

GENE-SWitCH

The regulatory GENome of SWine and CHicken: functional annotation during development

Protocol WP2

Chicken gene annotation with Isoseq sequencing using isoseq nextflow pipeline and downstream analysis.

Authors: Sébastien Guizard (UEDIN)

Workpackage: WP2

Version: 1.0

Protocol associated with Deliverable(s):	D2.1 D2.2
Submission date to FAANG:	23/03/2022

Research and Innovation Action, SFS-30-2018-2019-2020 Agri-Aqua Labs
Duration of the project: 01 July 2019 – 30 June 2023, 48 months



Table of contents

1	Summary.....	4
2	Protocol description	5
2.1	Genome annotation	5
2.2	Read count computing	7
2.3	Filtering low confidence annotations	10
2.4	Remove fragment annotations.....	11
2.5	Filter internal priming artefacts.....	11
2.6	Detect and remove annotations produce by NMD and RT-switching	12
2.7	Merge annotation to redefine genes	13
2.8	Updating read count.....	13
2.9	Long Non Coding RNA detection	14
2.10	Merge FEELnc biotype to read support file	15
2.11	Add gene and transcripts sources to read support file	16
2.12	Add block information from bed file	17
2.13	Discard Ensembl annotations	18



Table of figures

Figure 1: ISOseq nexflow pipeline⁵

Figure 2: Read counting process⁸



1 Summary

GENE-SWitCH aims at identifying functional elements located in the genomes of the pig and chicken working on seven different tissues at three different developmental stages.

The seven tissues analyzed in GENE-SWitCH are:

- Cerebellum
- Lung
- Kidney
- Dorsal skin
- Small intestine
- Liver
- Skeletal muscle

The three developmental stages are:

- Early organogenesis (E8 chick embryo and D30 pig fetuses)
- Late organogenesis (E15 chick embryo and D70 pig fetuses)
- Newborn piglets and hatched chicks

For each tissue at each time point, an Isoseq long read sequencing has been done. The raw subreads needs to be processed to generates definitive consensus sequences. The reads are mapped on the reference genome with uLTRA. The gene models are cleaned with TAMA. The resulting annotation files have to be processed to convert to the GFF format and add information (read count, annotation confidence).

2 Protocol description

The following protocol has been applied on either chicken and pig isoseq datasets.

2.1 Genome annotation

The isoseq raw subreads processing, the genome mapping and the gene model annotation creation are made using the [isoseq nextflow pipeline](#).

First the pipeline generates Circular Consensus Sequences (CCS) using ccs Pacbio's program. The raw data are divided into batches of sequences that are processed in parallel. On each ccs batch, the program LIMA (from Pacbio) select CCS with appropriate primers pairs and removes them from the sequence. The resulting sequences are then processed by Pacbio's isoseq3 refine. It selects non-chimeric sequences with poly(A) tail. The sequences in BAM format are converted into FASTA format using bamtools convert. The program polyacleanup from TAMA is then apply to remove remaining polyA tails. At the end of the cleaning process, sequences are called Full Length Non Chimeric (FLNC).

Each batch of FLNC is aligned on the reference genome with the program uLTRA.

Each alignment is processed by TAMA collapse, for false positive and redundancy removing. A bed annotation file is generated for each batch of sequences. Those batches are regrouped based on their sample of origin and merge into one annotation by TAMA merge.

