on the structural characteristics of the chromosomes in three-dimensional space, suitable clustering models were selected to anchor non-localized scaffolds on the chromosomes, and the correct order and direction of the scaffolds were determined by the corresponding sorting algorithm, and the whole genome sequence at the chromosome level was assembled with the software LACHESIS(<https://sourceforge.net/projects/>

Evaluation of *A. cannabinum* genome assembly quality

CEGMA (Core Eukaryotic Genes Mapping Approach) and BUSCO (Benchmarking Universal Single-Copy Orthologs) [20] were used to evaluate the assembly quality of the *A. cannabinum* genome. CEGMA evaluation selects conserved genes (248 genes) from six eukaryotic model organisms to construct a core gene library and uses tblastn, genewise, and geneid for genome evaluation. In BUSCO, 1440 orthologous single-copy genes from the plant database were used to evaluate the integrity of the assembled genome. To evaluate the assembly accuracy, small fragment library reads were aligned to the assembled genome using BWA software [21], and the distribution of the alignment rate, genome coverage, and depth of coverage of the reads were calculated to assess assembly integrity and sequencing uniformity.

Genome annotation analysis

Repeat sequence annotation

The default parameters of LTR_FINDER [22](http://tlife.fudan.edu.cn/ltr_finder/; with -C -w 2 -s), RepeatScout[23] (<http://www.repeatmasker.org/>;with-sequence-freq -output), and RepeatModeler [24] (<http://www.repeatmasker.org/RepeatModeler.html>; with-database-engine ncbi-pa 15) were used to construct a de novo repeat sequence database of the *A. cannabinum* genome. The homologous repeat sequence database RepBase [25](<http://www.girinst.org/repbase/>) was then used for integration, and RepeatMasker [26] (<http://www.repeatmasker.org/>) was used for repeat sequence annotation of the *A. cannabinum* genome.

Annotation of non-coding RNA

The tRNAscan-SE [27] (<http://lowelab.ucsc.edu/tRNAscan-SE/>) software was used to identify tRNA sequences in the genome based on the structural characteristics of the tRNAs. As rRNAs are highly conserved, the rRNA sequences of phylogenetically related species were used as reference sequences, and blast alignment was used to identify rRNAs in the genome. Rfam covariance models were used along with the INFERNAL [28] (<http://infernal.janelia.org/>) in Rfam to predict miRNA and snRNA sequence information in the *A. cannabinum* genome.

Gene structure prediction

First, the repeat sequence mask obtained from the annotation by RepeatMasker (with -nolow -no_is -norna -parallel 1 -lib -s) was used for *de novo* prediction of the *A. cannabinum* genome using Augustus [29] (<http://bioinf.uni-greifswald.de/augustus/>; with-species = pasa1 -uniqueGenelid = true-nolnFrameStop = true-gff3 = on-genemodel = complete-strand = both) and GlimmerHMM[30] (<http://ccb.jhu.edu/software/glimmerhmm/>). Blast (<http://blast.ncbi.nlm.nih.gov/Blast.cgi>) was used to align the protein sequences of four phylogenetically related species, namely crown flower (*Calotropis gigantea*), Madagascar periwinkle (*Catharanthus roseus*), Arabian coffee (*Coffea arabica*), and flax (*Linum usitatissimum*), to the *A. cannabinum* genome to predict the structure sets of homologous genes.

Following this, blat [31] (<http://genome.ucsc.edu/cgi-bin/hgBlat>) was used for the alignment of the expressed sequence tag (EST) data to predict gene structures. Finally, EVIDENCEModeler [24] (<http://evidencemodeler.sourceforge.net/>) was used in combination with the transcriptome alignment data to combine the gene sets predicted by the different methods into a non-redundant and intact gene set.

Functional annotation of the genes

The gene sets obtained from the prediction of *A. cannabinum* gene structures were aligned using the Swiss-Prot [32] (<http://www.uniprot.org/>) and Non-redundant (Nr) [33] (<http://www.ncbi.nlm.nih.gov/protein>) databases using alignment software, and the best alignment result was used to determine gene function. The *A. cannabinum* genome was then aligned with Kyoto Encyclopedia of Genes and Genomes (KEGG) [34] metabolic pathways to determine the functions of the genes in the metabolic network and signaling pathways. In addition, InterPro [35] (<https://www.ebi.ac.uk/interpro/>) software was used for Gene Ontology (GO) [36] annotation and prediction of gene motifs and domains.

Comparative genome analysis

In this study, we selected the genomes of 15 plants for homologous gene analysis. These genomes included *C. roseus*, *C. gigantea*, *Medicago truncatula*, *Populus trichocarpa*, *L. usitatissimum*, *Gossypium raimondii*, *Gossypium hirsutumD*, *Gossypium hirsutumA*, *Gossypium barbadenseA*, *Morus notabilis*, *Arabidopsis thaliana*, *Glycyrrhiza uralensis*, *Gossypium barbadenseD*, *Phyllostachys heterocycla*, and *Corchorus capsularis*. The CDS sequences of the aforementioned plants were obtained from various databases (NCBI, Ensembl, Phytozome). Before analysis, redundancy was removed from all data and only the longest protein sequence was retained for each gene. OrthoMCL [37] was used for homologous gene analysis. MUSCLE [38] (<http://www.drive5.com/muscle/>) alignment was carried out on various families by using the sequences of all single-copy genes. RaxML [39] (<http://sco.h-its.org/exelixis/web/software/raxml/index.html>) was used for the construction of phylogenetic trees based on the sequence alignment results using the maximum likelihood method (ML TREE). Finally, CAFÉ [40] (<http://sourceforge.net/projects/cafehahnlab/>) was used for gene family expansion and contraction analysis.

Results And Statistics

Sequencing data statistics

The PacBio platform were used for whole genome sequencing the *A. cannabinum* genome. Table 1 represents total sequencing volume was 117.13 G, and the coverage was 490.04 X based on calculations of the genome size (239.02 M) estimated from a survey [41]. In addition, the Illumina platform was used for the sequencing of a constructed second-generation small fragment library.

