

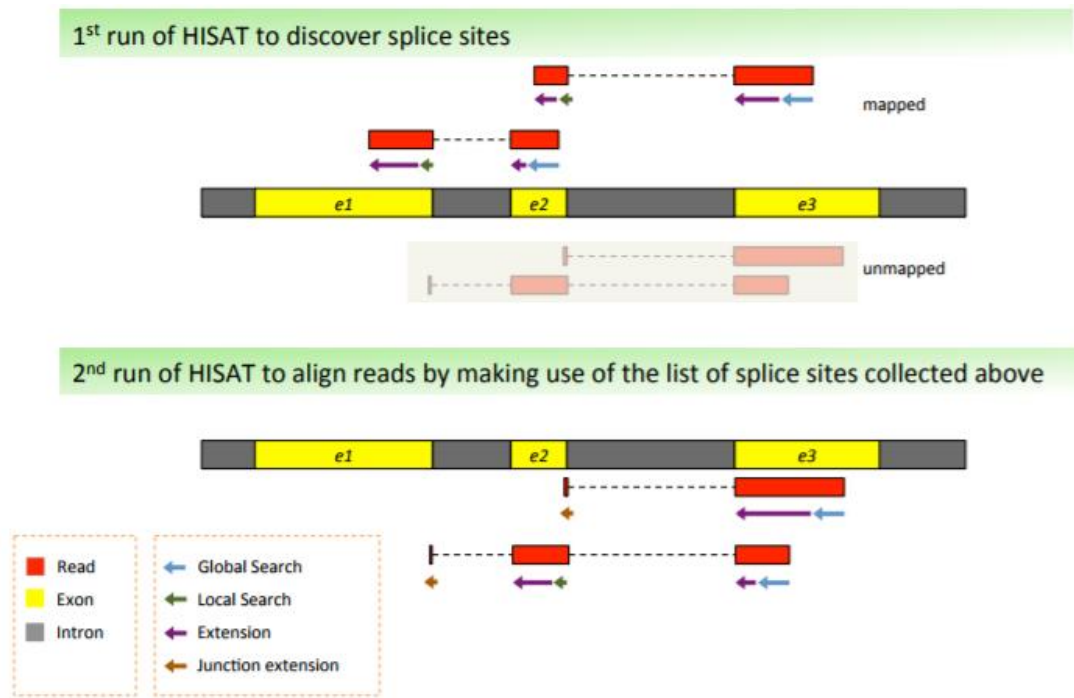


Figure S1 HISAT2 alignment principle

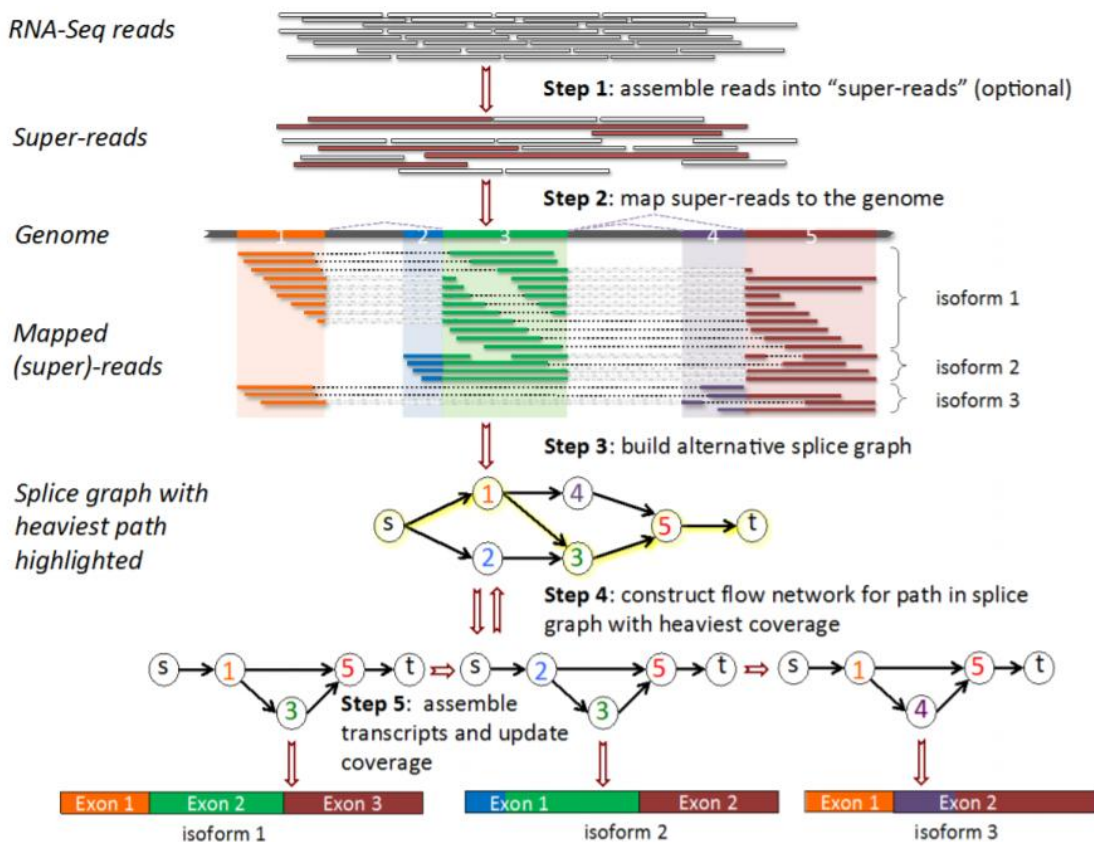


Figure S2 The principle of transcript assembly for Stringtie

Table S1 Gene differential expression analysis method

Type	Software	Standardized methods	PValue calculation model	FDR calculation method	Differential gene screening criteria
Biological duplication	DESeq2(Anders et al, 2014)	DESeq	Negative binomial distribution	BH	$ \log_2(FC) > 1, FDR < 0.05$
No biological duplication	edgeR(Robinson et al, 2010)	TMM	Negative binomial distribution	BH	$ \log_2(FC) > 1, FDR < 0.05$

Alternative Splicing Events

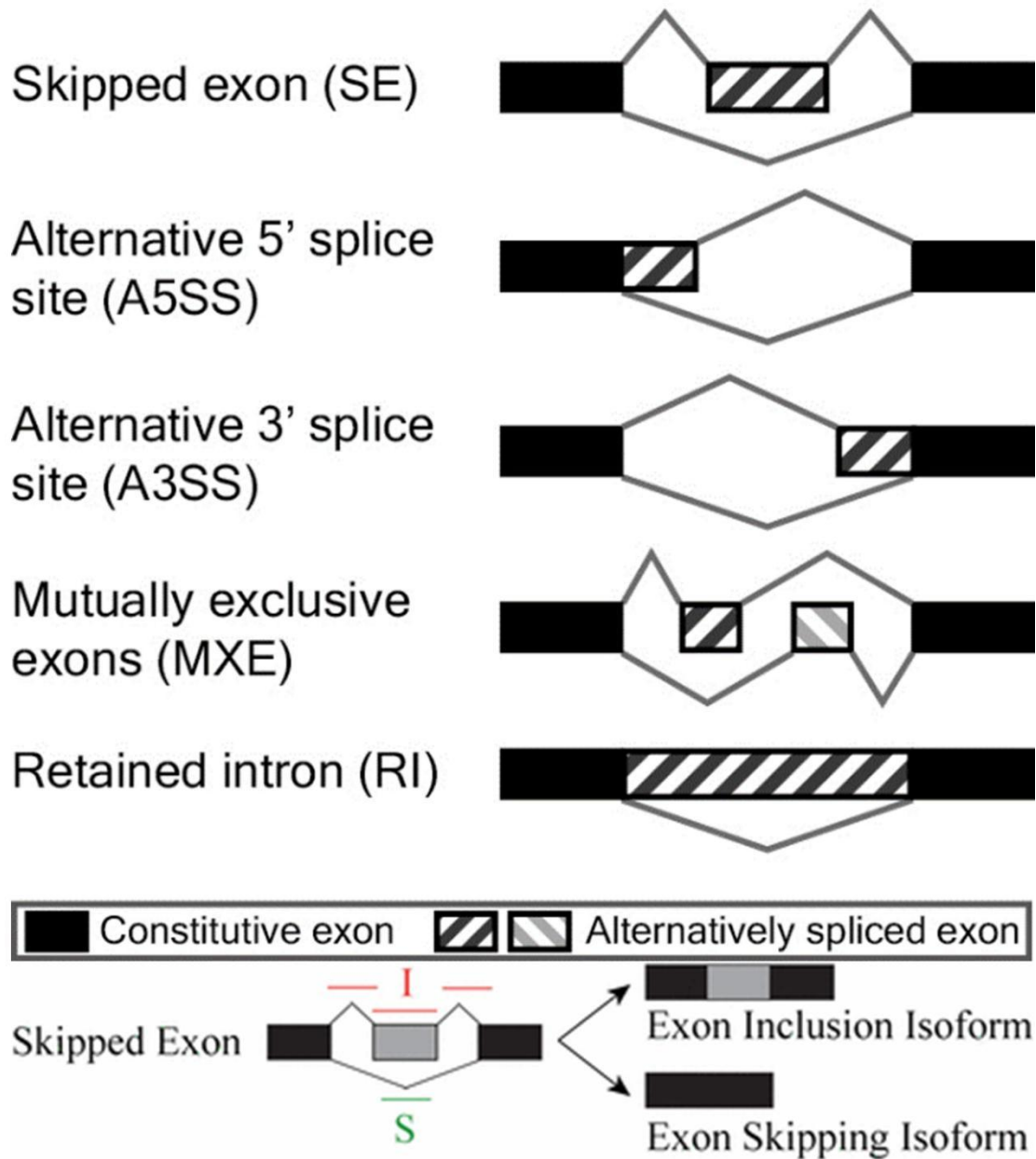


Figure S3 Schematic diagram of alternative splicing events

Code 1 Cross column duplicate filtering

```
import pandas as pd
```

```
from collections import defaultdict
```

```
# Read Excel file (assuming file name is input.xlsx)
```

```
df = pd.read_excel("input.xlsx", header=0)
```

```
# Create a mapping dictionary from values to column labels (automatic deduplication)
```

```
col_map = defaultdict(set)
```

```

for col in df.columns:
    for val in df[col].dropna().astype(str).unique():
        col_map[val].add(col)

# Create a mapping dictionary from column sets to values (automatic grouping)
group_map = defaultdict(set)
for val, cols in col_map.items():
    if len(cols) >= 2:
        key = ','.join(sorted(cols)) # Sort column labels alphabetically
        group_map[key].add(val)

# Generate the final result
result = [[cols, ','.join(sorted(values))]
           for cols, values in group_map.items()]

# Export to a new Excel file
pd.DataFrame(result, columns=["Comparison group", "Common
items"]).to_excel("output.xlsx", index=False)

```

Code 2 Q value search on demand

```

import pandas as pd

# Read Excel file
df = pd.read_excel("input.xlsx")

# Create a dictionary (with keys in the first column and values in the second column)
key_value = df.set_index(df.columns[0])[df.columns[1]].to_dict()

# Process the third column query and generate results
result = df[df.columns[2]].apply(
    lambda x: ','.join([str(key_value.get(k.strip(), "")) for k in str(x).split(",")])
)

# Save to a new Excel file
result.to_excel("output.xlsx", index=False, header=["Qvalue"])

```

Code 3 Merge with 3 decimal places retained

```

import pandas as pd

def format_row(go_str, val_str):
    # Force conversion to string processing
    go_ids = str(go_str).split(',')
    values = [float(v) for v in str(val_str).split(',')]

```

```

# Process Scientific notation and format (use English brackets)
formatted = []
for g, v in zip(go_ids, values):
    # Change the parentheses here to English format
    formatted.append(f'{g.strip()}({v:.3f})')

return ','.join(formatted)

# Force all columns to be of string type when reading
df = pd.read_excel('input.xlsx', dtype=str)

# Process each row of data
df['combined'] = df.apply(
    lambda row: format_row(row.iloc[0], row.iloc[1]),
    axis=1
)

# Save the Results
df[['combined']].to_excel('output.xlsx', index=False)

```

Code 4 Delete duplicate values

```

import pandas as pd

def merge_columns(input_path, output_path):
    # Read data and retain header rows
    df = pd.read_excel(input_path, header=0)
    headers = df.columns.tolist() # Save original title

    # Processing data section (excluding header rows)
    df_data = df.map(lambda x: str(x).strip().replace(" ", "")) if pd.notna(x) else ""

    # Cell level deduplication (retaining first occurrence)
    seen = set()
    # Traverse the indexes (rows, column names) of all non empty cells
    for (row_idx, col_name) in df_data.stack().index:
        # Convert column names to integer indexes
        col_idx = df_data.columns.get_loc(col_name)
        val = df_data.iat[row_idx, col_idx]
        if val:
            if val in seen:
                df_data.iat[row_idx, col_idx] = ""
            else:
                seen.add(val)

```

```

# Extract valid data from each column and merge them
merged_data = []
# Traverse column index (integer)
for col_idx in range(df_data.shape[1]):
    valid_values = [
        df_data.iat[row_idx, col_idx]
        for row_idx in df_data.index
        if df_data.iat[row_idx, col_idx]
    ]
    merged_data.append(", ".join(valid_values))

# Build result (title row+merged data row)
result_df = pd.DataFrame([headers, merged_data])

# Save to Excel
result_df.to_excel(output_path, index=False, header=False)

# Execute processing
merge_columns("input.xlsx", "output.xlsx")