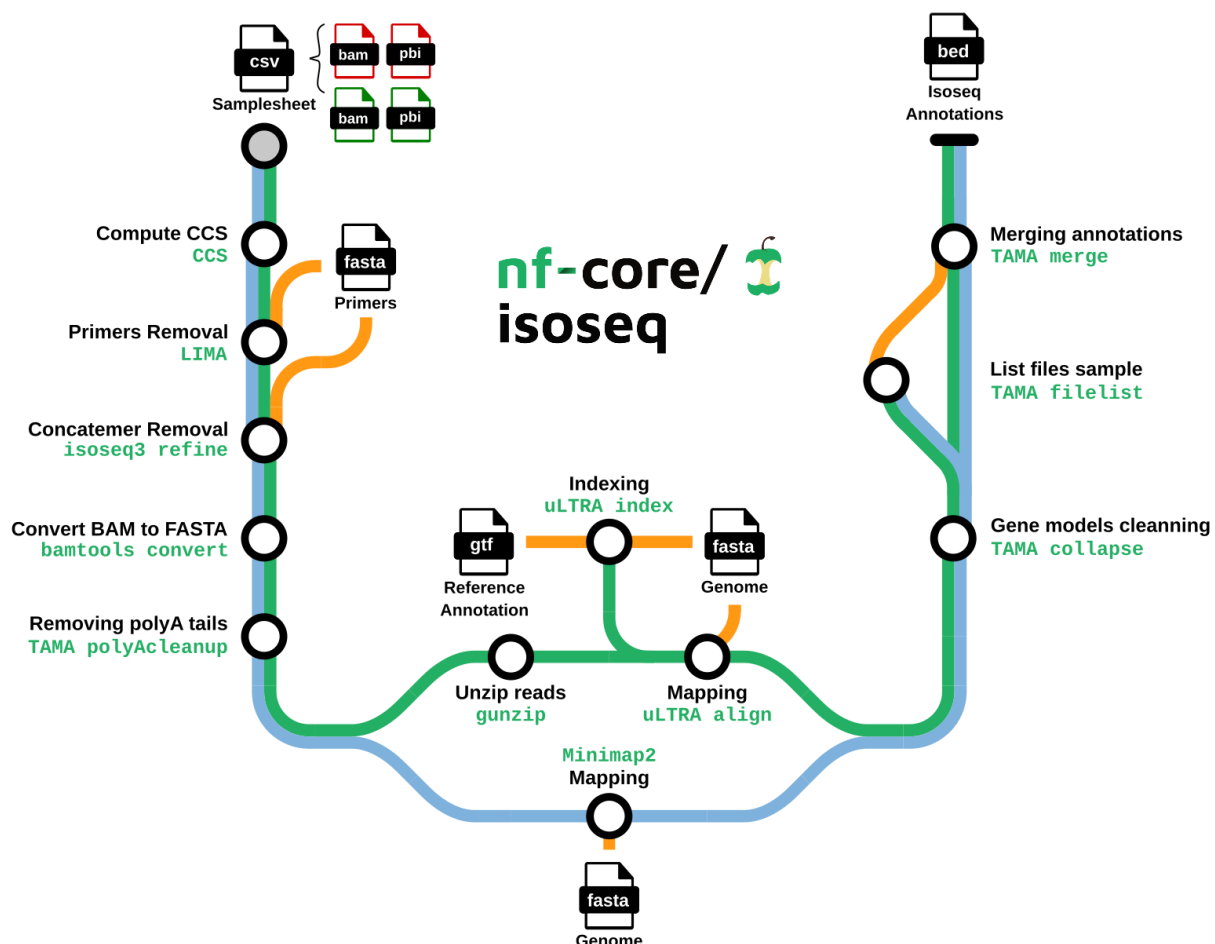


Figure 1: ISOseq nexflow pipeline



To run the pipeline, a data directory is created and the following input files are stored in it:

- Samples subreads (.bam)
- Pacbio index files (.bam.pbi)
- Reference genome (.fasta)
- Primer sequences (.fasta)
- Ensembl gene annotation (.gtf)

The genome used is Gallus Gallus 6 (GRCg6a) and its gene annotation from Ensembl release-105. Contrary to PacBio's instructions, the primer sequence remains unmodified and the polyA sequence is conserved:

```
>5p
TGGATTGATATGTAATACGACTCACTATAG
>3p
AAAAAAAAAAAAAAAAAACGCCTGAGA
```

The following command line has been used to run the pipeline on each samples:

```
nextflow run sguizard/isoseq \
-r dev-rkv-eddie \
--input ../data/ \
--primers ../data/primers_complete.fasta \
--fasta ../data/Gallus_gallus.GRCg6a.dna.toplevel.fa \
--chunk 300 \
--rq 0.90 \
--min_passes 2 \
--five_prime 2000 \
--splice_junction 10 \
--three_prime 2000 \
--capped \
--ultra \
--gtf ../data/Gallus_gallus.GRCg6a.105.gtf \
-profile singularity,eddie
```

The files 30 bed files obtained for the 21 pairs of tissues/development stages have been merged together using TAMA merge.

```
tama_merge.py -f inputChicken.tsv -d merge_dup -a 2000 -m 10 -z 2000 -p chicken 1>
merge.log 2> merge.err
```



The “inputChicken.tsv” file is a four-column tabulation separated file. It contains the list of annotation to merge (column 1), the indication if the sequenced RNAs were capped (column 2), the priority annotation merging (column 3), and an id (column 4).