Table 1
Apocynum cannabinum genome sequencing data statistics

Pair-end libraries	Insert size	Total data (G)	Read length (bp)	Sequence coverage (X)
Illumina reads	350 bp	31.94	150	133.63
PacBio reads	29.8 kb	8.76	11734	162.16
10X_Genomics	500–700 bp	46.43	150	194.25
Total	-	117.13	-	490.04

Assembly results and statistics

The genome size was 260 Mb, the contig N50 was 3.11 Mb, and the scaffold N50 was 4.19 Mb (scaffolds of 100 bp and above were selected for assembly results). The GC analysis showed that the GC content of *A. cannabinum* was 33.26%. Through Hi-C-assisted assembly, the final genome sequence was anchored to 11 chromosome. The draft genome scaffold chromosome anchoring results showed that the constructed super-scaffold N50 was 21.16 Mb and the number of super-scaffolds was six, which was significantly decreased. The number of scattered scaffolds was low, showing good assembly results. Among the scaffolds, N90 was 11, indicating that 90% of sequences were on the 11 chromosome. A comparison of N50, N60, N70, N80, and N90 showed that the Hi-C assembled *A. cannabinum* genome result was significantly improved from the original draft genome (Table 2).

Table 2
Apocynum cannabinum genome assembly results

Sample ID		Length		Number	
		Contig (bp)	Scaffold (bp)	Contig (bp)	Scaffold (bp)
10X_Genomics	Total	259,856,155	260,145,681	479	439
	Max	9,567,299	15,154,524	-	-
	Number > = 2000	-	-	479	439
	N50	3,107,135	4,189,930	26	19
	N60	2,442,000	3,120,689	35	27
	N70	1,722,976	2,442,000	48	36
	N80	1,116,808	1,547,882	66	49
	N90	257,161	339,954	109	80
Hi-C	Total	259,856,155	260,165,581	569	330
	Max	9,567,299	26,508,341	-	-
	Number > = 2000	-	-	567	330
	N50	2,442,000	21,162,547	33	6
	N60	1,938,819	20,272,879	45	7
	N70	1,256,579	19,617,059	61	9
	N80	786,196	19,407,958	87	10
	N90	207,482	17,325,238	152	11

A heatmap was used to validate the Hi-C assembly results (Fig. 1). From the heatmap, we can see that the interaction relationships near the diagonal were far stronger than those far away from the diagonal. The overall results showed that there was no significant clustering error between the *A. cannabinum* chromosomes.

Evaluation of assembly results

BUSCO and CEGMA software were used to evaluate the assembly results for *A. cannabinum*. The results showed that of the 1440 orthologous single-copy genes in the plant database, 1348 genes (C = 93.6%) were completely covered, of which 1258 single copies were covered (S = 87.4%), 89 multiple copies were covered (D = 6.2%), 30 genes were not completely covered (F = 2.1%), and 62 genes were missing (M = 4.3%). This suggests that the vast majority of single-copy genes was completely assembled, that there was no over-assembly or re-assembly, and that the assembly accuracy and integrity were good. At the same time, 227 genes were assembled from 248 core eukaryotic genes (CEGs), indicating that the

assembly results were relatively complete. Alignment using BWA software showed that the alignment rate of all small fragment reads to the genome was 94.80% and the coverage was 96.95%, demonstrating that the reads and the assembled genome possess good consistency.

Repeat sequence annotation

Repeat sequences can be classified as tandem repeats and interspersed repeats. Tandem repeats include microsatellite sequences and minisatellite sequences. Interspersed repeats are also known as transposon elements and include DNA transposons that use DNA-DNA transposition and retrotransposons. Commonly seen retrotransposons include long terminal repeats (LTRs), long interspersed elements (LINEs), and short interspersed elements (SINEs). Around 36.62% of sequences in the *A. cannabinum* genome were identified as repeat sequences, of which DNA transposons, LTRs, LINEs, and SINEs accounted for 2.22%, 28.45%, 1.01%, and 0%, respectively (Table 3).

Table 3
Statistics of repeat sequence classification results for *A. cannabinum*

	Denovo + Repbase length (bp)	% in genome	TE protein length (bp)	% in genome	Combined TE length (bp)	Type
DNA	5,282,438	2.03	765,646	0.29	5,763,303	2.22
LINE	1,744,782	0.67	1,157,517	0.44	2,630,319	1.01
SINE	1,597	0.00	0	0	1,597	0.00
LTR	72,896,704	28.02	16,110,716	6.19	74,004,538	28.45
Unknown	14,826,765	5.70	0	0	14,826,765	5.70
Total	92,189,164	35.44	18,033,358	6.93	93,170,210	35.81
Note: Denovo + Repbase are transposon elements obtained after the databases predicted by RepeatModeler, RepeatScout, Piler, and LTR_FINDER were combined with the RepBase nucleic acid database. This was followed by integration using Uclust software according to the 80-80-80 principle, following which RepeatMasker was used for genome annotation. The TE proteins are transposon elements that were obtained when the RepeatProteinMask software was used to annotate genomes in the RepBase protein database. Combined TEs is the result obtained by combining the results obtained from these two methods and following the removal of redundant genes. Unknown indicates that this repeat sequence could not be classified by RepeatMasker.						

Table 3
Statistics of repeat sequence classification results for *Apocynum cannabinum*

	Denovo + Repbase length (bp)	% in genome	TE protein length (bp)	% in genome	Combined TE length (bp)	Type
DNA	5,282,438	2.03	765,646	0.29	5,763,303	2.22
LINE	1,744,782	0.67	1,157,517	0.44	2,630,319	1.01
SINE	1,597	0.00	0	0	1,597	0.00
LTR	72,896,704	28.02	16,110,716	6.19	74,004,538	28.45
Unknown	14,826,765	5.70	0	0	14,826,765	5.70
Total	92,189,164	35.44	18,033,358	6.93	93,170,210	35.81
Note: Denovo + Repbase are transposon elements obtained after the databases predicted by RepeatModeler, RepeatScout, Piler, and LTR_FINDER were combined with the RepBase nucleic acid database. This was followed by integration using Uclust software according to the 80-80-80 principle, following which RepeatMasker was used for genome annotation. The TE proteins are transposon elements that were obtained when the RepeatProteinMask software was used to annotate genomes in the RepBase protein database. Combined TEs is the result obtained by combining the results obtained from these two methods and following the removal of redundant genes. Unknown indicates that this repeat sequence could not be classified by RepeatMasker.						