```

Table S2 Data filtering statistics table

Sample	RawData as	CleanData(%)	Adapter(%)	LowQuality(%)	polyA(%)	N(%)
H167-1	4173115	41597026	20384	113726	0	18
	4	(99.68%)	(0.05%)	(0.27%)	(0.00%)	(0.00%)
H167-2	4170718	41500934	71170	135026	0	58
	8	(99.51%)	(0.17%)	(0.32%)	(0.00%)	(0.00%)
H167-3	3746515	37334922	21768	108344	0	120
	4	(99.65%)	(0.06%)	(0.29%)	(0.00%)	(0.00%)
H74-1	6829054	67974512	39216	276596	0	220
	4	(99.54%)	(0.06%)	(0.41%)	(0.00%)	(0.00%)
H74-2	4806825	47833516	63468	171110	0	160
	4	(99.51%)	(0.13%)	(0.36%)	(0.00%)	(0.00%)
H74-3	7158956	71355148	36824	197576	0	20
	8	(99.67%)	(0.05%)	(0.28%)	(0.00%)	(0.00%)
U3423-1	4095154	40829476	15596	106454	0	16
	2	(99.70%)	(0.04%)	(0.26%)	(0.00%)	(0.00%)
U3423-2	4385412	43714232	23252	116612	0	24
	0	(99.68%)	(0.05%)	(0.27%)	(0.00%)	(0.00%)
U3423-3	3995531	39827988	23998	103326	0	0 (0.00%)
	2	(99.68%)	(0.06%)	(0.26%)	(0.00%)	
U6-1	4097434	40833328	25276	115730	0	14
	8	(99.66%)	(0.06%)	(0.28%)	(0.00%)	(0.00%)

U6-2	4669716	46548458	22422	126256	0	24
	0	(99.68%)	(0.05%)	(0.27%)	(0.00%)	(0.00%)
U6-3	5406426	53885742	41442	136928	0	154
	6	(99.67%)	(0.08%)	(0.25%)	(0.00%)	(0.00%)
T15-1	4860588	48456232	39572	110056	0	20
	0	(99.69%)	(0.08%)	(0.23%)	(0.00%)	(0.00%)
T15-2	4905240	48898474	35946	117960	0	22
	2	(99.69%)	(0.07%)	(0.24%)	(0.00%)	(0.00%)
T15-3	4246613	42346634	27200	92264	0	38
	6	(99.72%)	(0.06%)	(0.22%)	(0.00%)	(0.00%)

Note: Sample is the sample ID; RawDatas is the number of raw reads downloaded from the device; CleanData (%) is the number of high-quality reads and percentage (based on RawReads); Adapter (%) is the number of reads containing adapters and the percentage (based on RawReads); LowQuality (%) refers to low-quality reads. The number and percentage of reads with a quality value $Q \leq 20$ for more than 50% of the bases in single ended reads (based on RawReads); PolyA (%) is the number of polyA reads and the percentage (based on RawReads); N (%) is the number of high filtering reads with N ratio and the percentage (based on RawReads).

Table S3 Base Information Statistics Table

Samp le	RawDat a(bp)	BF_Q20(%)	BF_Q30(%)	BF_N(%)	BF_GC(%)	CleanDat a(bp)	AF_Q20(%)	AF_Q30(%)	AF_N(%)	AF_GC(%)
H167	6259673	6189581836	6050856636	60194	3055811345	6185077	6127653797	5998875952	59226	3015901019
-1	100	(98.88%)	(96.66%)	(0.00%)	(48.82%)	958	(99.07%)	(96.99%)	(0.00%)	(48.76%)
H167	6256078	6153039286	5975092353	65297	3064908972	6161966	6074518585	5908080224	63533	3012122635
-2	200	(98.35%)	(95.51%)	(0.00%)	(48.99%)	425	(98.58%)	(95.88%)	(0.00%)	(48.88%)
H167	5619773	5552542762	5425154466	123754	2764272705	5539957	5486318795	5369523766	120312	2720642458
-3	100	(98.80%)	(96.54%)	(0.00%)	(49.19%)	095	(99.03%)	(96.92%)	(0.00%)	(49.11%)
H74-	1024358	10097286928	9821551154	217967	5025907051	1009853	9980122362	9725393894	211820	4945980394
1	1600	(98.57%)	(95.88%)	(0.00%)	(49.06%)	7704	(98.83%)	(96.30%)	(0.00%)	(48.98%)
H74-	7210238	7075893617	6852030220	126841	3551467499	7127759	7009231706	6795273701	123114	3505271837
2	100	(98.14%)	(95.03%)	(0.00%)	(49.26%)	892	(98.34%)	(95.34%)	(0.00%)	(49.18%)
H74-	1073843	10616472974	10377496330	104009	5290083207	1060074	10501869374	10281441382	102239	5216618184
3	5200	(98.86%)	(96.64%)	(0.00%)	(49.26%)	3444	(99.07%)	(96.99%)	(0.00%)	(49.21%)
U342	6142731	6077287743	5947827232	60795	3017205003	6080516	6026173016	5905251034	59897	2983609459
3-1	300	(98.93%)	(96.83%)	(0.00%)	(49.12%)	160	(99.11%)	(97.12%)	(0.00%)	(49.07%)
U342	6578118	6506123438	6363618920	63500	3238195302	6506267	6446938668	6314505927	62424	3198585350
3-2	000	(98.91%)	(96.74%)	(0.00%)	(49.23%)	046	(99.09%)	(97.05%)	(0.00%)	(49.16%)
U342	5993296	5925251185	5790699606	28928	2951668987	5906445	5852060538	5728511221	28463	2904117281
3-3	800	(98.86%)	(96.62%)	(0.00%)	(49.25%)	356	(99.08%)	(96.99%)	(0.00%)	(49.17%)
U6-1	6146152	6074057077	5935335968	58520	3048285451	6056906	5999976167	5873075356	57459	2999408106
	200	(98.83%)	(96.57%)	(0.00%)	(49.60%)	306	(99.06%)	(96.96%)	(0.00%)	(49.52%)
U6-2	7004574	6927416575	6776844968	68928	3466225327	6916462	6854130204	6715138463	67698	3418790869
	000	(98.90%)	(96.75%)	(0.00%)	(49.49%)	007	(99.10%)	(97.09%)	(0.00%)	(49.43%)

U6-3	8109639	8012910696	7829774866	180432	4024188190	7980871	7904592790	7738213336	175358	3952569103
	900	(98.81%)	(96.55%)	(0.00%)	(49.62%)	613	(99.04%)	(96.96%)	(0.00%)	(49.53%)
T15-1	7290882	7209573336	7052843553	71971	3553187354	7178515	7115368740	6973527095	70462	3492596029
	000	(98.88%)	(96.74%)	(0.00%)	(48.73%)	214	(99.12%)	(97.14%)	(0.00%)	(48.65%)
T15-2	7357860	7277578241	7119993735	71863	3579686349	7263831	7199415393	7054867368	70667	3527741181
	300	(98.91%)	(96.77%)	(0.00%)	(48.65%)	866	(99.11%)	(97.12%)	(0.00%)	(48.57%)
T15-3	6369920	6301435899	6166379715	102415	3089957953	6293877	6237676937	6112427139	100736	3048545361
	400	(98.92%)	(96.80%)	(0.00%)	(48.51%)	921	(99.11%)	(97.12%)	(0.00%)	(48.44%)

Note: Sample is the name of the sample; BF (Before filter) is the base information of the sample before filtering; AF (After filter) is the base information of the filtered sample; RawData (bp) is the total number of bases in the downloaded data (Unit bp); CleanData (bp) is the total number of bases in high-quality filtered data (Unit bp); Q20 (%) is the number of bases with sequencing base quality values above Q20 and the percentage of them in RawData (or CleanData); Q30 (%) is the number of bases with sequencing base quality values above Q30 and the percentage of them in RawData (or CleanData); N (%) is the number of N bases in a single ended read and the percentage of it in RawData (or CleanData); GC (%) is the proportion of GC bases in the sequence before (after) filtering.

Table S4 Alignment ribosome statistics

Sample	clean_reads	Mapped_Reads(%)	Unmapped_Reads(%)
H167-1	41597026	199528 (0.48%)	41397498 (99.52%)
H167-2	41500934	207520 (0.50%)	41293414 (99.50%)
H167-3	37334922	186710 (0.50%)	37148212 (99.50%)
H74-1	67974512	417820 (0.61%)	67556692 (99.39%)
H74-2	47833516	302714 (0.63%)	47530802 (99.37%)
H74-3	71355148	645324 (0.90%)	70709824 (99.10%)
U3423-1	40829476	238938 (0.59%)	40590538 (99.41%)
U3423-2	43714232	221606 (0.51%)	43492626 (99.49%)
U3423-3	39827988	254076 (0.64%)	39573912 (99.36%)
U6-1	40833328	362020 (0.89%)	40471308 (99.11%)
U6-2	46548458	378124 (0.81%)	46170334 (99.19%)
U6-3	53885742	1067772 (1.98%)	52817970 (98.02%)
T15-1	48456232	972886 (2.01%)	47483346 (97.99%)
T15-2	48898474	1031604 (2.11%)	47866870 (97.89%)
T15-3	42346634	318322 (0.75%)	42028312 (99.25%)

Note: Sample is the sample name; Clean_reads is the number of high-quality reads; Mapped_Reads (%) is the number and percentage of reads (based on clean reads) compared to the ribosomes of this species; Unmapped_Reads (%) is the number of reads that cannot be aligned with ribosomes and the percentage based on clean reads.

Table S5 Alignment of reference statistics

Sample	Total	Unmapped(%)	Unique_Mapped (%)	Multiple_Mapped (%)	Total_Mapped(%)
H167-1	413974	6638723 (16.04%)	33522923 (80.98%)	1235852 (2.99%)	34758775 (83.96%)
H167-2	412934	6508677 (15.76%)	33618896 (81.41%)	1165841 (2.82%)	34784737 (84.24%)
H167-3	371482	5531447 (14.89%)	30492378 (82.08%)	1124387 (3.03%)	31616765 (85.11%)
H74-1	675566	10058113 (14.89%)	55428636 (82.05%)	2069943 (3.06%)	57498579 (85.11%)
H74-2	475308	6778032 (14.26%)	39313709 (82.71%)	1439061 (3.03%)	40752770 (85.74%)
H74-3	707098	10113427 (14.30%)	58305820 (82.46%)	2290577 (3.24%)	60596397 (85.70%)
U3423-1	405905	5647225 (13.91%)	33753890 (83.16%)	1189423 (2.93%)	34943313 (86.09%)
U3423-2	434926	5938218 (13.65%)	36266688 (83.39%)	1287720 (2.96%)	37554408 (86.35%)

U3423	395739	5398141	32990426		34175771
-3	12	(13.64%)	(83.36%)	1185345 (3.00%)	(86.36%)
U6-1	404713	5946955	33296587		34524353
	08	(14.69%)	(82.27%)	1227766 (3.03%)	(85.31%)
U6-2	461703	6795909	37957579		39374425
	34	(14.72%)	(82.21%)	1416846 (3.07%)	(85.28%)
U6-3	528179	7767967	43130980		45050003
	70	(14.71%)	(81.66%)	1919023 (3.63%)	(85.29%)
T15-1	474833	7816610	37933323		39666736
	46	(16.46%)	(79.89%)	1733413 (3.65%)	(83.54%)
T15-2	478668	7912497	38165331		39954373
	70	(16.53%)	(79.73%)	1789042 (3.74%)	(83.47%)
T15-3	420283	6977344	33753784		35050968
	12	(16.60%)	(80.31%)	1297184 (3.09%)	(83.40%)

Note: Sample is the sample name; Total is the number of reads after filtering ribosomes, known as effective reads; Unmapped (%) refers to the number of reads that have not been aligned with the reference genome and the proportion of effective reads; Unique_Mapped (%) is the number of reads and the proportion of valid reads that are uniquely aligned with the reference genome; Multiple_Mapped (%) refers to the number of reads and the proportion of effective reads on the reference genome compared at multiple locations; Total_Mapped (%) is the total number of reads that can be located on the genome and the proportion of effective reads.

Table S6 Align the statistical situation of the reference area

sample	exon	intron	intergenic
H167-1	31032085 (89.28%)	2975531 (8.56%)	751159 (2.16%)
H167-2	31173635 (89.62%)	2879186 (8.28%)	731916 (2.10%)
H167-3	28436201 (89.94%)	2520824 (7.97%)	659740 (2.09%)
H74-1	50962373 (88.63%)	5254562 (9.14%)	1281644 (2.23%)
H74-2	36442928 (89.42%)	3431200 (8.42%)	878642 (2.16%)
H74-3	54206766 (89.46%)	5003492 (8.26%)	1386139 (2.29%)
U3423-1	31305294 (89.59%)	2872505 (8.22%)	765514 (2.19%)
U3423-2	33691031 (89.71%)	3055756 (8.14%)	807621 (2.15%)
U3423-3	30622322 (89.60%)	2800963 (8.20%)	752486 (2.20%)
U6-1	30949705 (89.65%)	2827539 (8.19%)	747109 (2.16%)
U6-2	35300008 (89.65%)	3226347 (8.19%)	848070 (2.15%)
U6-3	39690801 (88.10%)	4151243 (9.21%)	1207959 (2.68%)
T15-1	34731960 (87.56%)	3760363 (9.48%)	1174413 (2.96%)
T15-2	34838736 (87.20%)	3890991 (9.74%)	1224646 (3.07%)
T15-3	31249454 (89.15%)	2972689 (8.48%)	828825 (2.36%)

Table S7 New gene annotation table

Gene ID	Description	K	K	P	K	GO Component	GO Function	GO Process	Taxonomy
MSR001	KAK3444190.1 taxid=71139 hypothetical protein	-	-	-	-	-	GO:0003677// DNA binding	GO:0006355//regulation of transcription, DNA-templated	-
MSR002	EUGRSUZ_A000 32 [Eucalyptus grandis]	-	-	-	-	-	GO:0033926//g lycopeptide alpha-N-acetyl galactosaminid ase activity	-	-

[illegible]

[illegible]

ribosomal
subunit;GO:0005840//ribosome;GO:0015935//
small ribosomal
subunit;GO:0031974//membrane-enclosed
lumen;GO:0032991//macromolecular
complex;GO:0043226//organelle;GO:0043227
//membrane-bounded
organelle;GO:0043228//non-membrane-bound
ed organelle;GO:0043229//intracellular
organelle;GO:0043231//intracellular
membrane-bounded
organelle;GO:0043232//intracellular
non-membrane-bounded
organelle;GO:0043233//organelle
lumen;GO:0044391//ribosomal
subunit;GO:0044422//organelle
part;GO:0044424//intracellular
part;GO:0044429//mitochondrial
part;GO:0044444//cytoplasmic
part;GO:0044446//intracellular organelle
part;GO:0044464//cell
part;GO:0070013//intracellular organelle
lumen;GO:0098798//mitochondrial protein
complex;GO:1990904//ribonucleoprotein
complex

expression;GO:0019538//protein metabolic
process;GO:0032543//mitochondrial
translation;GO:0034641//cellular nitrogen
compound metabolic
process;GO:0034645//cellular macromolecule
biosynthetic process;GO:0043043//peptide
biosynthetic
process;GO:0043170//macromolecule
metabolic process;GO:0043603//cellular amide
metabolic process;GO:0043604//amide
biosynthetic process;GO:0044237//cellular
metabolic process;GO:0044238//primary
metabolic process;GO:0044249//cellular
biosynthetic process;GO:0044260//cellular
macromolecule metabolic
process;GO:0044267//cellular protein
metabolic process;GO:0044271//cellular
nitrogen compound biosynthetic
process;GO:0071704//organic substance
metabolic
process;GO:1901564//organonitrogen
compound metabolic
process;GO:1901566//organonitrogen
compound biosynthetic
process;GO:1901576//organic substance
biosynthetic process

MSR	GO	GO ID	GO Description	GO Type
XP_039163050.1	taxid=71139 LOW			
QUALITY				
PROTEIN: G-type				
lectin				
S-receptor-like	-	-	-	-
serine/threonine-pr				
otein kinase				
LECRK3				
[Eucalyptus				
grandis]				
MSR				
GO	-	-	-	-
.1	-	-	-	-
0				
0				
7				
6				
MS				
T	-	-	-	-
R	-	-	-	-
G				

.1
0
0
8
9

Note: GeneID is a new gene ID; Symbol is the name of a gene; Description is an Nr annotation; KEGG_A_class is the first level annotation of KEGG; KEGG_B_class is the second level annotation of KEGG; Pathway is a KEGG pathway annotation; GO Component is a component annotation for GO cells; GO Function is a functional annotation of GO molecules; GO Process is a GO biological process annotation; TF_family is a transcription factor annotation (This table only displays the top ten new genes. Please refer to the complete Excel spreadsheet 'New Gene Annotation' for all new genes.).