E15_Skin_P2_cleaned.bed	capped	1,1,1	E15_Skin_P2
HC_Lung_cleaned.bed	capped	1,1,1	HC_Lung
E8_Brain_cleaned.bed	capped	1,1,1	E8_Brain
HC_Skin_P1_cleaned.bed	capped	1,1,1	HC_Skin_P1
E15_Skin_P1_cleaned.bed	capped	1,1,1	E15_Skin_P1
E15_Muscle_P1_cleaned.bed	capped	1,1,1	E15_Muscle_P1
HC_Liver_P2_cleaned.bed	capped	1,1,1	HC_Liver_P2
HC_Brain_P1_cleaned.bed	capped	1,1,1	HC_Brain_P1
E15_Liver_cleaned.bed	capped	1,1,1	E15_Liver
HC_Muscle_cleaned.bed	capped	1,1,1	HC_Muscle
E8_Liver_cleaned.bed	capped	1,1,1	E8_Liver
E8_Lung_cleaned.bed	capped	1,1,1	E8_Lung
HC_Skin_P2_cleaned.bed	capped	1,1,1	HC_Skin_P2
E8_Muscle_cleaned.bed	capped	1,1,1	E8_Muscle
HC_Liver_P1_cleaned.bed	capped	1,1,1	HC_Liver_P1
E15_Brain_P1_cleaned.bed	capped	1,1,1	E15_Brain_P1
E8_Ileum_cleaned.bed	capped	1,1,1	E8_Ileum
E15_Kidney_P1_cleaned.bed	capped	1,1,1	E15_Kidney_P1
E8_Skin_cleaned.bed	capped	1,1,1	E8_Skin
E15_Lung_P1_cleaned.bed	capped	1,1,1	E15_Lung_P1
E15_Lung_P2_cleaned.bed	capped	1,1,1	E15_Lung_P2
HC_Ileum_P1_cleaned.bed	capped	1,1,1	HC_Ileum_P1
HC_Brain_P2_cleaned.bed	capped	1,1,1	HC_Brain_P2
E15_Kidney_P2_cleaned.bed	capped	1,1,1	E15_Kidney_P2
E15_Muscle_P2_cleaned.bed	capped	1,1,1	E15_Muscle_P2
E15_Ileum_cleaned.bed	capped	1,1,1	E15_Ileum
HC_Kidney_cleaned.bed	capped	1,1,1	HC_Kidney
HC_Ileum_P2_cleaned.bed	capped	1,1,1	HC_Ileum_P2
E8_Kidney_cleaned.bed	capped	1,1,1	E8_Kidney
E15_Brain_P2_cleaned.bed	capped	1,1,1	E15_Brain_P2

The combined bed file has been converted to GTF format is using TAMA script `tama_convert_bed_gtf_ensembl_no_cds.py`.

```
tama_convert_bed_gtf_ensembl_no_cds chicken.bed chicken.gtf
```

The GTF file has been modified to add:

- The number of reads supporting each transcript annotation or ‘read count’
- The results of lncRNA detection with feelnc
- Ensembl ids for already annotated genes and transcripts
- Annotations sample of origin

2.2 Read count computing

TAMA include a python script dedicated to read count computing. By running it after each TAMA collapse or TAMA merge, it's possible to keep track of read count through the cleaning/merging process. In this case, `tama_read_support_levels.py` has been run after pipeline's TAMA collapse (one bed file per read batch), pipeline's TAMA merge (one bed file per sample), and the last TAMA merge (one bed file for all samples).