Gene structure annotation

Through the annotation of the *A. cannabinum* genome, a total of 22,793 non-redundant protein coding genes were predicted. The mean gene sequence length of the gene was 3532 bp, the mean length of coding sequence was 1206 bp, and the exons and introns were 231 bp and 551 bp, respectively. The mean number of exons in every gene was 5.22. Comparison of the lengths of genes, coding sequences, introns, exons, and the number of exons in *A. cannabinum* with *C. gigantea*, *Catharanthus roseus*, *C. arabica*, and *L. usitatissimum* showed that the lengths of the coding sequences and exons and the number of exons in every gene in *A. cannabinum* were closest to *C. gigantea* and *L. usitatissimum*, while *C. roseus* had lower numbers than these three species and *C. arabica* had higher numbers than these species. The lengths of the genes and introns of *A. cannabinum* were most similar to *C. arabica*, while the lengths of the genes and introns of *L. usitatissimum* were significantly lower than the other four plant species. The lengths of almost all introns in the five plant species were within 1000 bp, and introns with lengths greater than 2000 bp were extremely rare or absent (Fig. 2).

Annotation of gene function and non-coding RNAs

A total of 78.70% of genes could be aligned to known metabolic pathways; 12,823 functional genes could be annotated by GO function; and 4.4% of genes had unknown functions. Among the annotated protein-coding genes, 21,780 genes had homologous sequences in the protein database, accounting for 95.6% of

annotated genes. Furthermore, 95.2% of genes were aligned to non-redundant protein databases, of which 92.30% contained conserved protein domains, showing that the functions of most *A. cannabinum* genes were relatively conserved. In addition, 689 tRNAs, 3934 rRNAs, 592 snRNAs, and 159 miRNAs were annotated in the *A. cannabinum* genome (Fig. 3).

Comparative genomic analysis of *A. cannabinum*

Clustering analysis of the gene families indicated that 43,837 gene families were clustered in 17 species, of which 3084 gene families were common to all these species, while there were 55 common single-copy gene families in the different species. The clustering results of *A. cannabinum*, *A. venetum*, *C. roseus*, *C. gigantea*, and *C. capsularis* were used to plot a Venn diagram (Fig. 4).

From the Venn diagram, we can see that *A. cannabinum* contains 283 species-specific gene families and 449 genes compared with other species. The gene family expansion and contraction analysis results showed that the most recent common ancestor (MRCA) contained 43,823 gene families, of which 399 gene families and 166 genes were expanded in *A. cannabinum*, while 492 gene families and 231 genes were contracted (Fig. 5).

The gene family clustering analysis results showed that *C. roseus*, *A. venetum*, *A. cannabinum*, and *C. gigantea* clustered in the same taxon. The Markov chain Monte Carlo (MCMC) program in PAML software indicated that the speciation time for *A. cannabinum* was around 35.9 (26.7–46.8) million years ago.

Discussion

In this study, we employed Illumina HiSeq, single-molecule real-time sequencing, and 10X genomics as well as chromatin interaction (Hi-C), a new assembly technique, to successfully assemble the whole draft genome for *A. cannabinum* for the first time. The whole-genome sequence of *A. cannabinum* was 260 Mb, the size of the genome super-scaffold N50 was 21.16 Mb, and the scaffold N50 was 4.19 Mb. The combination of the Hi-C sequencing results and high-density genetic maps resulted in the scaffolds being anchored to 11 chromosome. The super-scaffold N50 length was 21.16 Mb, which is currently the best result among all assembled genomes from the Apocynaceae family. There are four species genomes of Apocynaceae family (*Catharanthus roseus*, *Calotropis gigantea*, *Rhazya stricta* and *Asclepias syriaca*) have been sequenced and published. However, since the contig N50 of *C. roseus* and *A. syriaca* is all less than 10 kb, it is easy to cause fairly big errors in genome annotation and gene localization. In addition, the published genomes from the Apocynaceae family are at the chromosome level, while the present assembled *A. cannabinum* genome exceeds the chromosome level. Gene annotation refers to the specific role of biomolecules in the process of life. Genome annotation is an essential link connecting the sequence information of genes with specific biological processes. Therefore, the functional annotation process of all genes in the whole genome sequence of plants is also considered as the "metabolic reconstruction" of plants. The purpose of metabolic reconstruction is to identify metabolic pathways and which genes have this function in plants. However, the plant genome sequence contains a large number of repetitive regions, pseudogenes, as well as many new protein-coding and non-coding genes. The

repetitive sequence can even be as high as 80% in plant genome. Although the gene duplication plays an important role in maintaining the regulation of gene expression, the spatial structure of chromosomes and genetic recombination, etc., it brings out many false positives in BLAST results, which increases the pressure of gene structure prediction and may lead to the high error rate of gene annotation [44]. This provides a foundation for future studies on *A. cannabinum*. Gene annotation in *A. cannabinum* revealed a total of 22,793 non-redundant protein-coding genes, of which 21,780 have homologous sequences in the protein database, accounting for 95.6% of annotated genes. This prediction result is similar to the number of genes (21,164) in *Rhazya stricta* from the Apocynaceae family [43], but is 12,049 lower than in *C. roseus* (33,829) [44]. Conversely, the number of genes in *Asclepias syriaca* (14,474) [45] and *C. gigantea* (18,197) [46] is lower than in *A. cannabinum*. A previous study also showed that no direct relationship exists between genome size and the number of genes in plants [47]. The reasons for the large differences in the number of genes include: the selection of overly high threshold values for annotation, resulting in marginal data being excluded; the selection of only a few major public databases, resulting in incomplete coverage; a lack of homology annotation information for certain specific genes, which require functional validation in the future; and excessive structural annotation, resulting in the annotation of false positive genes [47]. According to our study, *A. cannabinum* has better drought tolerance than *A. venetum* [1, 4], we believe that it will be well explained when the genome is analyzed in the future.