Table S8 Genetic testing statistics

samp le	Refer_ Genes	sequenced_Refer_ _Genes(%)	Novel_ Genes	sequenced_Novel_ _Genes(%)	Total_ Genes	sequenced_Total_ _Genes(%)
all	36005	30215 (83.92%)	3071	3071 (100.00%)	39076	33286 (85.18%)
H16 7-1	36005	24903 (69.17%)	3071	2260 (73.59%)	39076	27163 (69.51%)
H16 7-2	36005	24404 (67.78%)	3071	2228 (72.55%)	39076	26632 (68.15%)
H16 7-3	36005	24662 (68.50%)	3071	2233 (72.71%)	39076	26895 (68.83%)
H74- 1	36005	24577 (68.26%)	3071	2124 (69.16%)	39076	26701 (68.33%)
H74- 2	36005	25539 (70.93%)	3071	2292 (74.63%)	39076	27831 (71.22%)
H74- 3	36005	24079 (66.88%)	3071	2163 (70.43%)	39076	26242 (67.16%)
U34 23-1	36005	23151 (64.30%)	3071	2123 (69.13%)	39076	25274 (64.68%)
U34 23-2	36005	23213 (64.47%)	3071	2095 (68.22%)	39076	25308 (64.77%)
U34 23-3	36005	23112 (64.19%)	3071	2101 (68.41%)	39076	25213 (64.52%)
U6-1	36005	24650 (68.46%)	3071	2132 (69.42%)	39076	26782 (68.54%)
U6-2	36005	24724 (68.67%)	3071	2148 (69.94%)	39076	26872 (68.77%)
U6-3	36005	25033 (69.53%)	3071	2183 (71.08%)	39076	27216 (69.65%)
T15- 1	36005	24301 (67.49%)	3071	2166 (70.53%)	39076	26467 (67.73%)
T15- 2	36005	24098 (66.93%)	3071	2129 (69.33%)	39076	26227 (67.12%)
T15- 3	36005	23943 (66.50%)	3071	2089 (68.02%)	39076	26032 (66.62%)

Note: Refer_Genes is the total number of genes in the reference genome (or set of reference genes), and when there are no new genes, there are no results in this column; Sequenced_Refer_Genes (%) is the total number and percentage of Referr_Genes detected in the

sequencing results. When there are no new genes present, there is no result in this column; Novel_Genes is the number of new genes discovered in the project. When there are no new genes, there will be no results in this column; Sequenced_Novel_Genes(%) is the total number and percentage of Novel Genes detected in the sequencing results. When there are no new genes present, there will be no results in this column; Total_Genes is the total number of genes, including the reference genome and new genes; Sequenced_Total_Genes(%) is the total number of genes detected in the sequencing results and the percentage of them to All_Genes.

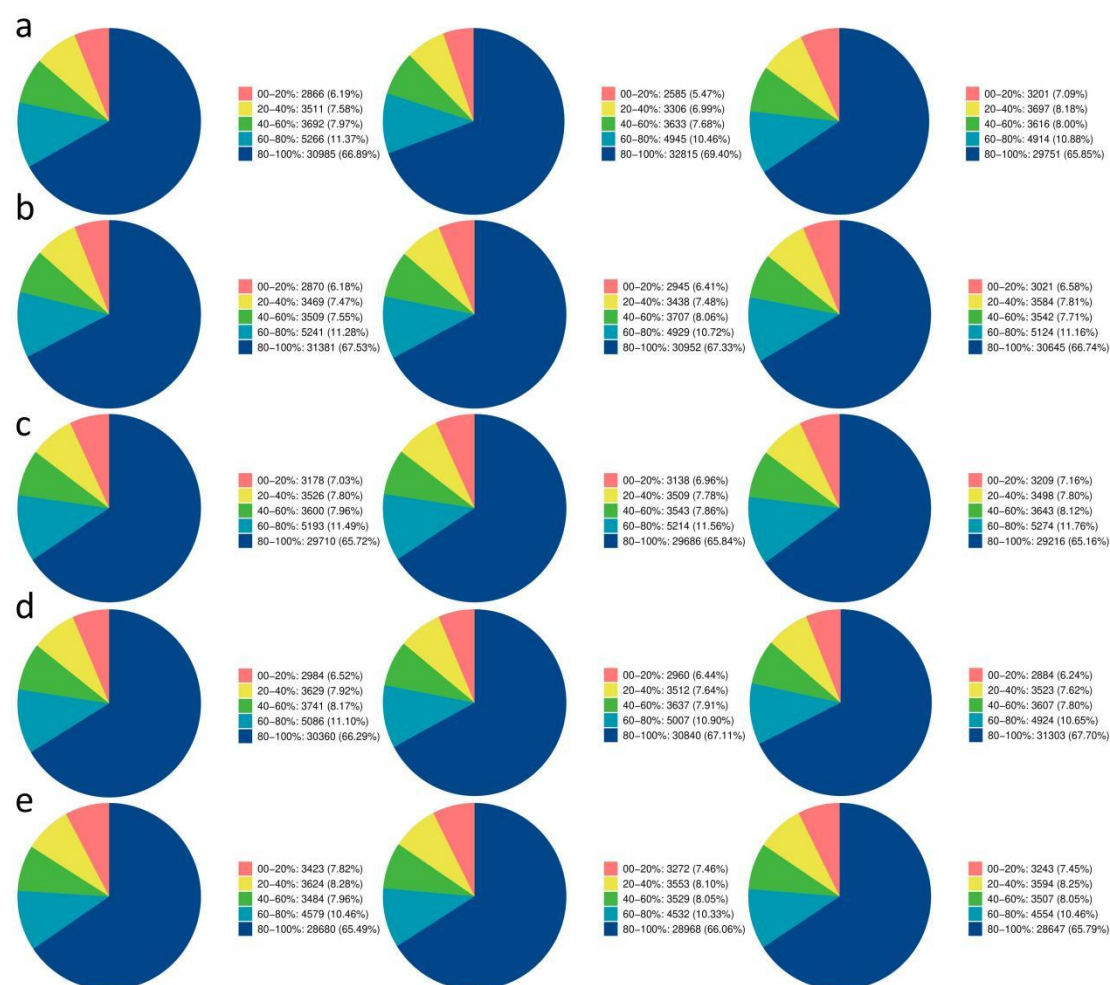


Figure S4 Gene coverage distribution (Note: The different colors in the figure represent the proportion of genes with a certain coverage range. a: H74; b: H167; c: T15; d: U6; e: U3423.)

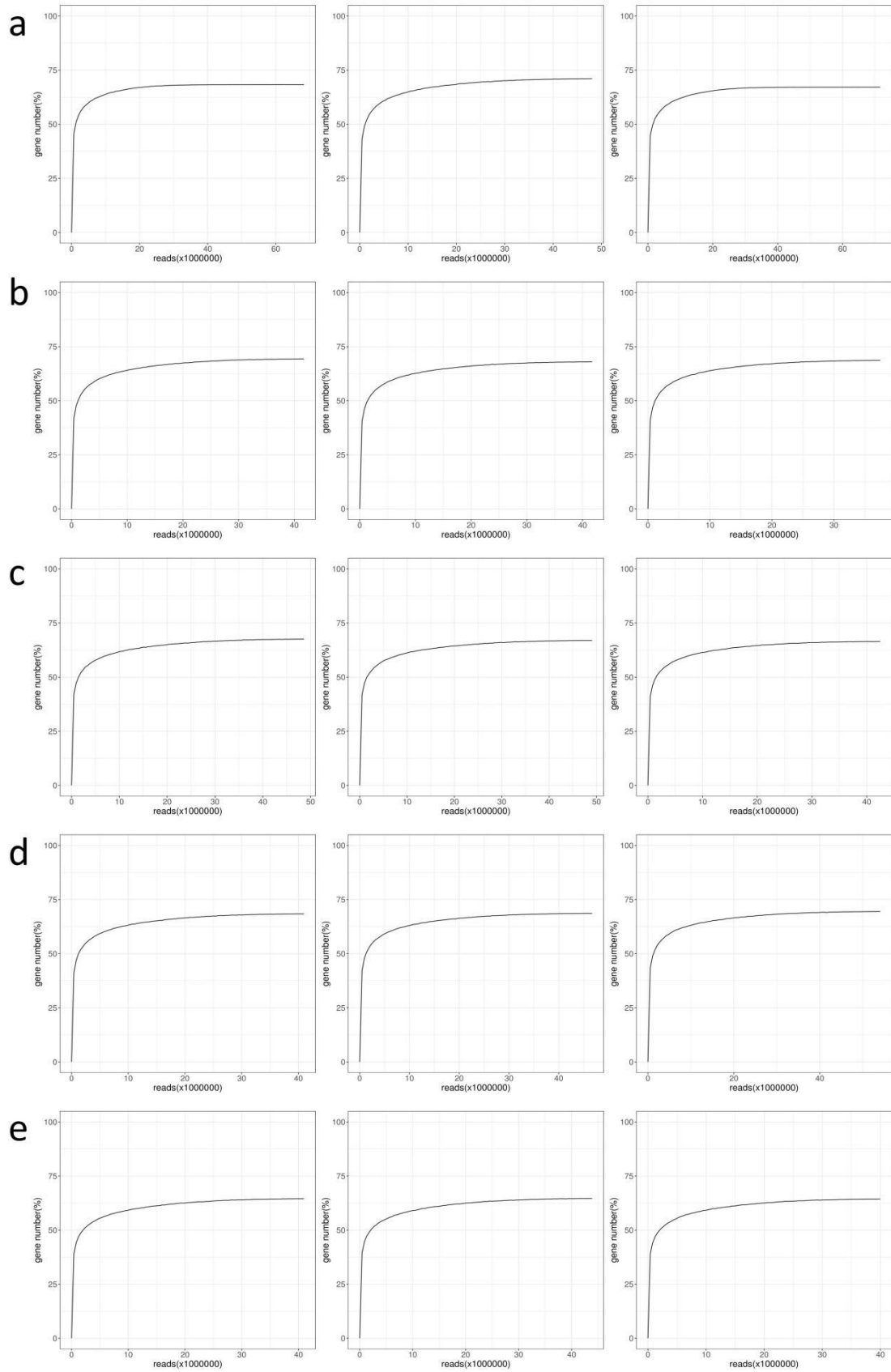


Figure S5 Sequencing saturation distribution (Note: a: H74; b: H167; c: T15; d: U6; e: U3423.)