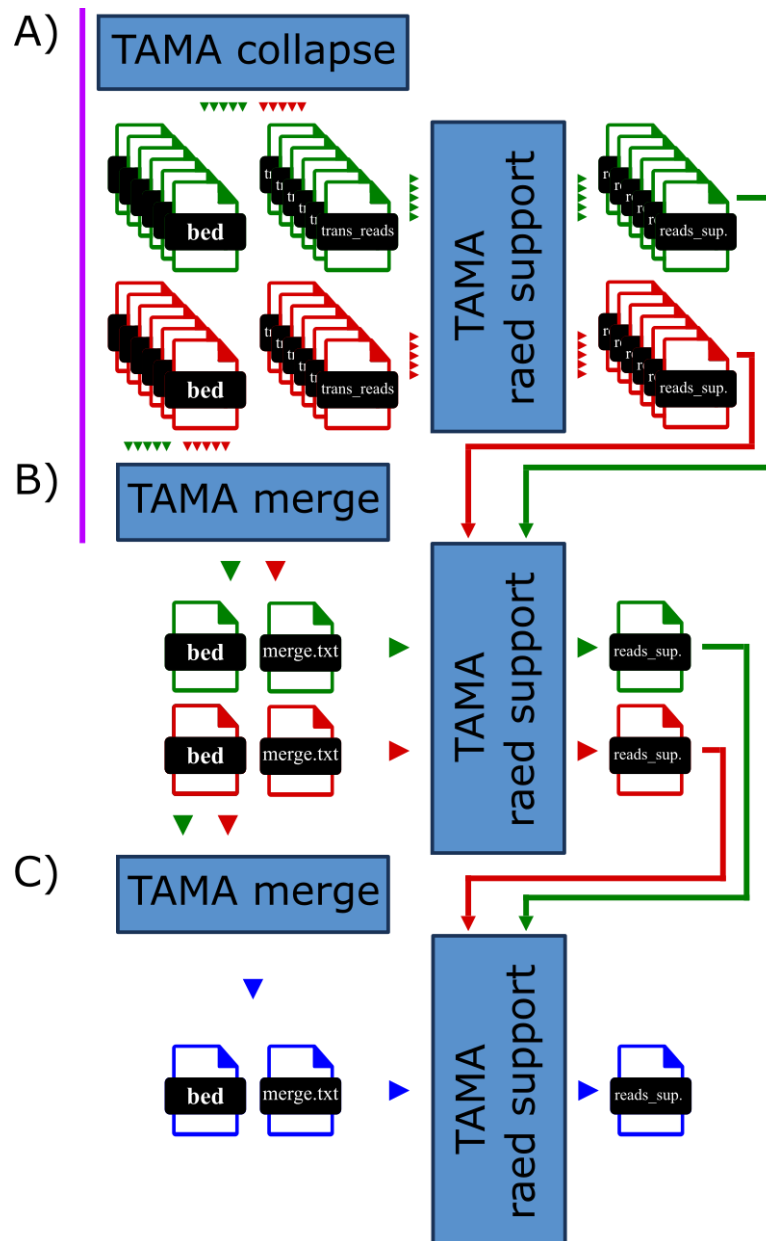


Figure 2: Read counting process – Purple bar indicates programs launch by the pipeline; others are launched manually. A) TAMA read support is launched on each `trans_reads.bed` files produced by TAMA collapse. B) TAMA read support is run on each sample `merge.txt` their associated chromosome `read_support` file. C) TAMA read support is run on the `merge.txt` of all samples with their associated sample `read_support` files



The first counting needs to be run on each `trans_read.bed` files produced by TAMA collapse (Fig 2. A). For each of them, a `read_support` file is created. If all `trans_read.bed` files are gathered in the same directory, the following fish shell loop will run the script on each of them:

```
for FILE in (find 09_GSTAMA_COLLAPSE/ -name '*_trans_read.bed')
  set FILE_SHORT (string replace -r '^./' '' $FILE)
  set ID (string replace "_trans_read.bed" "" $FILE_SHORT)
  set OUTFILE $ID"_filelist.txt"
  echo -e "$ID\t$FILE\ttrans_read" > $OUTFILE
  set CMD "tama_read_support_levels.py -f $OUTFILE -m no_merge -o $ID"
  echo $CMD
  eval $CMD
end
```

At each loop iteration, it will create a `filelist.tsv` input file composed of three tabulations separated columns (an ID, the `trans_read.bed` file and the file type) and run `tama_read_support_levels.py`.