Comparative genomic analysis of *A. cannabinum* indicated that a total of 43,837 gene families were clustered, of which 3084 were common gene families, 399 gene families and 1666 genes were expanded, and 492 gene families and 231 genes were contracted. This may be due to the varying degrees of acquisition and loss of genes (families) of each species during evolution [48]. Additionally, *A. cannabinum* speciated around 35.9 (26.7–46.8) million years ago. Gene family clustering found that *C. roseus*, *A. venetum*, *A. cannabinum*, and *C. gigantea* were clustered in the same taxon, which is highly consistent with the angiosperm classification system IV (APG IV).

Conclusions

Whole-genome sequencing of *A. cannabinum* was carried out, and 117.13 G of raw data was obtained. The sequencing depth was 490.04 X, the assembled genome size was 260 Mb, and the contig N50 was 3.11 Mb. Multiple methods were used for the evaluation of the assembled versions. The results showed that *A. cannabinum*'s genome version had high consistency, integrity, and accuracy. A total of 22,793 protein-coding genes were predicted by genome annotation, and the ratio of repeat sequences was 35.81%, of which functions were predicted for 21,780 genes (95.6%). Clustering analysis of the 17 species was used to obtain 43,837 gene families, of which 3084 constituted common gene families. Comparative analysis of the gene families showed that *A. cannabinum* speciated around 35.8 (27.0–46.9) million years ago.

Declarations

Acknowledgments This work is supported by Ningxia Key Research and Development Project (grant number: 2022BEG03006), and Ningxia Natural Science Foundation (grant number: 2020AAC03076). Especially, we heartily thank professor Dongmei Ma, Mr. Boxun Xie and Sheng Xie, they all gave us some useful revision advices.

Authors Contributions Guoqi Li conceived and designed the experiment, Lixiao Song and Jinfeng Che conducted experiment and data analysis, Guoqi Li, Lixiao Song, Jinfeng Che and Yanyun Chen wrote, discussed and rewrote the manuscript together.

Conflict of interest The authors declared that they have no competing interest.

Ethical approval This article does not contain any studies with human participants or animals performed by any of the authors.

Data availability All research data are available upon request.

References

1. Li G Q, Chen Y Y. (2012) *Physioecology of Apocynum*. Science Press, 2012 (in Chinese)
2. Xie W Y, Zhang X Y, Wang T, Hu J J. (2012). Botany, traditional uses, phytochemistry and pharmacology of *Apocynum venetum* L. (Luobuma): A review. *Journal of Ethnopharmacology*, 2012, 141 (1): 1-8. <https://doi.org/10.1016/j.jep.2012.02.003>
3. Li G Q, Zhao P P, Shao W S. (2019) Cash crop halophytes of China. In: Gul B, Böer B, Khan M A, Clüsener-Godt, Hameed A. (eds) *Sabkha Ecosystems*(Volume VI): Asia/Pacific. Springer, Nature Switzerland, 2019: 497-504. <https://doi.org/10.1007/978-3-030-04417-6>
4. Wang D Q, Li G Q, Wang L. (2012) Daily dynamics of photosynthesis and water physiological characteristics of *Apocynum venetum* and *A. cannabinum* under drought stress. *Acta Botanica Boreali-Occidentalia Sinica*, 32 (6), 1198-1205. <https://doi.org/10.3969/j.issn.1000-4025.2012.06.020> (in Chinese)
5. DiTommaso A, Clements D R, Darbyshire S J, Dauer J T. (2009) The Biology of Canadian Weeds, 143, *Apocynum cannabinum* L. *Can J Plant Sci* 89(5), 977-992. <https://doi.org/10.1639/0007-2745-112.3.614>
6. Keever C. (1979) Mechanisms of plant succession on old fields of Lancaster County, Pennsylvania. *Bulletin of the Torrey Botanical Club*, 106 (4), 299-308. <https://doi.org/10.2307/2560356>
7. Mulhouse J M, Galatowitsch S M. (2003) Revegetation of prairie pothole wetlands in the mid-Continental US: twelve years post-reflooding. *Plant Ecol* 169 (1), 143-159. <https://doi.org/10.2307/20146504>
8. Duprey A. (1905) A case of mitral incompetency and ascites treated with *Apocynum cannabinum*. *Lancet* 166 (4283), 955-956. [https://doi.org/10.1016/S0140-6736\(01\)12615-2](https://doi.org/10.1016/S0140-6736(01)12615-2)