[illegible]

RSUZ
_A00
032
[Euca
lyptus
grandi
s]

GO:
0033
926//
glyc
opep
tide
alph
a-N-
acety
lgala
ctosa
mini
dase
activ
ity

M
S
T
R
G
.
1
0
0
0

M
S
T
R

3 0 3 1 1 7 0 0 1 0 1 7 8 6 3 5 3 1 8 6 6 1 1 1 1 8
0 1 0 1 0 3 1 8 . 4 4 2 3 . 0 3 . . . 1 0 3
0 5 1 3 8 5 4 0 2 0 1 0 3 1 8 . 9 . . . 1 0 3
1 1 1 2 4 7 2 9 4 3 7 5 5 1 7 4 7 9 8 5 4 8 2 5 4 3 1 7
0 1 2 7 1 2 4 7 3 5 8 4 2 5 8 3 3 9 2 2 5

1 C XP_0
O 18727
D 310.2
M taxid=

GO:0005886//p
lasma
membrane

-
-
-
-
-
-
-
-

[illegible]

M
S
T
R
G
. 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 0 3 7 0 0 0 0 0 0 0 0 0 0 0 0 0 0 3 1 2
1
0
0
2
6

1 0 0
. . .
4 5 9
8 4 7

EUG
RSUZ
_C04
396
[Euca
lyptus
grandi
s]
KAK
34395
35.1
taxid=
71139
hypot
hetica
l
protei
n
EUG
RSUZ
_C04
401
[Euca
lyptus
grandi
s]

- - - - -
- - - - -
- - - - -

ial	O:0005737//cyt	const	bolic
[Euca	oplasm;GO:00	ituen	process;GO:00
lyptus	05739//mitocho	t of	09058//biosynt
grandi	ndrion;GO:000	ribos	hetic
s]	5759//mitochon	ome;	process;GO:00
	drial	GO:	09059//macrom
	matrix;GO:000	0005	olecule
	5761//mitochon	198//	biosynthetic
	drial	struc	process;GO:00
	ribosome;GO:0	tural	09987//cellular
	005763//mitoch	mole	process;GO:00
	ondrial small	cule	10467//gene
	ribosomal	activ	expression;GO:
	subunit;GO:00	ity	0019538//prote
	05840//ribosom		in metabolic
	e;GO:0015935/		process;GO:00
	/small		32543//mitocho
	ribosomal		ndrial
	subunit;GO:00		translation;GO:
	31974//membra		0034641//cellul
	ne-enclosed		ar nitrogen
	lumen;GO:003		compound
	2991//macromo		metabolic
	lecular		process;GO:00
	complex;GO:0		34645//cellular
	043226//organe		macromolecule

lle;GO:004322	biosynthetic
7//membrane-b	process;GO:00
ounded	43043//peptide
organelle;GO:0	biosynthetic
043228//non-m	process;GO:00
embrane-bound	43170//macrom
ed	olecule
organelle;GO:0	metabolic
043229//intrac	process;GO:00
ellular	43603//cellular
organelle;GO:0	amide
043231//intrac	metabolic
ellular	process;GO:00
membrane-bou	43604//amide
nded	biosynthetic
organelle;GO:0	process;GO:00
043232//intrac	44237//cellular
ellular	metabolic
non-membrane	process;GO:00
-bounded	44238//primary
organelle;GO:0	metabolic
043233//organe	process;GO:00
lle	44249//cellular
lumen;GO:004	biosynthetic
4391//ribosoma	process;GO:00
l	44260//cellular

subunit;GO:0044422//organelle	macromolecule metabolic process;GO:0044267//cellular protein metabolic process;GO:0044271//cellular nitrogen compound biosynthetic process;GO:0071704//organic substance metabolic process;GO:1901564//organonitrogen compound metabolic process;GO:1901566//organonitrogen compound biosynthetic process;GO:19
-------------------------------	---

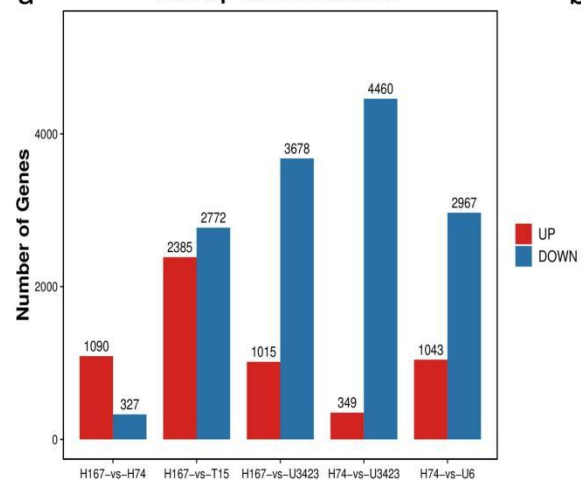
[illegible]

K3
[Euca
lyptus
grandi
s]

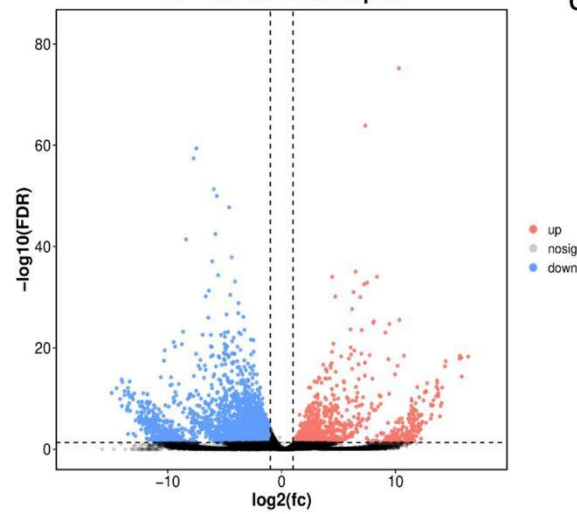
[illegible]

Note: Id is the gene ID; Xx_tpm is the tpm expression level of the corresponding gene for each sample; Xx_count is the original reads count of the corresponding genes for each sample; Symbol is the name of a gene; Description is an Nr annotation; KEGG_A_class is the first level annotation of KEGG; KEGG_B_class is the second level annotation of KEGG; Pathway is a KEGG pathway annotation; GO Component is a component annotation for GO cells; GO Function is a functional annotation of GO molecules; GO Process is a GO biological process annotation; TF_family is a transcription factor annotation (This table only displays the top ten genes. Please refer to the complete Excel spreadsheet 'Gene expression level table' for all new genes.).

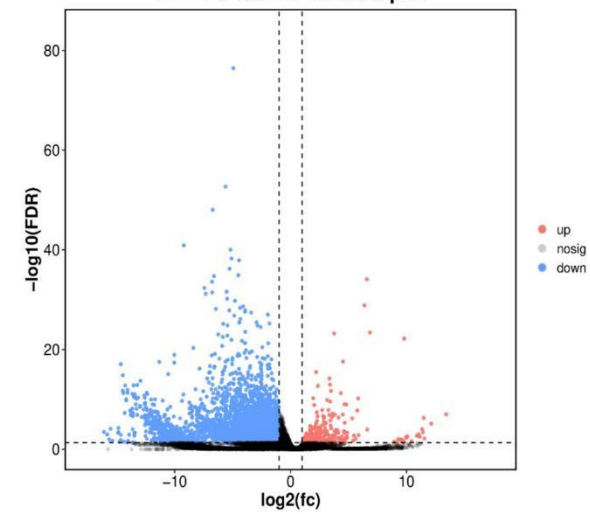
a DiffExp Genes Statistics



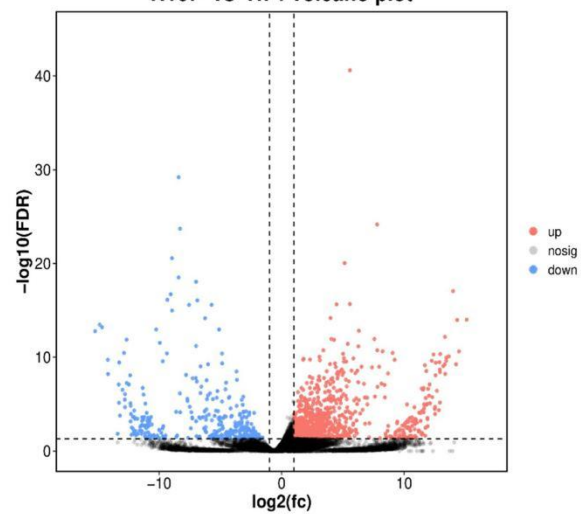
b H74-vs-U6 volcano plot



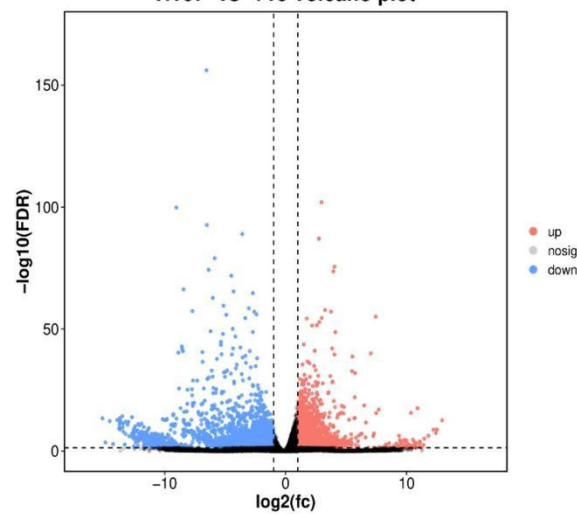
c H74-vs-U3423 volcano plot



d H167-vs-H74 volcano plot



e H167-vs-T15 volcano plot



f H167-vs-U3423 volcano plot

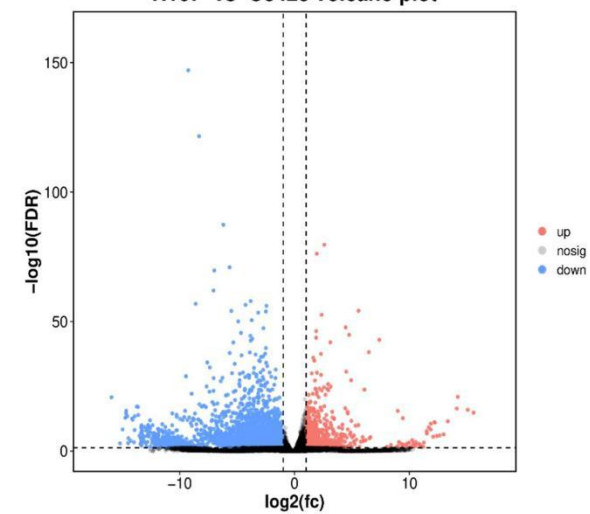


Figure S6 Analysis of inter group differences (Note: a: Differential gene statistical chart (X-axis: Samples compared in pairs; Y-axis: number of differentially expressed genes; Red represents upregulated differentially expressed genes; Blue represents differentially expressed genes that are downregulated.); b-f: Comparison of volcano maps with differences (The X-axis represents the logarithmic value of the difference multiple between two groups, and the Y-axis represents the negative Log10 value of the FDR of the difference between the two groups. The red (up-regulated expression level of group_2 relative to group_1) and blue (down regulated expression level) dots indicate differences in gene expression levels (with $FDR < 0.05$ and a difference factor of more than twice), while the black dots indicate no difference.).)

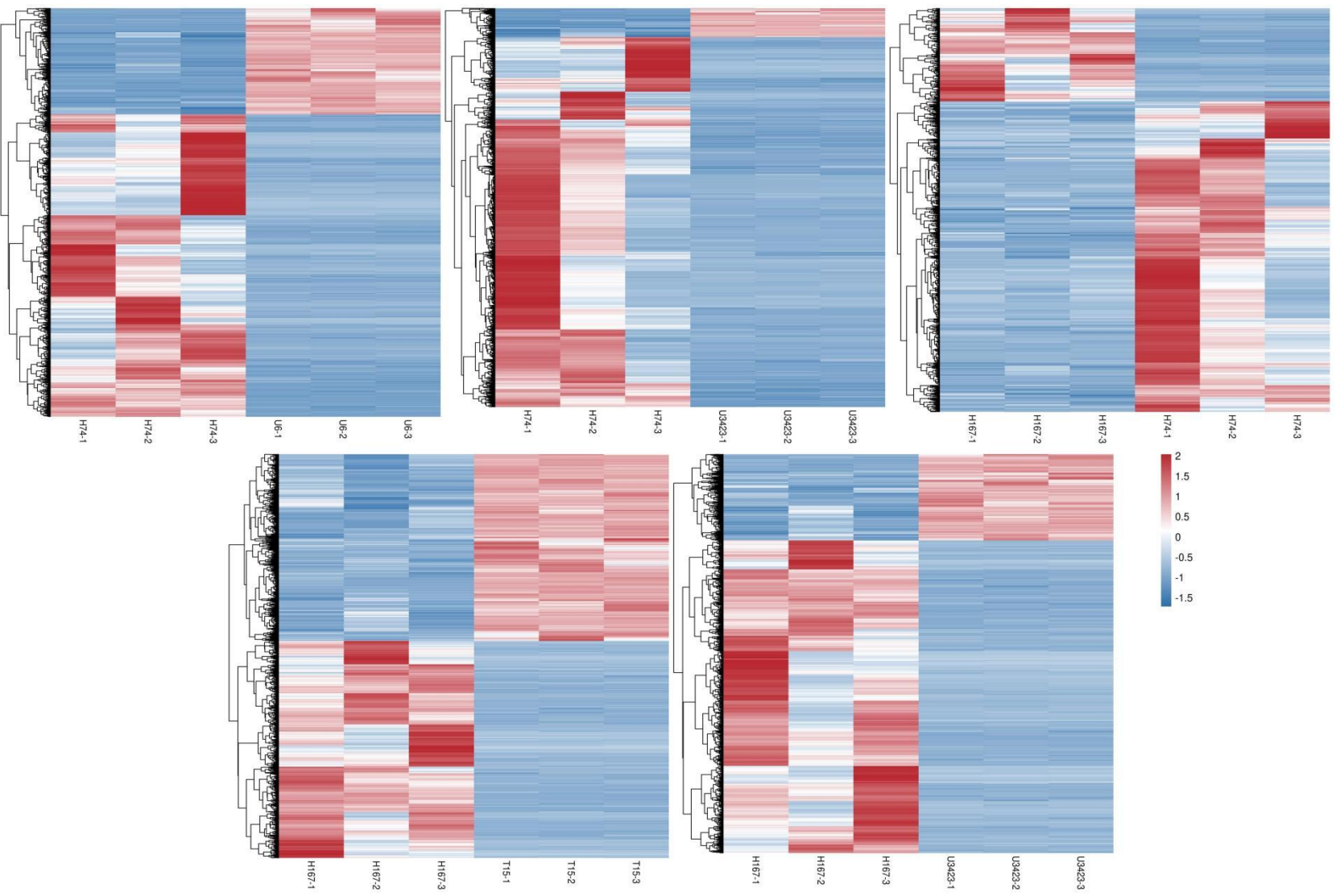


Figure S7 Differential comparison clustering heatmap (Note: Each column in the figure represents a sample, and each row represents a gene. The expression levels of genes in different samples are represented by different colors, with red indicating higher expression levels and blue indicating lower expression levels.)