Next, the generated `read_support` files are used to compute the read count at sample level (first TAMA merge). For each sample, a `filelist.tsv` file listing `read_support` files must be created. The first column is the ID, the second is the file path and the third one contains `read_support` as file type. The script can be started with following command:

```
for ID in (ls /*_read_support.txt|perl -pe 's/./\//; s/\..+txt//g'|sort|uniq)
  set INPUT $ID"_input.tsv"

  for RS in (ls ../01-read_support_collapse/$ID*_read_support.txt)
    set ID2 (echo $RS|perl -pe 's/\.\.\./01-read_support_collapse\//; s/_read_support.txt/_collapsed.bed/g')
    echo -e "$ID2\t$RS\tread_support" >> $INPUT
  end

  set MERGE $ID"_merge.txt"
  set CMD "ln -sf ../05-merging_files/02-merge.txt_by_samples/"$MERGE
  eval $CMD
  set CMD "tama_read_support_levels.py -f $INPUT -m $MERGE -o $ID"
  echo $CMD
  eval $CMD
end
```

Finally, the global read count can be computed using the `read_support` files generated for each sample. A fish shell loop can be used to generate the file list file. The `tama_read_support_levels.py` script is run using the resulting `filelist` file and the `merge.txt` generated during TAMA merge process.

```
for RS in (ls ../02-read_support_merge1/*_read_support.txt)
  set ID (string replace -r '_read_support.txt' '' $RS|string replace -r '\.\.\./02-read_support_merge1/' '')
  echo -e "$ID\t$RS\tread_support" >> chicken_input.tsv
end

ln -s ../04-merge_annotations/chicken_merge.txt
tama_read_support_levels.py -f chicken_input.tsv -m chicken_merge.txt -o chicken
```



2.3 Filtering low confidence annotations

To reduce false positive results, annotations supported by only 1 or 2 reads have been discarded. This is done by reading data from read_support file generated by `tama_read_support_levels.py` and listing high confidence annotation's ids.

```
cd ~/gp_vol/Projects/iseqseq_rkv_v2/chicken_mat/
mkdir 07-filter_under_3_reads
cd 07-filter_under_3_reads
ln -s ../06-read_count/03-read_support_merge2/chicken_mat_read_support.txt
ln -s ../04-merge_annotations/chicken_mat.bed
library(tidyverse)

d <-
  read_tsv(
    "chicken_read_support.txt",
    col_types = c(
      "merge_gene_id" = "c",
      "merge_trans_id" = "c",
      "gene_read_count" = "i",
      "trans_read_count" = "i",
      "source_line" = "c",
      "support_line" = "c"))

d %>%
  filter(trans_read_count >= 3) %>%
  write_tsv("chicken_read_support_hc.txt")

d %>%
  filter(trans_read_count < 3) %>%
  write_tsv("chicken_read_support_lc.txt")

hc_id <-
  d %>%
  filter(trans_read_count >= 3) %>%
  select(merge_trans_id)

lc_id <-
  d %>%
  filter(trans_read_count < 3) %>%
  select(merge_trans_id)

source("~/read_bed12.R")

bed <- read_bed12("chicken.bed")

bed %>%
  mutate(merge_trans_id = str_match(name, "G\\d+.(G\\d+\\.\\d+)")[,2]) %>%
  semi_join(hc_id) %>%
  select(-merge_trans_id) %>%
  write_tsv("chicken_min3reads.bed", col_names = FALSE)

bed %>%
  mutate(merge_trans_id = str_match(name, "G\\d+.(G\\d+\\.\\d+)")[,2]) %>%
  semi_join(lc_id) %>%
  select(-merge_trans_id) %>%
  write_tsv("chicken_under3reads.bed", col_names = FALSE)
```



2.4 Remove fragment annotations

At this point, the annotation still contain annotation which are fragments to larger to one. The `tama_remove_fragment_models.py` is used to removes those annotation.

```
tama_remove_fragment_models.py \
-f chicken_min3reads.bed \
-o chicken_min3reads_noFrag
```

2.5 Filter internal priming artefacts

The primer used for sequencing include a polyA stretch intended to bind to the polyA tail. However, if the mRNA contains internal polyA stretch, the primer can miss its target and generates artefacts. These artefacts transcripts can be detected by searching for polyA sequence following the last exon.