9. Leidy. (1884) Indian use of *Apocynum cannabinum* as a textile fibre. Proceedings of the Academy of Natural Sciences of Philadelphia, 36, 30-30. <https://doi.org/10.2307/4060952>
10. Li B T, Chen X M. (1997) Comparative review on Apocynaceae in flora reipublicae popularis sinicae and flora of China. Guihaia, 17 (4):299-305. (in Chinese)
11. Li P T, Leeuwenberg A J M, Middleton D J. (1995) Flora of China. Beijing: Science Press and St. Louis: Missouri Botanical Garden, 143-188
12. Goff S A, Ricke D, Lan T H, Presting G, Wang R L, Dunn M. (2002) A draft sequence of the rice genome (*Oryza sativa* L. ssp. japonica). Science 296 (5565), 92-100. <https://doi.org/10.1126/science.1068275>
13. Burton J N, Adey A, Patwardhan R P, Qiu R, Kitzman J O, Shendure J. (2013) Chromosome-scale scaffolding of de novo genome assemblies based on chromatin interactions. Nat Biotechnol 31(12), 1119-1125. <https://doi.org/10.1038/nbt.2727>
14. Eisenstein M. (2015) Startups use short-read data to expand long-read sequencing market. Nat Biotechnol 33(5), 433-435. <https://doi.org/10.1038/nbt0515-433>
15. Mascher M, Gundlach H, Himmelbach A, Beier S, Twardziok S O, Wicker T, et al. (2017) A chromosome conformation capture ordered sequence of the barley genome. Nature, 544(7651), 427-433. <https://doi.org/10.1038/nature22043>
16. Xu M, Sun Y G, Li H G. (2010) EST-SSRs development and paternity analysis for *Liriodendron* spp. New Forest, 40(3), 361-382. <https://doi.org/10.1007/s11056-010-9205-0>
17. Liu H L. (2018) The whole genome sequencing and analyzing of *Ginkgo biloba*. Nanjing Forestry University (in Chinese)
18. Hansen K D, Brenner S E, Dudoit S. (2010) Biases in illumina transcriptome sequencing caused by random hexamer priming. Nucleic Acids Res 38, 131-138. <https://doi.org/10.1093/nar/gkq224>
19. Jarvis D E, Ho Y S, Lightfoot D J, Schmöckel S M, Li B, Borm T J, Hajime Ohyanagi H, Mineta K, Michell C T, Saber N, Kharbatia N, Rupper R R, Sharp A R, Dally N, Boughton B, Woo Y, Gao G, Schijlen E, Guo X, Momin A A, Negrao S, Al-Babili S, Gehring C, Roessner U, Jung C, Murphy K, Arold S, Gojobori T, Linden C, Loo EV, Jellen R, Maughan J, Tester M. (2017) The genome of *Chenopodium quinoa*. Nature, 542 (7641), 307-312. <https://doi.org/10.1038/s41477-018-0166-1>
20. Hu W Q, Hou Y, Zhang F, Liu H D, Sun X. (2015) A Chromatin conformation analysis technology—Hi-C and extracting of chromatin conformation information. Genomics and Applied Biology, 34(11), 36-44. <https://doi.org/10.13417/j.gab.034.002319>
21. Li H. (2013) Aligning sequence reads, clone sequences and assembly contigs with BWA-MEM. arXiv, 1303.3997
22. Zhao, X. and Hao, W. (2007) LTR_FINDER: an efficient tool for the prediction of full-length LTR retrotransposons. *Nucleic Acids Res.* 35(Web Server issue), W265-8. <https://doi.org/10.1093/nar/gkm286>
23. Price A L, Jones N C, Pevzner P A. (2005) De novo identification of repeat families in large genomes. Bioinformatics, 21(suppl_1), 351-358. <https://doi.org/10.1093/bioinformatics/bti1018>

24. Haas B J, Salzberg S L, Zhu W, Pertea M, Allen JE, Orvis J, White O, Buell C R, Wortman JR. (2008) Automated eukaryotic gene structure annotation using EVIDENCEModeler and the Program to assemble spliced alignments. *Genome Biol* 9(1), R7. <https://doi.org/10.1186/gb-2008-9-1-r7>
25. Jurka J. (2000) Repbase update: a database and an electronic journal of repetitive elements. *Trends Genet* 16(9), 418-420. [https://doi.org/10.1016/S0168-9525\(00\)02093-X](https://doi.org/10.1016/S0168-9525(00)02093-X)
26. Chen N. (2004) Using RepeatMasker to identify repetitive elements in genomic sequences. *Current Protocols in Bioinformatics* 5(1). <https://doi.org/10.1002/0471250953.bi0410s05>
27. Kanehisa M Goto S. (2000) KEGG: Kyoto Encyclopedia of Genes and Genomes. *Nucleic Acids Res* 28(1), 27-30. <https://doi.org/10.1093/nar/28.1.27>
28. Nawrocki E P, Kolbe D L, Eddy S R. (2009) Infernal 1.0: inference of RNA alignments. *Bioinformatics*, 25(10), 1335–1337. <https://doi.org/10.1093/bioinformatics/btp157>
29. Stanke M, Steinkamp R, Waack S, Morgenstern B. (2004) AUGUSTUS: a web server for gene finding in eukaryotes. *Nucleic Acids Res* 32(Web Server issue), 309-12. <https://doi.org/10.1093/nar/gkh379>
30. Majoros W H, Pertea M, Salzberg S L. (2004) TigrScan and GlimmerHMM: Two open source ab initio eukaryotic gene-finders. *Bioinformatics*, 20(16), 2878-2879. <https://doi.org/10.1093/bioinformatics/bth315>
31. Kent W J. (2002) BLAT-the BLAST-like alignment tool. *Genome Res* 12(4), 656-664. <https://doi.org/10.1101/gr.229202>
32. Bairoch A. (2000) The SWISS-PROT protein sequence database and its supplement TrEMBL in 2000. *Nucleic Acids Res* 28(1), 45–48. <https://doi.org/10.1093/nar/28.1.45>
33. Griffiths-Jones S, Moxon S, Marshall M, Khanna A, Eddy S R, Bateman A. (2005) Rfam: annotating non-coding RNAs in complete genomes. *Nucleic Acids Res* 33, D121-124. <https://doi.org/10.1093/nar/gki081>
34. Ogata H, Goto S, Sato K, Fujibuchi W, Bono H, Kanehisa M. (1999) KEGG: kyoto encyclopedia of genes and genomes. *Nucleic Acids Res* 27(1), 29-34. <https://doi.org/10.1093/nar/27.1.29>
35. Zdobnov, E.M., & Apweiler, R. (2001). InterProScan-an integration platform for the signature-recognition methods in InterPro. *Bioinformatics*, 17(9): 847-848. <https://doi.org/10.1093/bioinformatics/17.9.847>
36. Ashburner M, Ball C A, Blake J A, Botstein D, Butler H, Cherry, J.M. (2000) Gene ontology: tool for the unification of biology. *Nat Genet* 25(1), 25-29. <https://doi.org/10.1038/75556>
37. Li L, Stoeckert C J, Roos D S. (2003) OrthoMCL: identification of ortholog groups for eukaryotic genomes. *Genome Res* 13(9), 2178-2189. <https://doi.org/10.1101/gr.1224503>
38. Robert C E. (2004) MUSCLE: multiple sequence alignment with high accuracy and high throughput. *Nucleic Acids Res* 32, 1792-1797. <https://doi.org/10.1093/nar/gkh340>
39. Guindon S, Gascuel O. (2003) A simple, fast, and accurate algorithm to estimate large phylogenies by maximum likelihood. *Syst Biol* 52(5), 696-704. <https://doi.org/10.1080/10635150390235520>