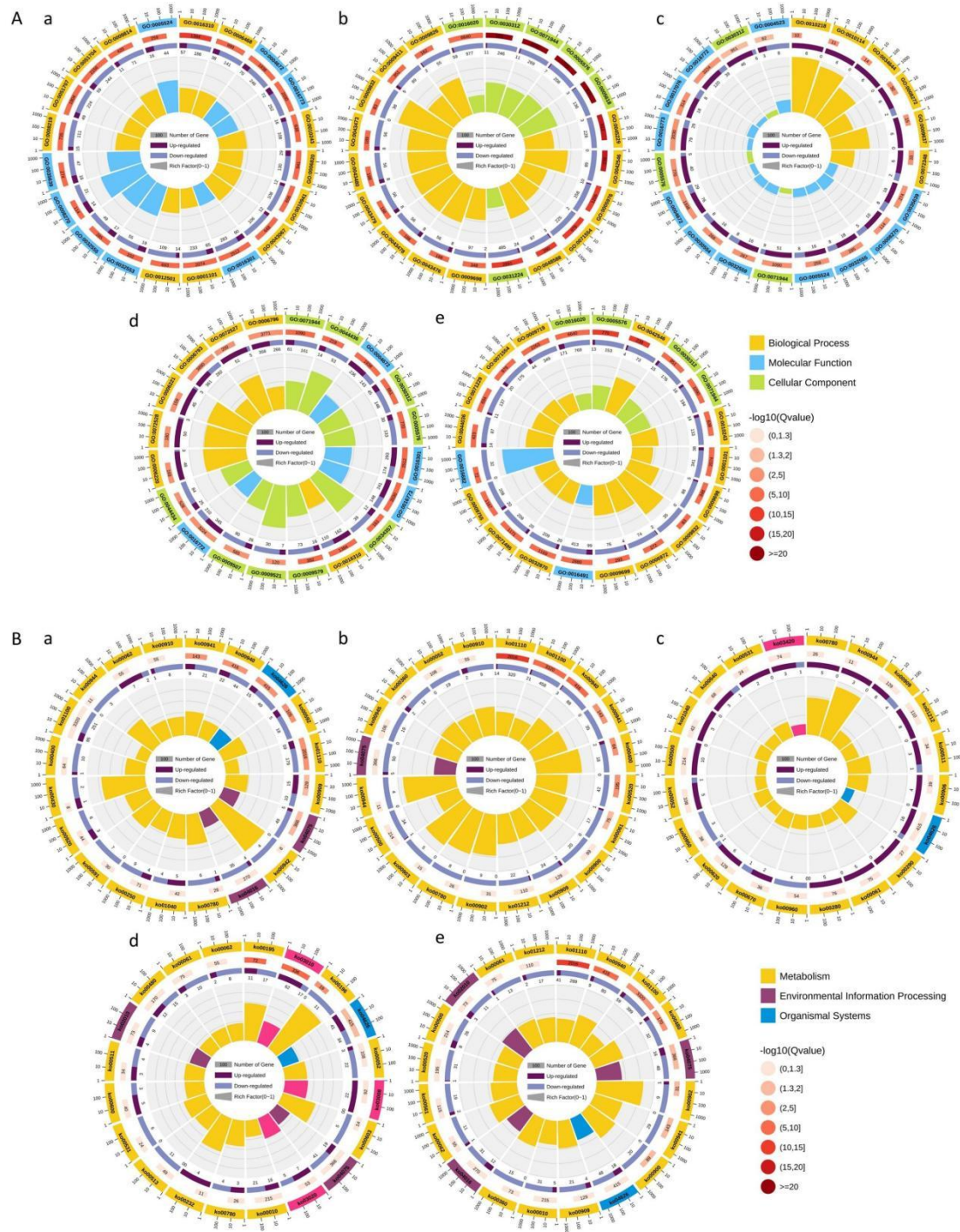


Figure S8 Enrichment circle diagram (Note: A: GO; B: KEGG. a: H74 vs U6; b: H74 vs U3423; c: H167 vs H74; d: H167 vs T15; e: H167 vs U3423. First circle: Enrich the top 20 GOterms (A)/pathway (B), and outside the circle is the coordinate ruler of the number of differentially expressed genes. Different colors represent different Ontology (A)/A class (B); Second circle: The number and Q value of the GOterm (A)/pathway (B) in the differential gene background. The more differential gene backgrounds there are, the longer the bars, and the smaller the Q value, the redder the color; Third circle: Bar chart of the proportion of upregulated and downregulated differentially expressed genes, with dark purple representing the proportion of upregulated differentially expressed genes and light purple representing the

proportion of downregulated differentially expressed genes; The specific numerical values are displayed below; Fourth circle: RichFactor values for each GOterm (A)/pathway (B) (the number of differentially expressed genes in the GOterm (A)/pathway (B) divided by all numbers in the GOterm (A)/pathway (B)), background grid lines, with each grid representing 0.1.)

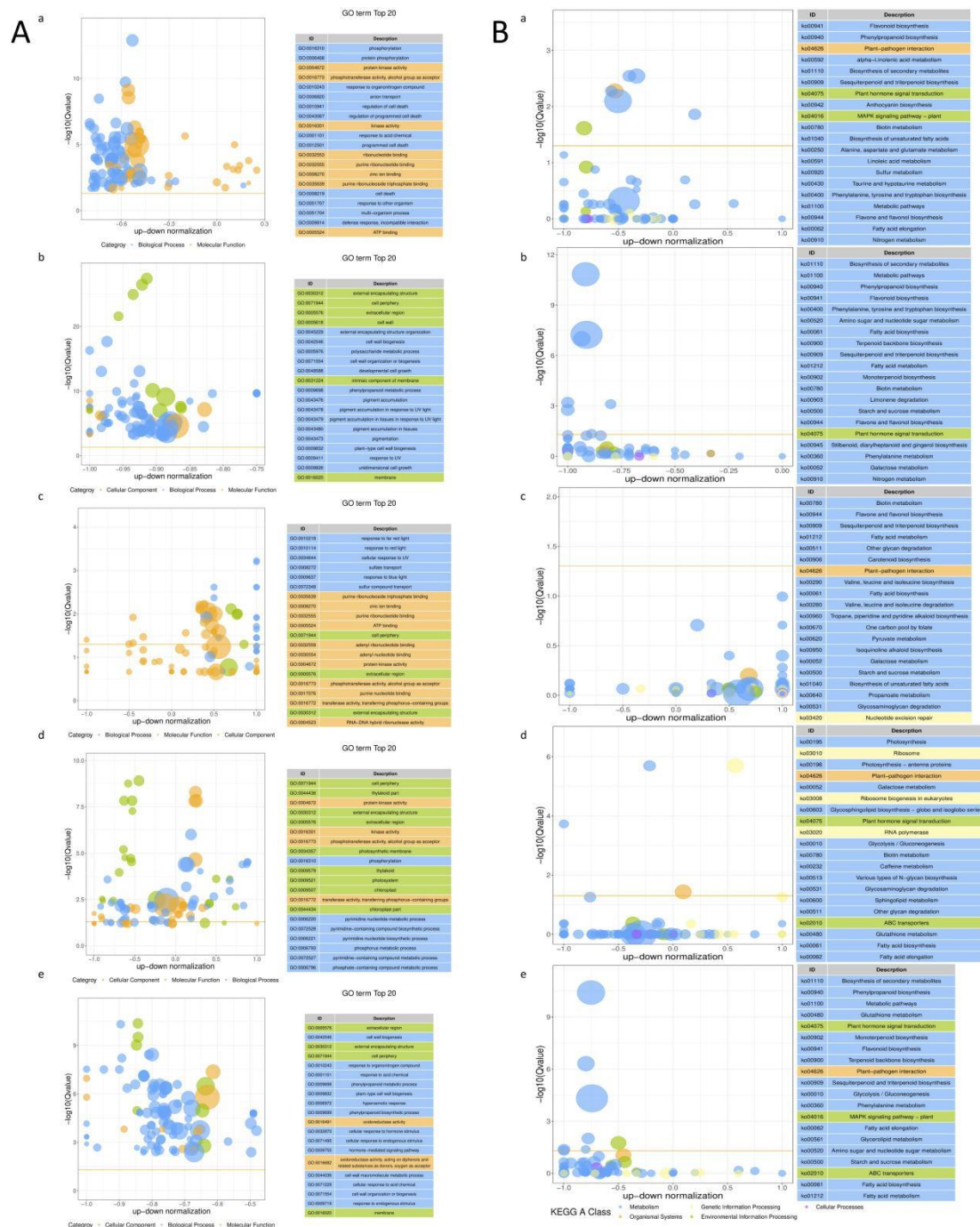


Figure S9 Enrichment difference bubble diagram (Note: A: GO; B: KEGG. a: H74 vs U6; b: H74 vs U3423; c: H167 vs H74; d: H167 vs T15; e: H167 vs U3423. The vertical axis is -log₁₀ (Qvalue), and the horizontal axis is the z-score value (the ratio of the difference between the number of up-regulated and down regulated differentially expressed genes to the

total differentially expressed genes). The yellow line represents the threshold of $Q_{\text{value}}=0.05$. On the right is the top 20 GOterm (A)/pathway (B) list of Q values. Different colors represent different Ontology (A)/A class (B); Draw separate graphs for different Ontology (A)/A class (B.).

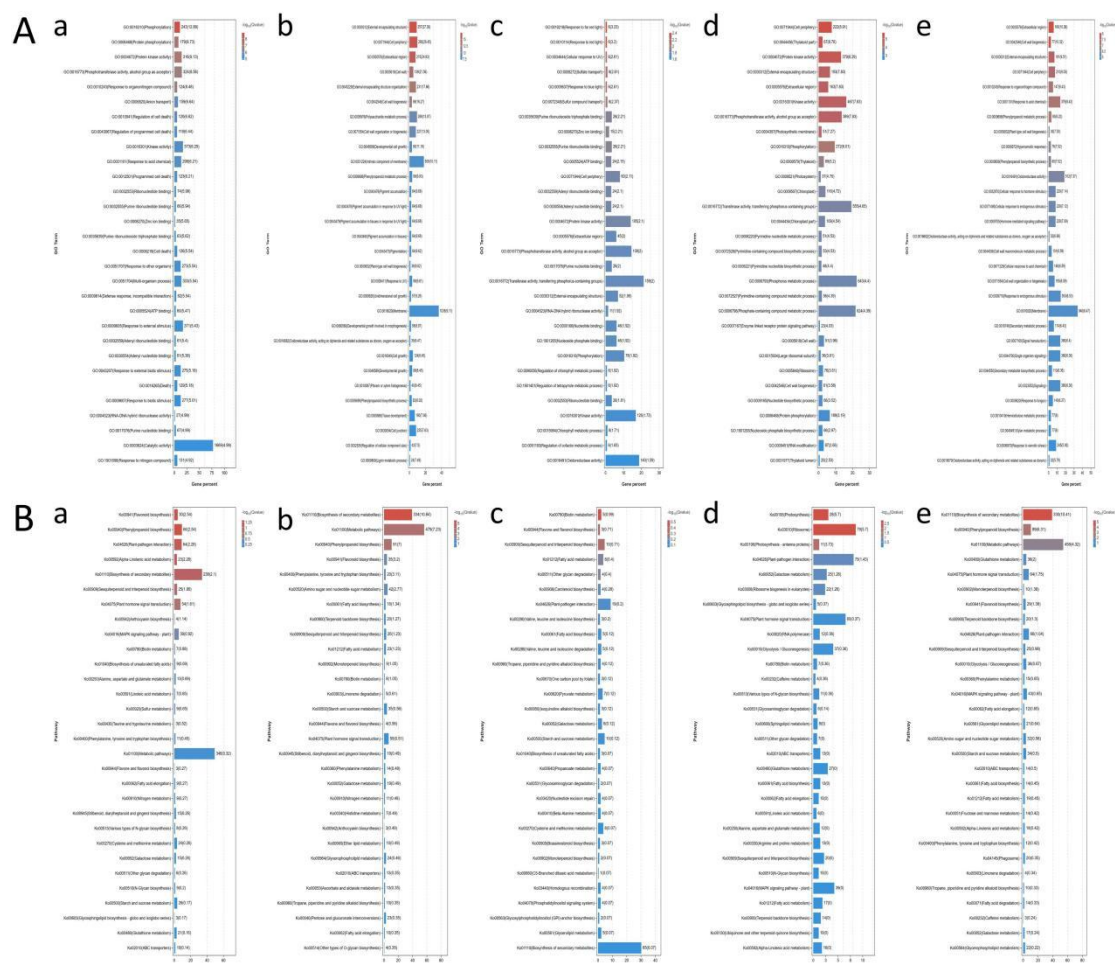


Figure S10 Enrichment bar diagram (Note: A: GO; B: KEGG. a: H74 vs U6; b: H74 vs U3423; c: H167 vs H74; d: H167 vs T15; e: H167 vs U3423. Use the top 20 GOterms (A)/pathway (B) with the lowest Q value to plot, with the vertical axis representing the number of GOterms (A)/pathway (B) and the horizontal axis representing the percentage of this GOterm (A)/pathway (B) to the total number of differentially expressed genes. The darker the color, the smaller the Q value. The values on the column represent the number of GOterms (A)/pathway (B) and their Q values.)