The annotations followed by A enriched sequence (> 75% in the 18 bp) are discarded.

First, the position of the 18bp sequences is computed.

```
library(tidyverse)

d <-
  read_tsv(
    file = "chicken_min3reads_noFrag.bed",
    col_names = c(
      "chrom",
      "chromStart",
      "chromEnd",
      "name",
      "score",
      "strand",
      "thickStart",
      "thickEnd",
      "itemRgb",
      "blockCount",
      "blockSizes",
      "blockStarts"),
    col_types = cols(
      chrom = "c",
      chromStart = "i",
      chromEnd = "i",
      name = "c",
      score = "i",
      strand = "c",
      thickStart = "i",
      thickEnd = "i",
      itemRgb = "c",
      blockCount = "i",
      blockSizes = "c",
      blockStarts = "c"))

d %>%
  mutate(
    start = ifelse(strand == "+", chromEnd, chromStart - 18),
    end = ifelse(strand == "+", chromEnd + 18, chromStart)) %>%
  mutate(start = ifelse(start < 0, 0, start)) %>%
  select(chrom, start, end, name, score, strand) %>%
  write_tsv("18bpExtension.bed", col_names = FALSE)
```



The 18bp sequences are extracted with bedtools. The resulting FASTA file is formatted for being easily loaded in R.

```
bedtools getfasta \
-fi Gallus_gallus.GRCg6a.dna.toplevel.fa \
-fo 18bpExtension.fa \
-bed 18bpExtension.bed \
-name -s

perl -pe 's/\n/\t/ if (/^>/)' 18bpExtension.fa > 18bpExtension.fa.tsv
```

The artifacts are then discarded.

```
pa <-
  read_tsv(
    '18bpExtension.fa.tsv',
    col_names = c('seqname_ext', 'seq_ext')) %>%
  mutate(
    name = str_extract(seqname_ext, '>G[^:]+') %>% str_replace('>', ''),
    A = str_count(seq_ext, 'A')/nchar(seq_ext)*100,
    T = str_count(seq_ext, 'T')/nchar(seq_ext)*100,
    G = str_count(seq_ext, 'G')/nchar(seq_ext)*100,
    C = str_count(seq_ext, 'C')/nchar(seq_ext)*100) %>%
  write_tsv("18bpExtension_pATGC.tsv")

d %>%
  left_join(pa) %>%
  write_tsv("Chicken_min3reads_noFrag_noInternPrim.bedCustom")

d %>%
  left_join(pa) %>%
  filter(A < 75) %>%
  select(1:12) %>%
  write_tsv(
    "chicken_min3reads_noFrag_noInternPrim.bed12",
    col_names = FALSE)
```

2.6 Detect and remove annotations produce by NMD and RT-switching

The SQANTI3 program can be used to classify and transcripts and evaluate if they are the product of mRNA Non Mediated Decay of RT-switching. First, the bed file must be converted into GTF format before being analyzed by SQANTI3.

```
tama_convert_bed_gtf_ensembl_no_cds.py \
chicken_min3reads_noFrag_noInternPrim.bed12 \
chicken_min3reads_noFrag_noInternPrim.gtf

python sqanti3_qc.py \
  chicken_min3reads_noFrag_noInternPrim.gtf \
  Gallus_gallus_gca000002315v5.GRCg6a.108.gtf \
  Gallus_gallus.GRCg6a.dna.toplevel.fa \
  -d SQANTI_chicken \
  -o SQANTI_chicken
```

To identify and discard NMD and RT-switching transcripts, the SQANTI3's file.

```
library(tidyverse)
```



```
source('~/.read_bed12.R')

d <- read_tsv("SQANTI_chicken_classification.txt")
bed <-
  read_bed12('~/chicken_min3reads_noFrag_noInternPrim.bed12') %>%
  mutate(isoform = str_match(name, 'G\\d+;(G\\d+\\.\\.\\d+)')[,2])

valid <-
  d %>%
  filter(RTS_stage == FALSE) %>%
  filter(predicted_NMD == FALSE | is.na(predicted_NMD)) %>%
  select(isoform)

bed %>%
  semi_join(valid) %>%
  select(-isoform) %>%
  write_tsv("../chicken_min3reads_noFrag_noInternPrim_noNMD_noRTS.bed12",
  col_names = FALSE)
```