40. Yang Z, Rannala B. (2012) Molecular phylogenetics: principles and practice. *Nat. Rev. Genet* 13(5), 303-314. [https://doi.org/ 10.1038/nrg3186](https://doi.org/10.1038/nrg3186)
41. Song L X, Li G Q, Jin C Q, Gong S P. (2019) Whole genome sequencing and development of SSR markers in *Apocynum cannabinum*. *Journal of Plant Genetic Resources*, 20(5), 1309-1316. [https://doi.org/ 10.13430/j.cnki.jpgr.20181218002](https://doi.org/10.13430/j.cnki.jpgr.20181218002) (in Chinese)
42. Huang J, Xu Q W, Lu Q. (2007) An RDF model of gene ontology and its associations. *Chinese Science Bulletin* (22): 40-46. [https://doi.org/ 10.3321/j.issn:0023-074x.2007.22.007](https://doi.org/10.3321/j.issn:0023-074x.2007.22.007)(in Chinese)
43. Sabir J S M, Jansen R K, Arasappan D, Calderon V, Noutahi E, Zheng C F, et al. (2016) The nuclear genome of *Rhazya stricta* and the evolution of alkaloid diversity in a medically relevant clade of Apocynaceae. *Sci. Rep-UK* 6(1), 33782. <https://doi.org/10.1038/srep33782>
44. Kellner F, Kim J, Clavijo B J, Hamilton J P, Childs K L, Vaillancourt B., Jason Cepela J, Habermann M, Steuernagel B, Catchpole L, Mclay K, Buell C R, o'connor S. (2015) Genome-guided investigation of plant natural product biosynthesis. *The Plant J.* 82(4), 680-692. <https://doi.org/10.1111/tpj.12827>
45. Weitemier K, Straub S, Fishbein M, Bailey C D, Cronn R C, Liston A. (2019) A draft genome and transcriptome of common milkweed (*Asclepias syriaca*) as resources for evolutionary, ecological, and molecular studies in milkweeds and Apocynaceae. *Peer J*, 7:e7649. <https://doi.org/10.7717/peerj.7649>
46. Hoopes G M, Hamilton J P, Kim J, Zhao D, Wiegert-Rininger K, Crisovan E, Buell C R. (2017) Genome assembly and annotation of the medicinal plant *Calotropis gigantea*, a producer of anticancer and Anti-malarial Cardenolides. *G3-Genes Genom Genet* 8(2), 385-391. <https://doi.org/10.1534/g3.117.300331>
47. Wang S. (2019) Whole genome sequencing and analysis of *Betula platyphylla*. Northeast Forestry University (in Chinese)
48. Lang K, Bi S D, Li F. (2018) Genome-wide analysis of expansion and contraction of gene families in parasitic wasps. *Journal of Anhui Agricultural University*, 45(5): 945-950. <https://doi.org/10.13610/j.cnki.1672-352x.20181023.025>