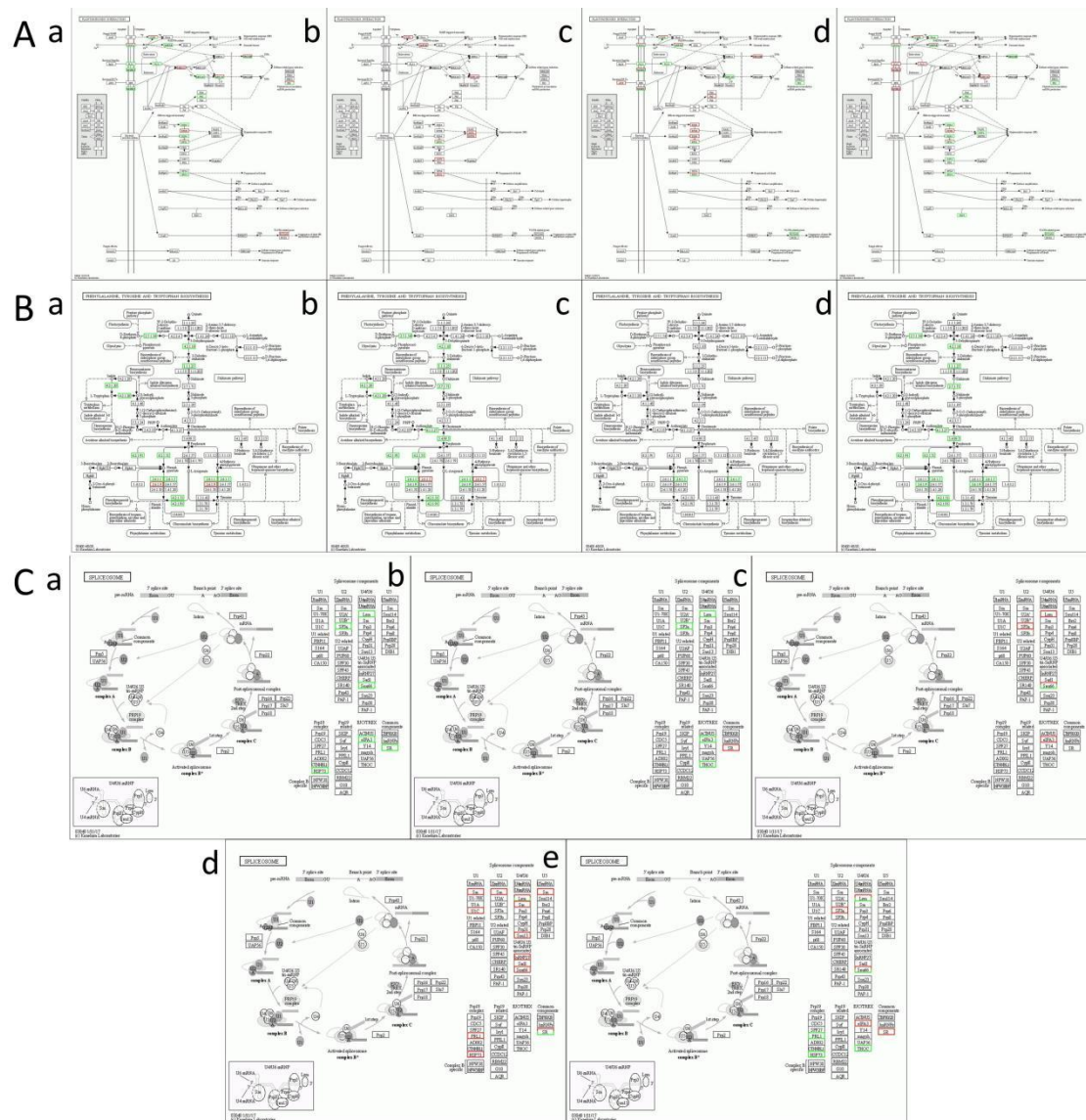


Figure S11 Detailed information on Pathways in KEGG database (Note: A: Plant-pathogen interaction (KO04626); a: H74 vs U6; b: H167 vs H74; c: H167 vs T15; d: H167 vs U3423; B: Phenylalanine, tyrosine and tryptophan biosynthesis (KO00400); a: H74 vs U6; b: H74 vs U3423; c: H167 vs H74; d: H167 vs U3423; C: Spliceosome (KO03040); a: H74 vs U6; b: H74 vs U3423; c: H167 vs H74; d: H167 vs T15; e: H167 vs U3423. The location of upregulated gene products is marked in red, while the location of downregulated gene products is marked in green.)

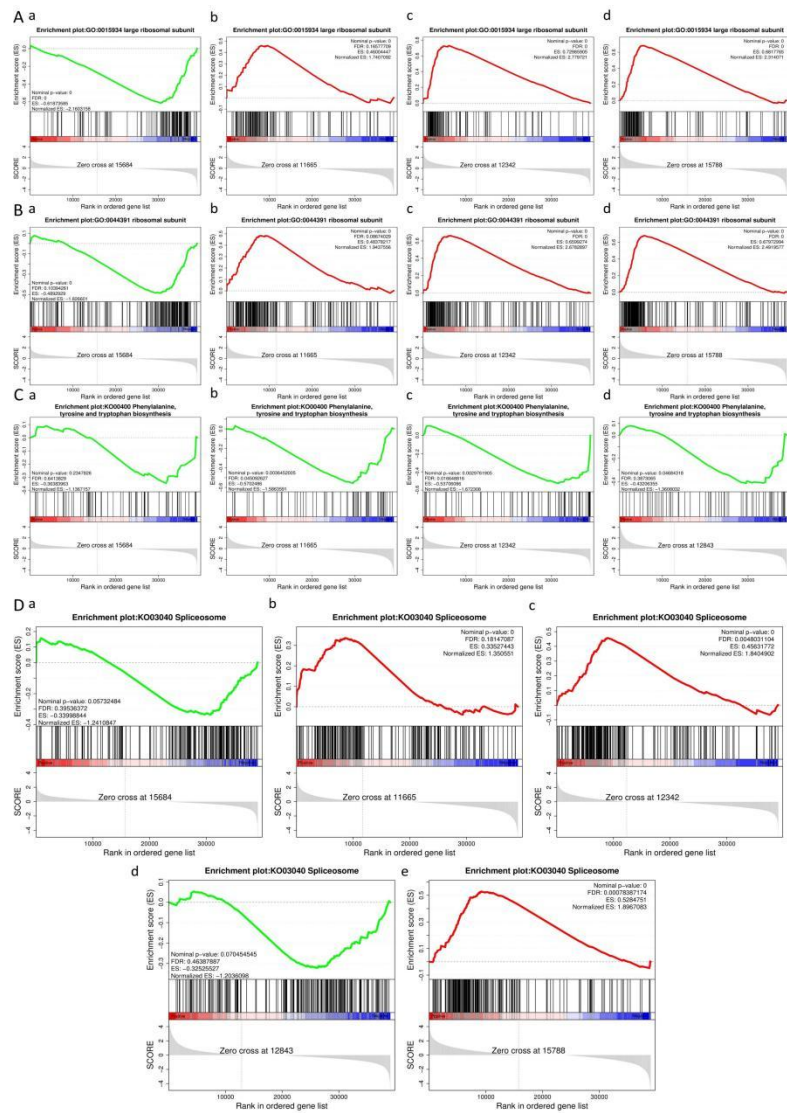


Figure S12 GSEA ES diagram (Note: A: GO:0015934; a: H167-H74, b: H167-U3423, c: H74-U3423, d: H167-T15. B: GO:0044391; a: H167-H74, b: H167-U3423, c: H74-U3423, d: H167-T15. C: KO00400; a: H167-H74, b: H167-U3423, c: H74-U3423, d: H74-U6. D: KO03040; a: H167-H74, b: H167-U3423, c: H74-U3423, d: H74-U6, e: H167-T15.)

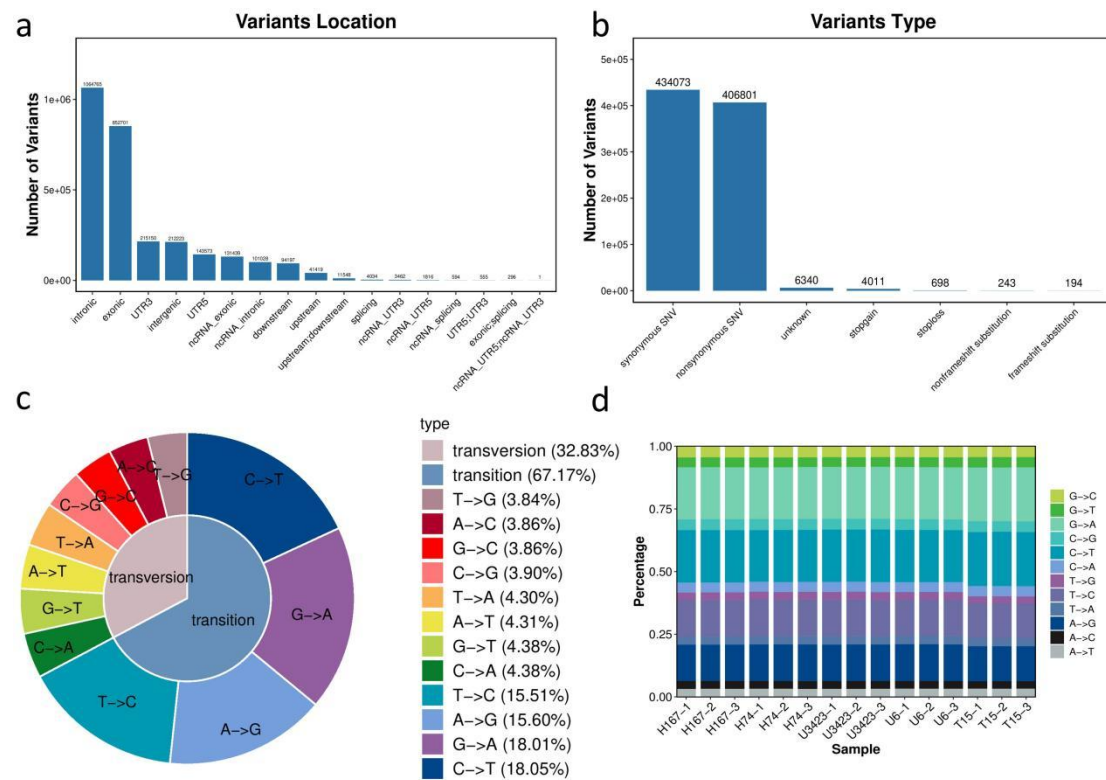


Figure S13 SNP/InDel analysis (Note: a: SNP location classification statistical chart; b: SNP functional classification statistical chart; c: SNP mutation statistical chart. Transversion refers to the substitution between purine and pyrimidine in base substitution (G↔T, C↔A, A↔T, C↔G), transition refers to the substitution of one purine with another purine or one pyrimidine with another pyrimidine (T↔C, A↔G). d: RNA Editing Classification Statistical Chart)

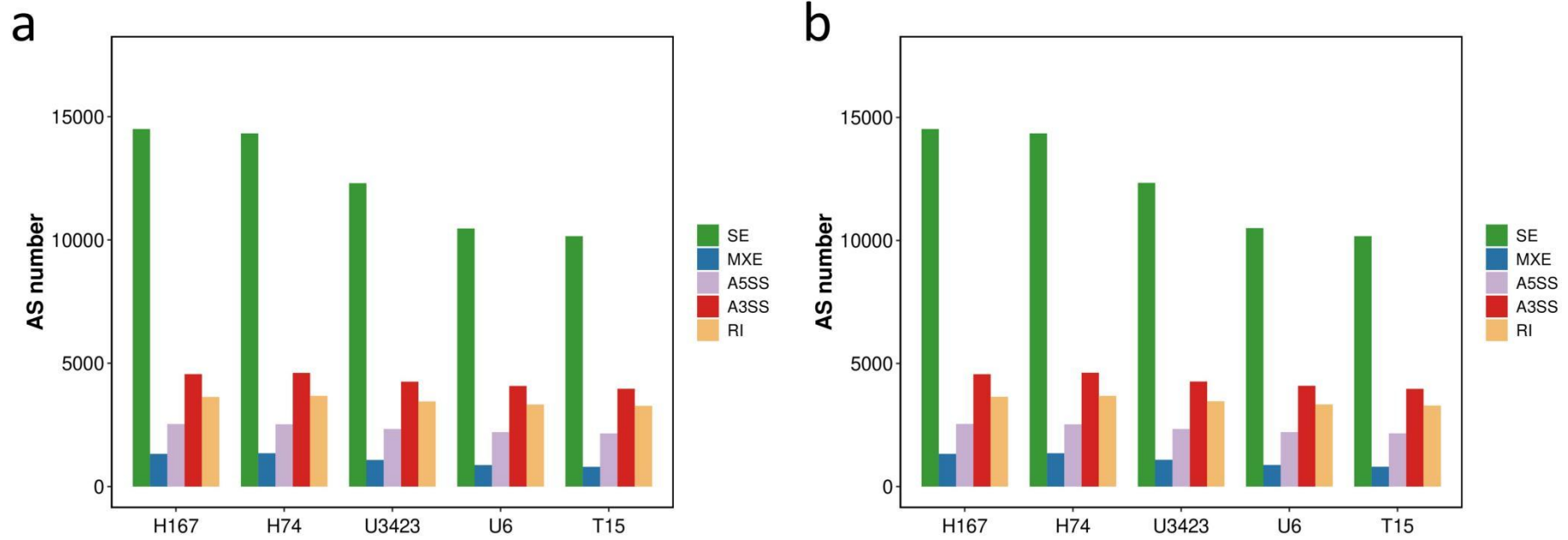


Figure S14 Variable cutting event statistics diagram (Note: The horizontal axis represents grouping, the vertical axis represents variable cutting quantity, and 5 columns correspond to 5 types of variable cutting events. a: JC Variable Shear Event Statistics Chart; b: JCEC Variable Shear Event Statistics Chart.)

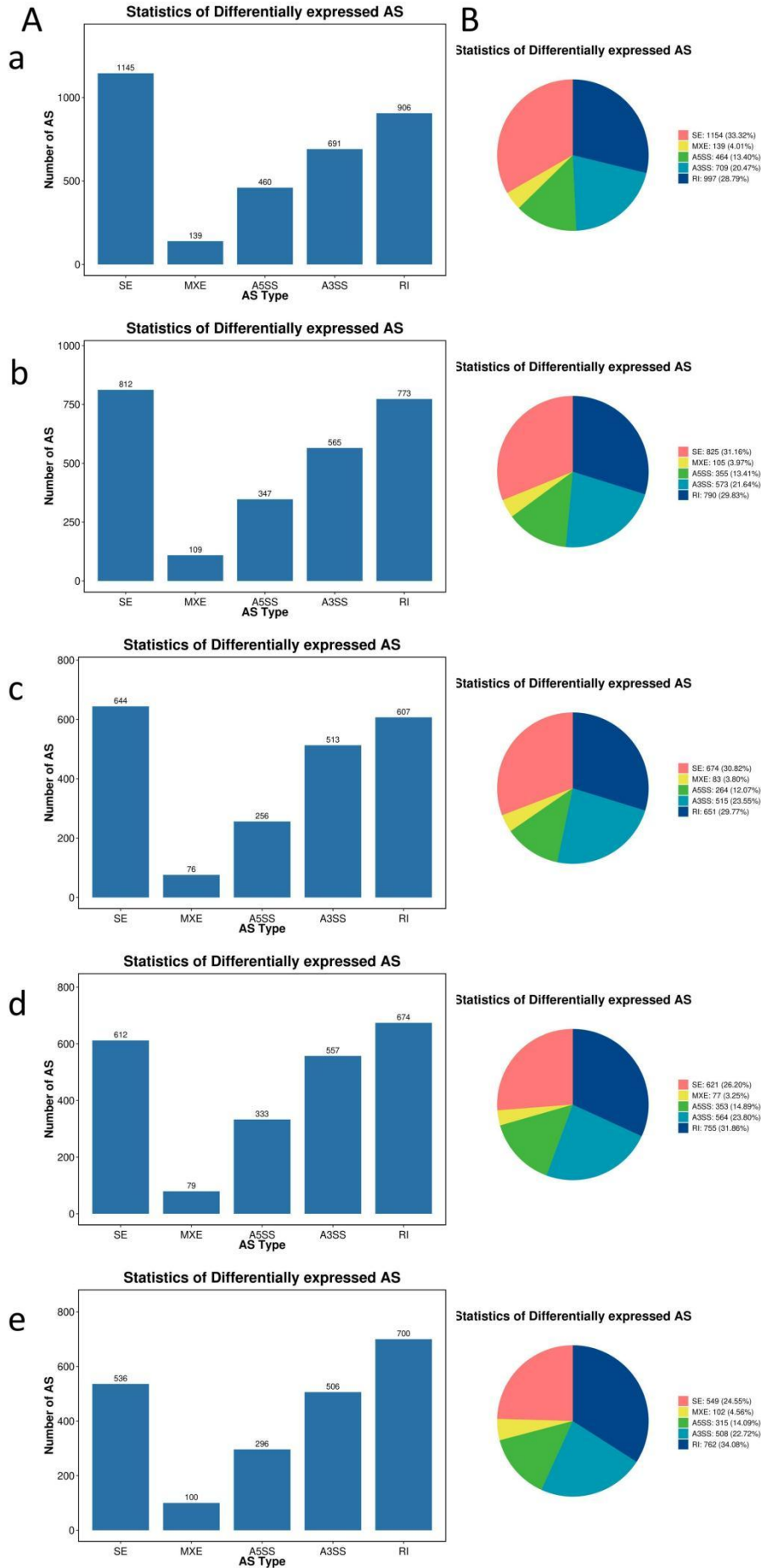


Figure S15 Differential AS classification and quantity statistical diagram (Note: A: JC only difference AS classification and quantity statistical chart; The expression level of the alternative splicing event Exon Skipping Isoform in the sample (the result of aligning reads between two constitutive exons, i.e. reads spanning upstream and downstream constitutive exons). B: JC+reads OnTarget difference AS classification and quantity statistics chart. The expression level of the alternative splicing event Exon Inclusion Isoform in the sample (reads the alternative splice junction between upstream constitutive exons and skipping exons; The reads are completely aligned to the skip type exon; Reads spans across the skip type exon and downstream constitutive exon alternative splice junction.) The number of reads for AS events in the sample is calculated using reads on target and junction counts. a: H74 vs U6; b: H74 vs U3423; c: H167 vs H74; d: H167 vs T15; e: H167 vs U3423.)

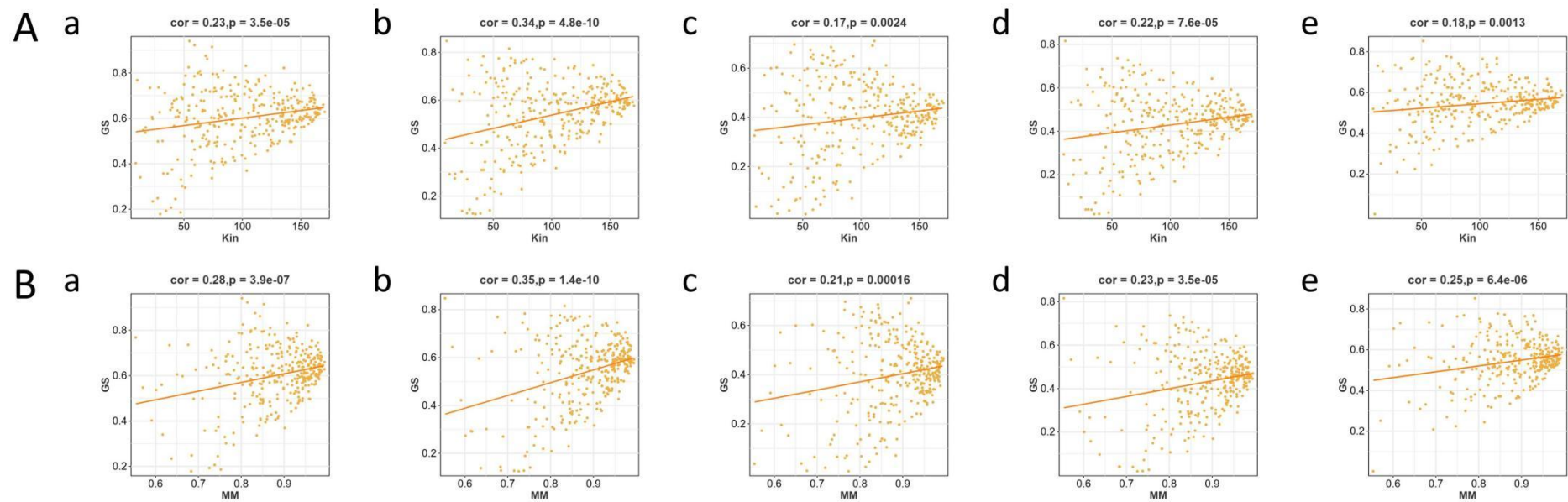


Figure S16 Scatter plot of correlation with GS (Note: A/B: Scatter plot of K.within/MM and GS correlation for the most positively correlated (left) and most negatively correlated (right) modules. The horizontal axis represents the K.within/MM value, the vertical axis represents the GS value, the dots represent genes, and the diagonal line represents the fitting line of correlation. a: HT; b: DBH; c: SUR; d: VOL; e: SS.)

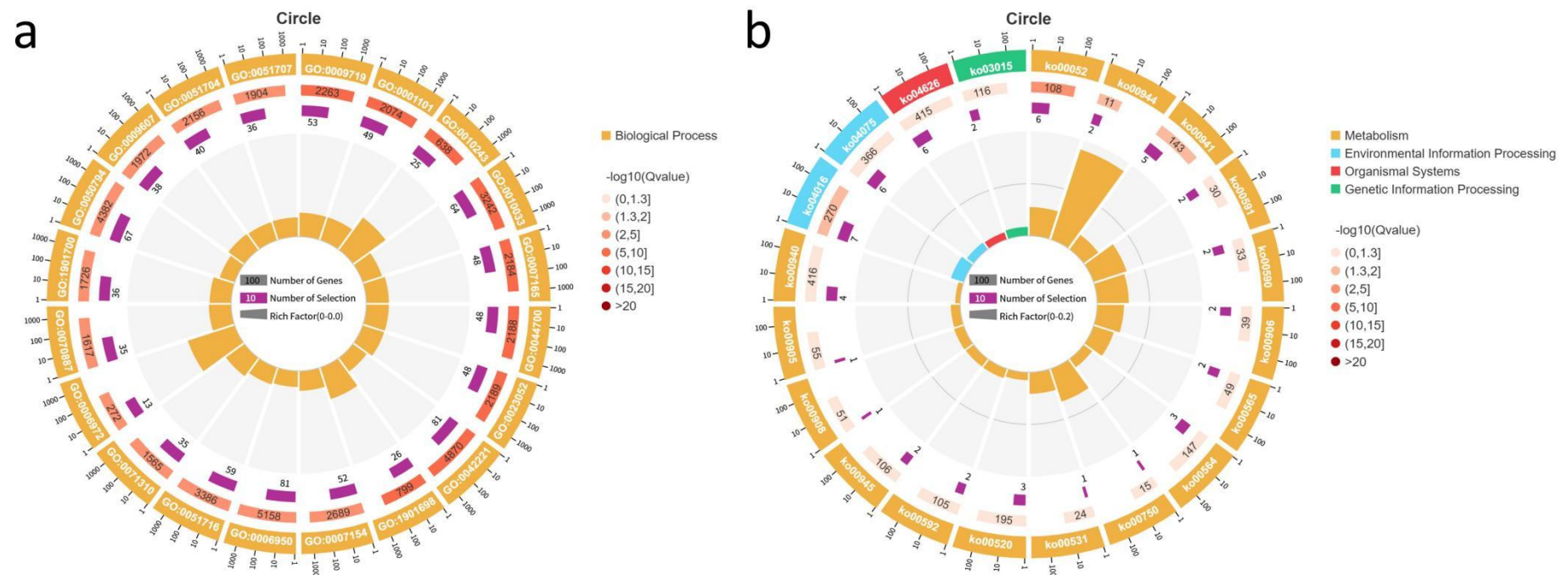


Figure S17 Enrichment circle diagram - WGCNA (Note: Tan module; a: GO; b: KEGG)

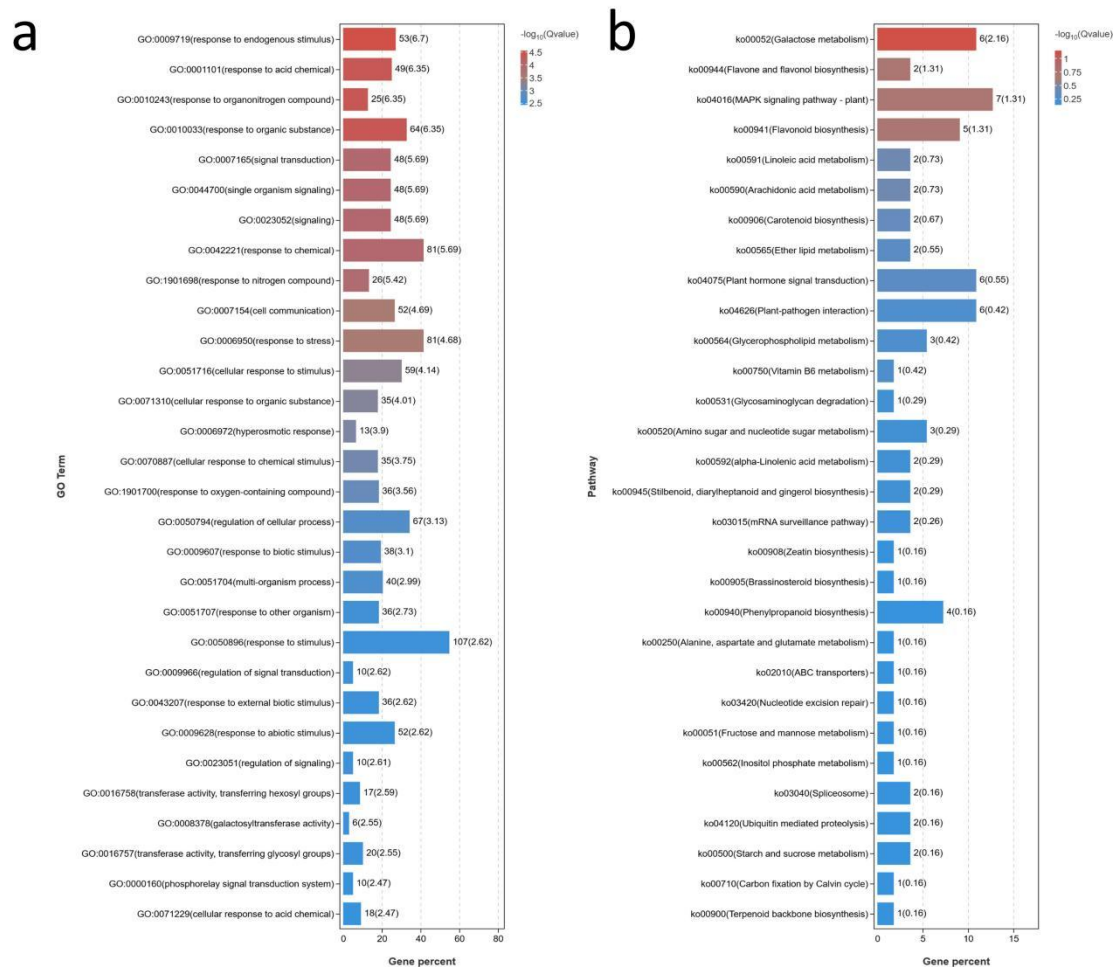


Figure S18 Enrichment bar diagram - WGCNA (Note: Tan module; a: GO; b: KEGG)