Figures

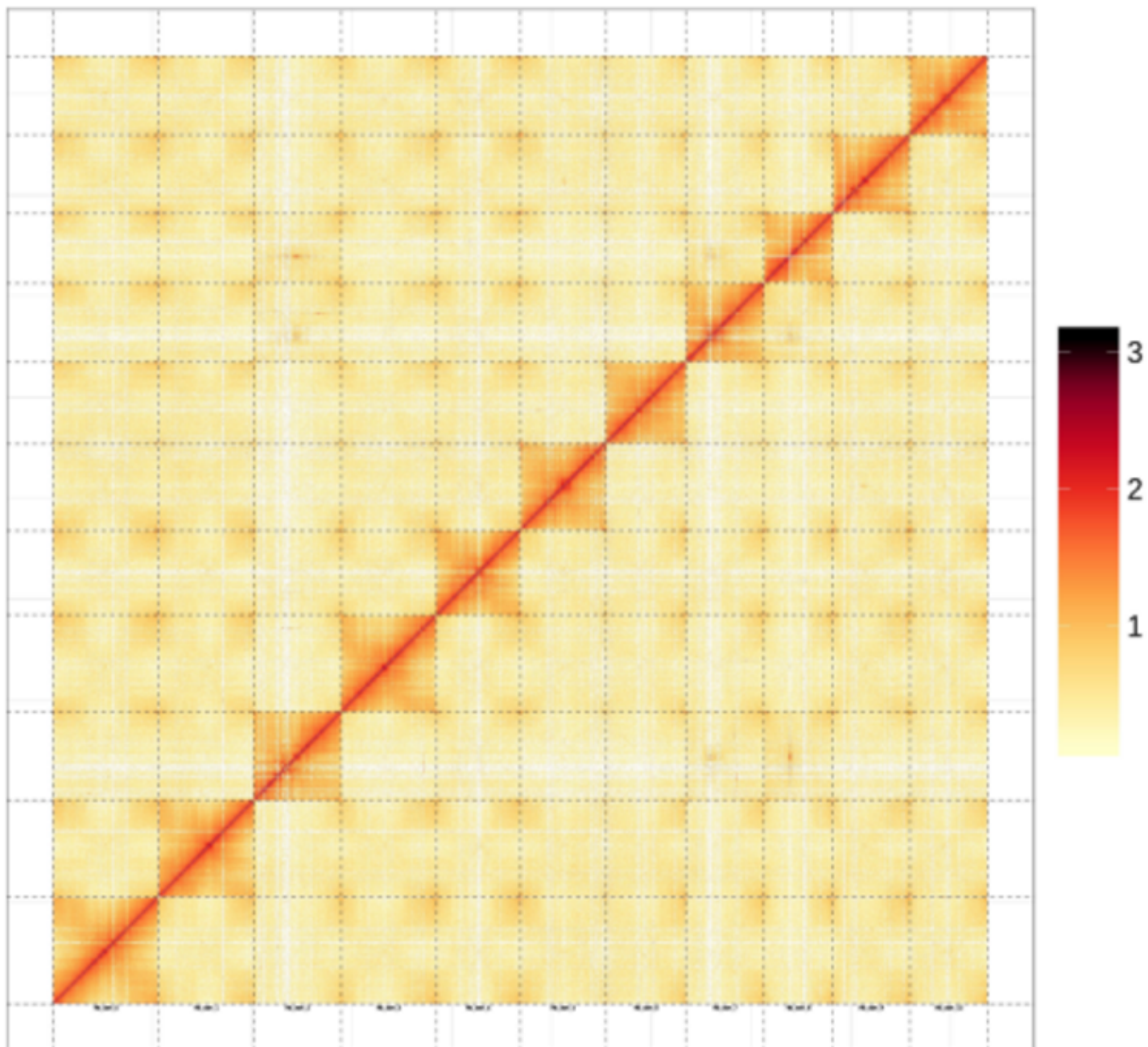


Figure 1

A. cannabinum chromosome interaction heatmap

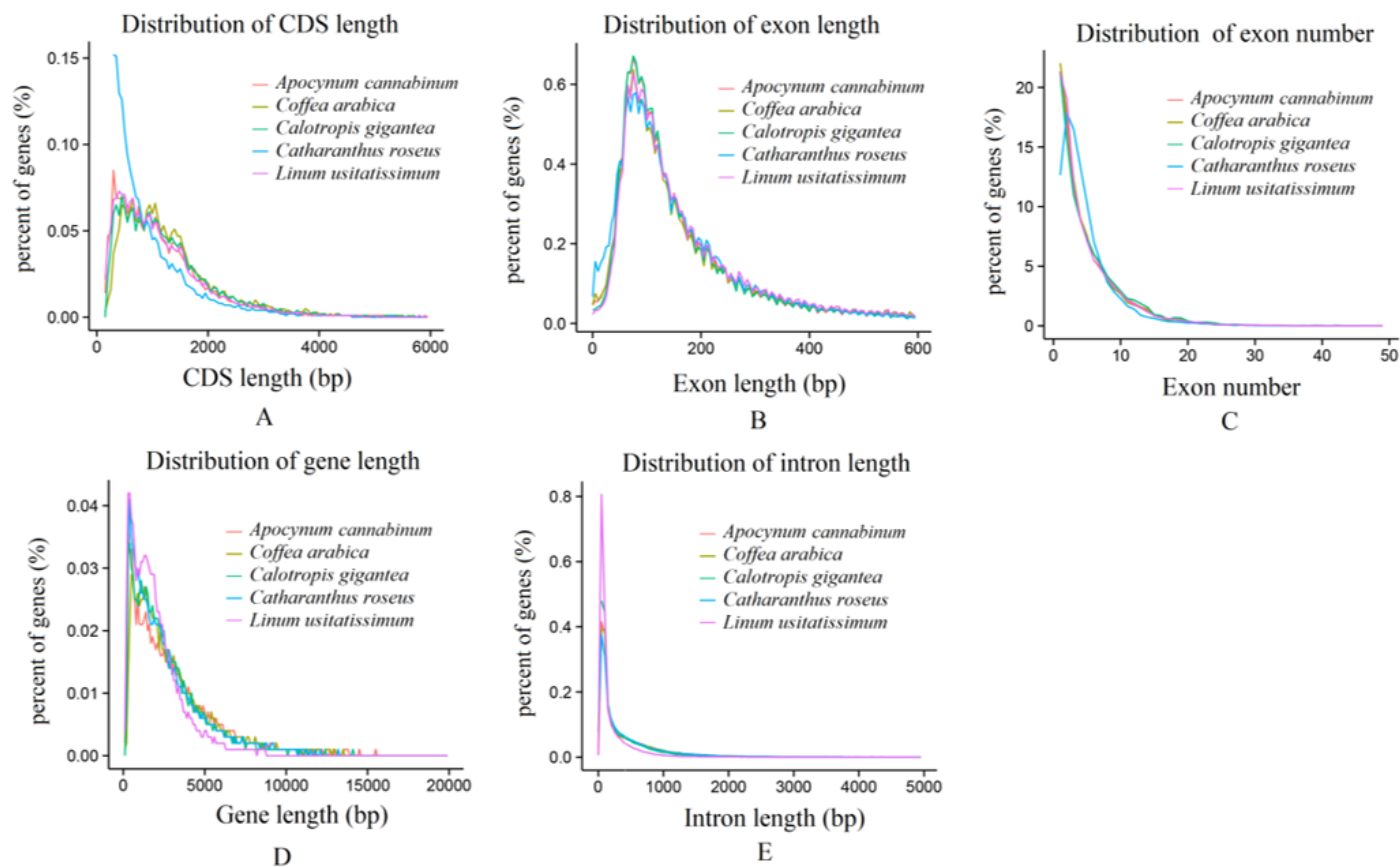


Figure 2

Comparison of various elements in phylogenetically related species

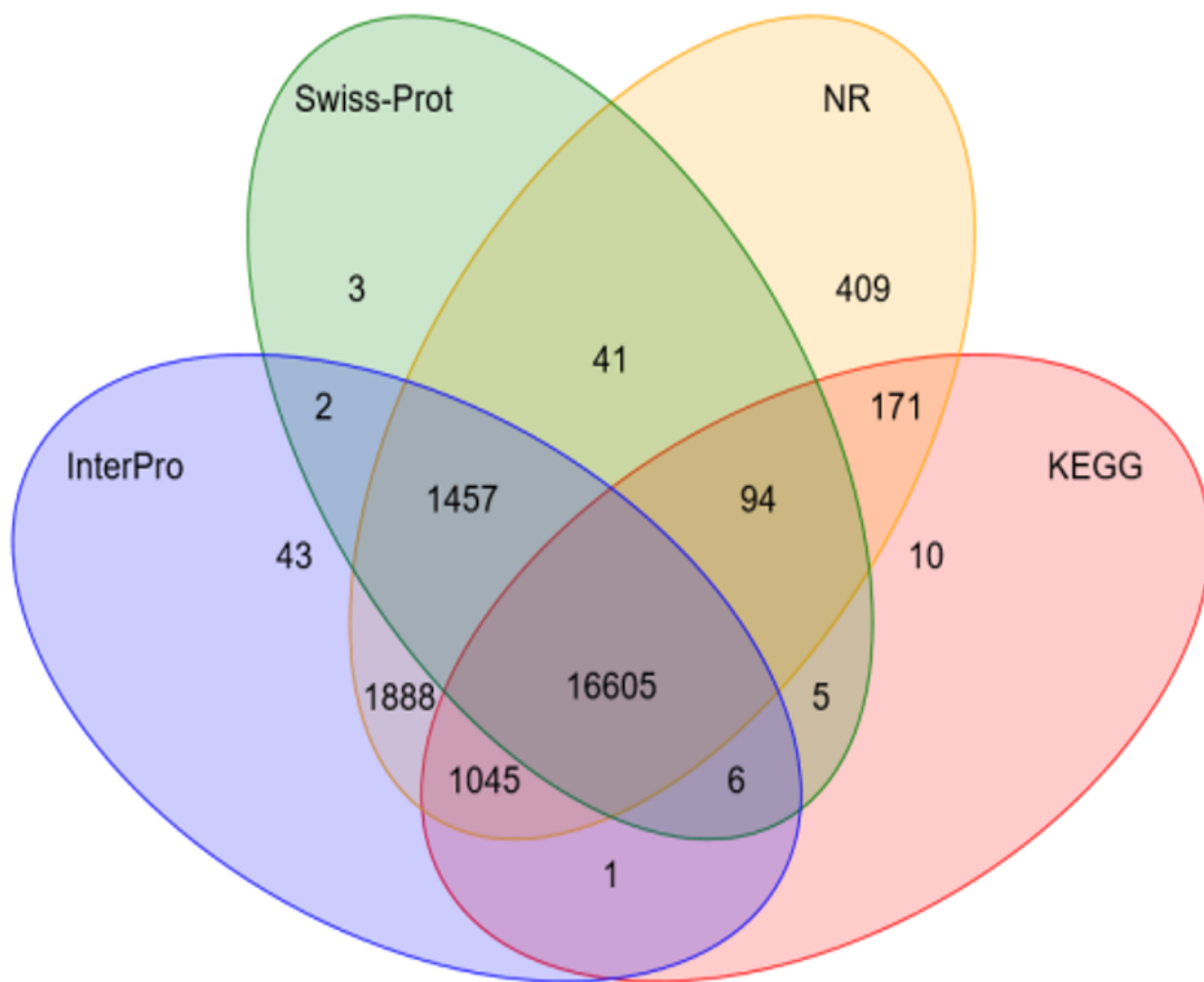


Figure 3