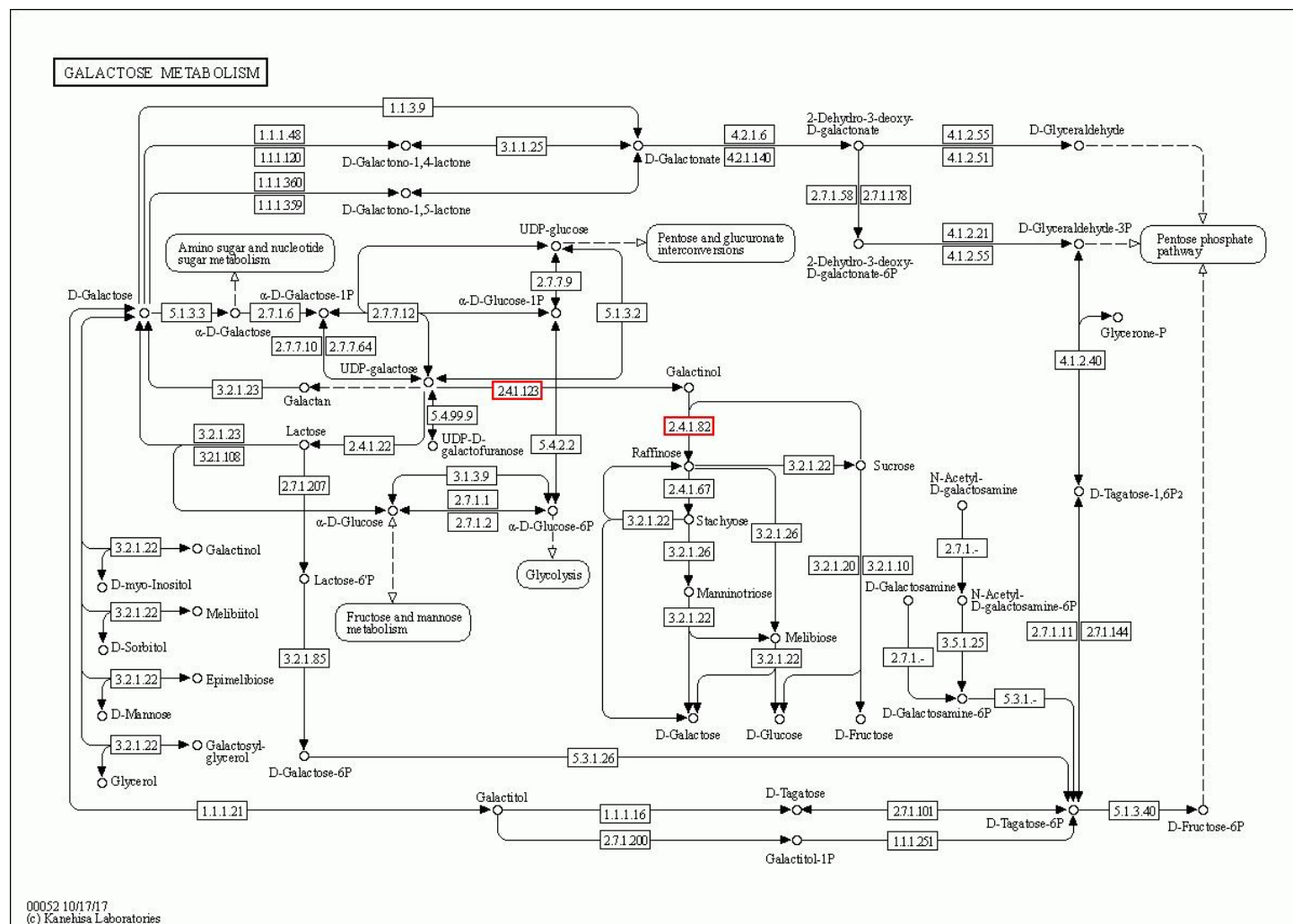
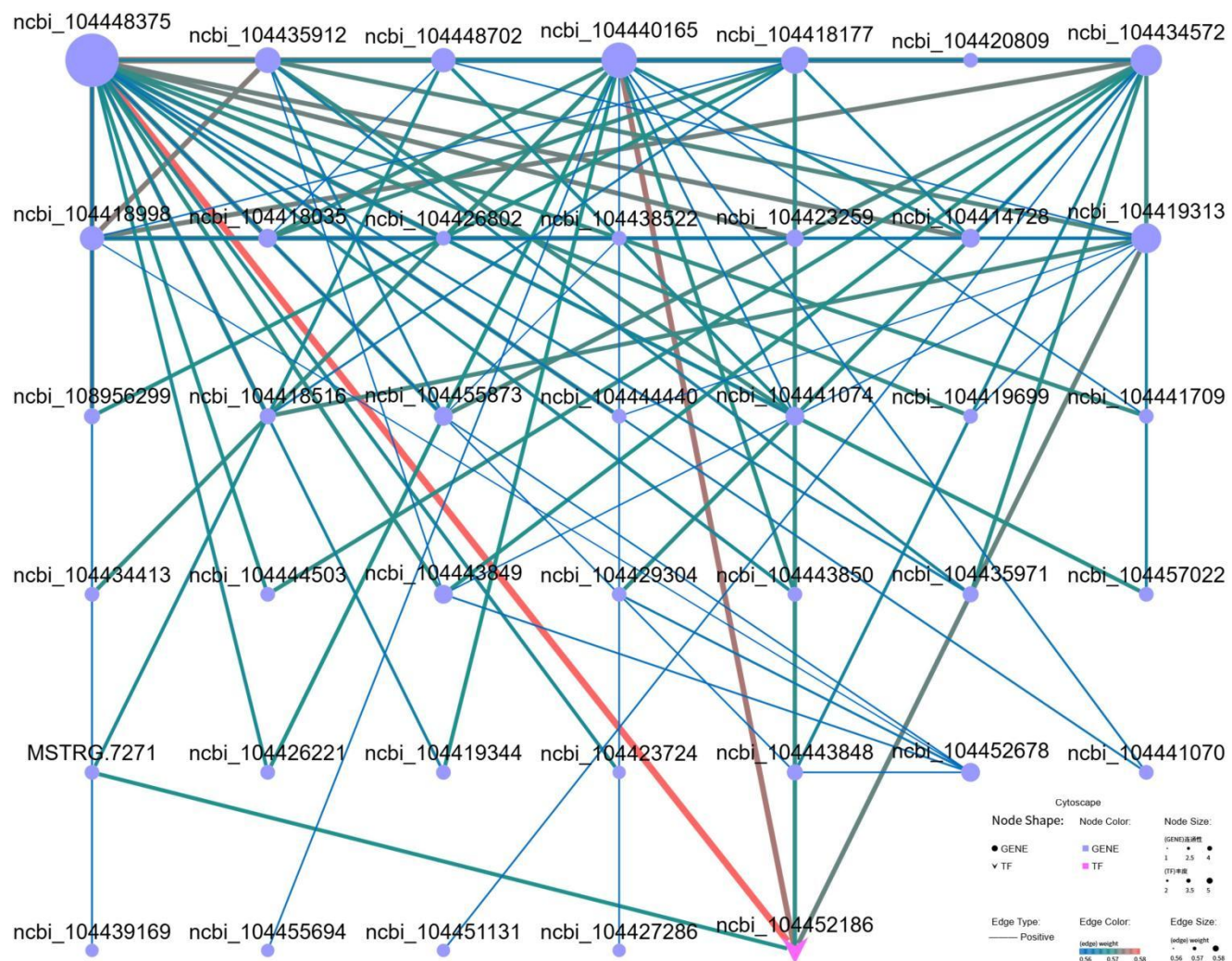


Figure S19 Detailed information on Pathways in KEGG database - WGCNA (Note: Tan module Galactose metabolism(KO00052))



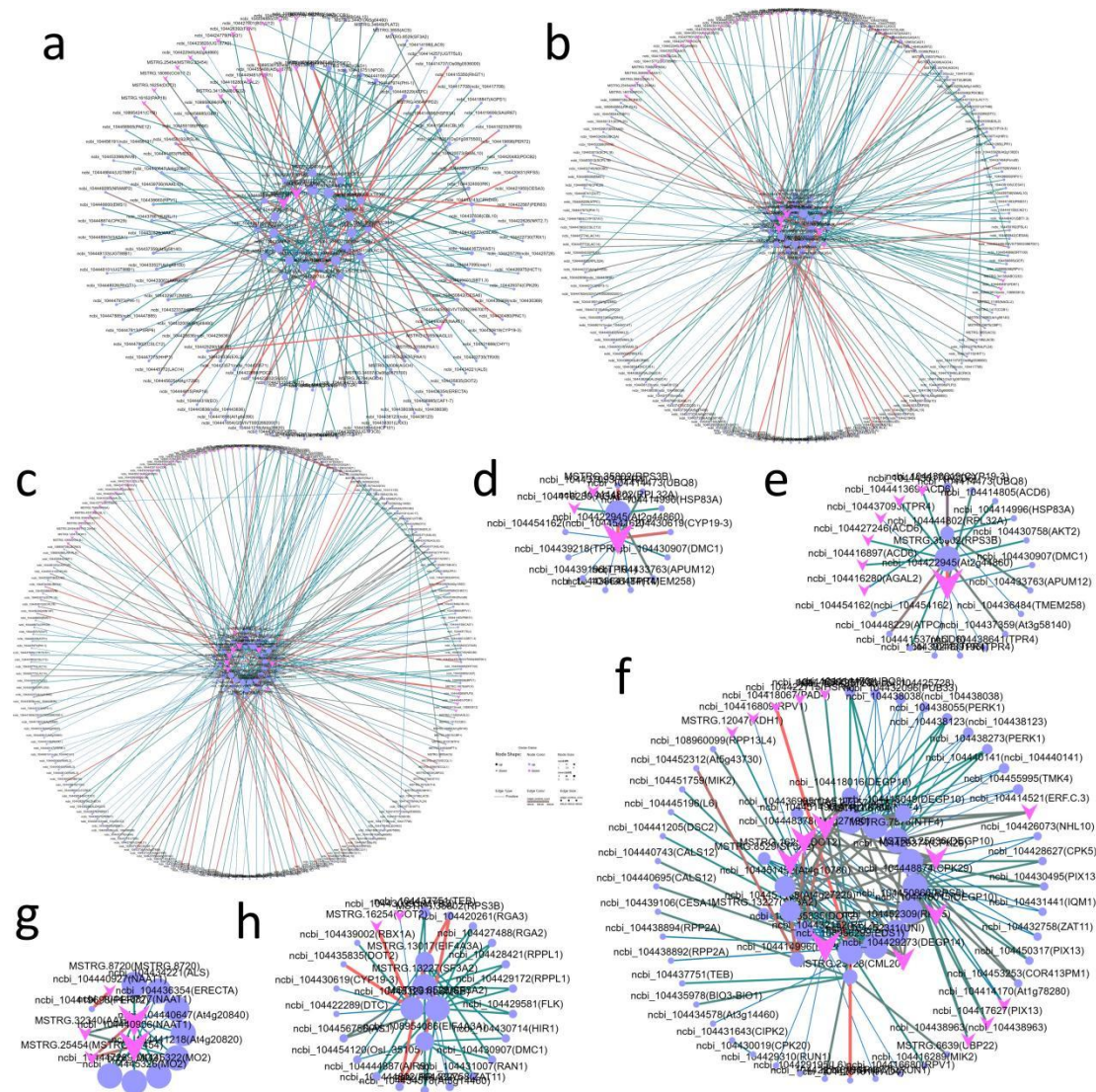


Figure S21 Preliminary target gene PPI network (Note: concentric diagram. a: extracellular region (GO:0005576); b: external encapsulating structure (GO:0030312); c: cell periphery (GO:0071944); d: large ribosomal subunit (GO:0015934); e: ribosomal subunit (GO:0044391); f: Plant-pathogen interaction (KO04626); g: Phenylalanine, tyrosine and tryptophan biosynthesis (KO00400); h: Spliceosome (KO03040).)

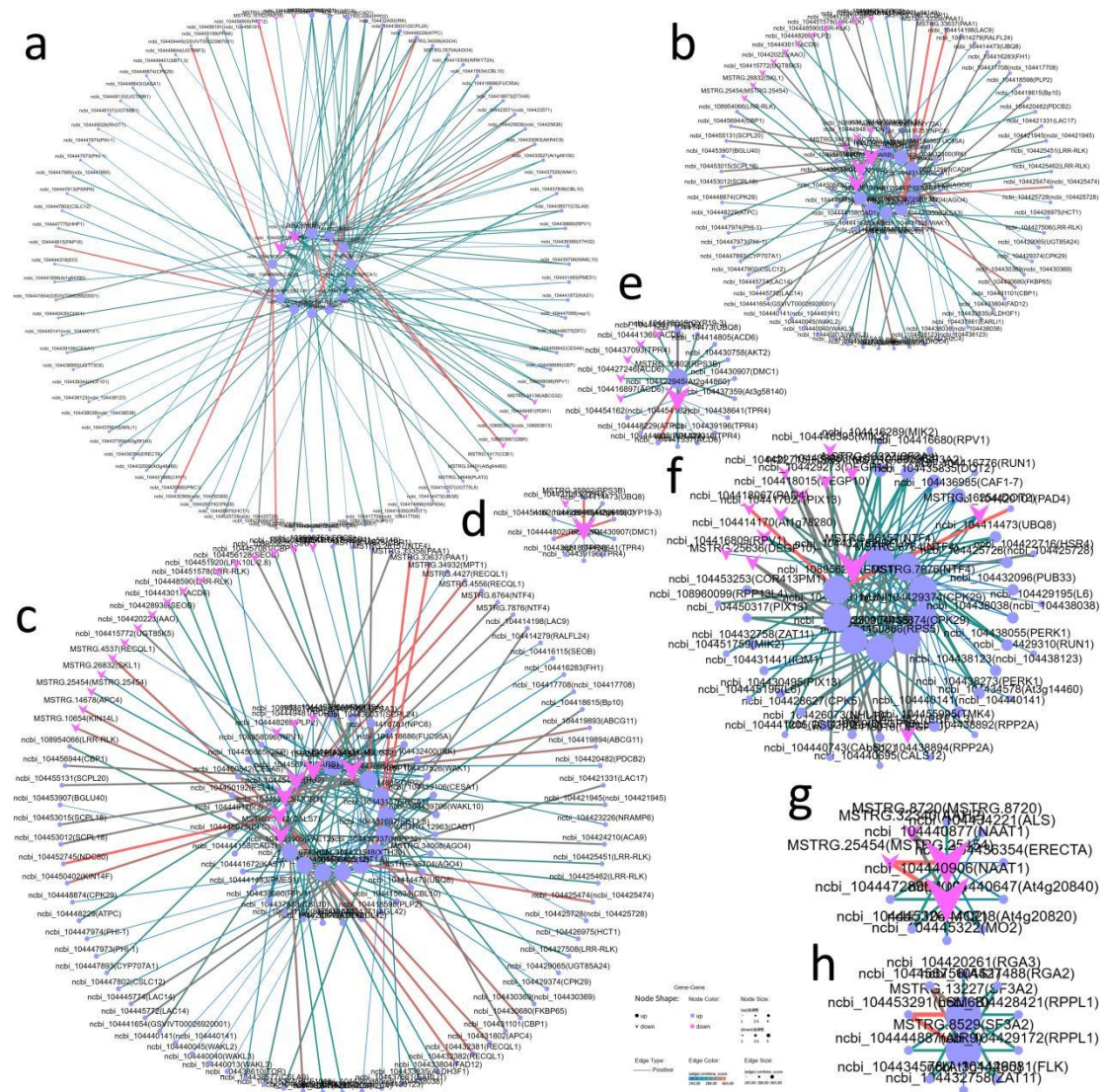


Figure S22 Target gene PPI network (Note: concentric diagram. a: extracellular region (GO:0005576); b: external encapsulating structure (GO:0030312); c: cell periphery (GO:0071944); d: large ribosomal subunit (GO:0015934); e: ribosomal subunit (GO:0044391); f: Plant-pathogen interaction (KO04626); g: Phenylalanine, tyrosine and tryptophan biosynthesis (KO00400); h: Spliceosome (KO03040).)