Gene function annotation results for *A. cannabinum*

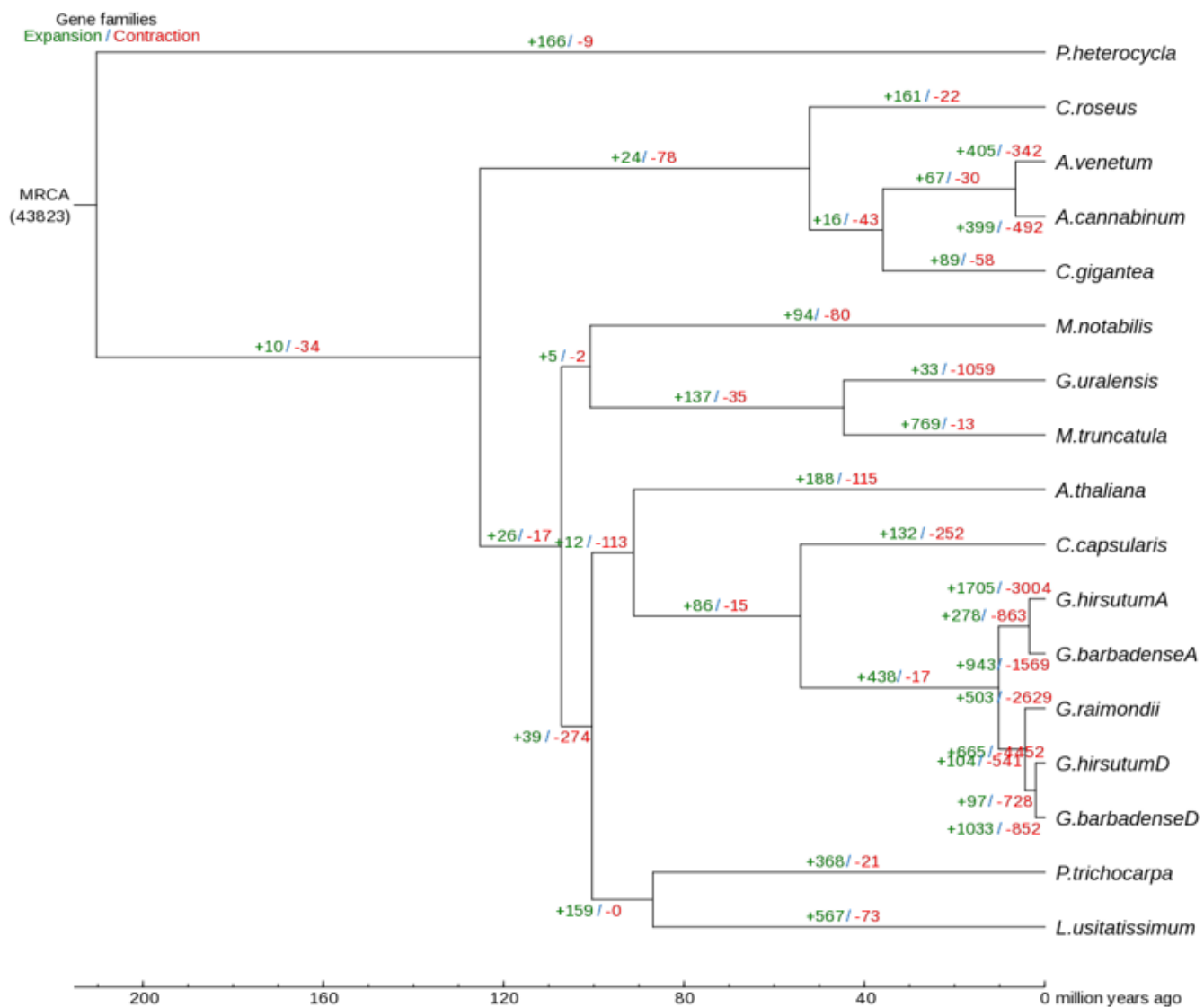


Figure 5

Annotation: as shown in the above figure, green numbers represent the number of gene families that expanded during evolution, and the red numbers represent the number of gene families that contracted during evolution.