2.7 Merge annotation to redefine genes

The extraction of low confidence transcripts reduces information and can influence the how transcripts can be merge. That why another merging steps is done to redefine genes. To add gene and transcripts ids from Ensembl, the ensemble annotation file is added to this merging step.

```
echo -e "chicken_clean.bed\tcapped\t1,1,1\tISOseq" >> inputChicken.tsv
echo -e "Gallus_gallus_gca000002315v5.GRCg6a.108.bed\tcapped\t1,1,1\tEnsembl" >>
inputChicken.tsv
tama_format_gtf_to_bed12_ensembl.py \
Gallus_gallus_gca000002315v5.GRCg6a.108.gtf \
Gallus_gallus_gca000002315v5.GRCg6a.108.bed > format.log > format.err

tama_merge.py \
-f inputChicken.tsv \
-d merge_dup \
-a 2000 \
-m 10 \
-z 2000 \
-p chicken_ensembl \
-cds Ensembl \
-s Ensembl 1> merge.log 2> merge.err
```

2.8 Updating read count

The last merging process updates the genes and transcripts id. Then, it's mandatory to update read count file with these new ids. Also, due to the discarded transcripts the count must be updated either.

```
library(tidyverse)

read_bed12 <- function(file, skip = 0) {require(readr); read_tsv(file, col_names =
c("chrom", "chromStart", "chromEnd", "name", "score", "strand", "thickStart",
"thickEnd", "itemRgb", "blockCount", "blockSizes", "blockStarts"), col_types =
cols(chrom = col_character(), chromStart = col_integer(), chromEnd = col_integer(),
name = col_character(), score = col_double(), strand = col_character(), thickStart
= col_integer(), thickEnd = col_integer(), itemRgb = col_character(), blockCount =
col_integer(), blockSizes = col_character(), blockStarts = col_character()), skip =
skip)}
```



```

m <-
  read_bed12(file = "chicken_ensembl_merge.txt") %>%
  transmute(
    merge_gene_id = str_match(name, "(G\\d+)\\.\\.\\d+;(Ensembl|ISOseq)_\\.+")[,2],
    merge_trans_id = str_match(name, "(G\\d+)\\.\\.\\d+;(Ensembl|ISOseq)_\\.+")[,2],
    source = str_match(name, "G\\d+\\.\\.\\d+;(Ensembl|ISOseq)_\\.+")[,2],
    previous_trans_id = str_match(name,
"G\\d+\\.\\.\\d+;(Ensembl|ISOseq)_\\.+")[,3]) %>%
  filter(source == "ISOseq") %>%
  select(-source)

read_rs <- function(file) {
  read_tsv(
    file = file,
    col_types = c(
      "merge_gene_id" = "c",
      "merge_trans_id" = "c",
      "gene_read_count" = "i",
      "trans_read_count" = "i",
      "source_line" = "c",
      "support_line" = "c"))
}

rs <-
  read_rs(file = "chicken_read_support.txt") %>%
  select(-merge_gene_id, -gene_read_count) %>%
  rename(previous_trans_id = merge_trans_id)

grc <-
  m %>%
  left_join(rs) %>%
  mutate(merge_gene_id = str_match(merge_trans_id, "^(G\\d+)\\.\\.")[,2]) %>%
  group_by(merge_gene_id) %>%
  summarise(gene_read_count = sum(trans_read_count))

m_rs_grc <-
  m %>%
  left_join(rs) %>%
  left_join(grc) %>%
  write_tsv("chicken_ensembl_read_support.txt")

```

2.9 Long Non Coding RNA detection

The program *feelnc* perform Long Non-Coding RNA (lncRNA) detection with an alignment-free method. By analyzing nucleotide compositional bias of coding and non-coding regions of the genome, the program can infer the transcript coding status.

A nextflow pipeline ([sguizard/nf-feelnc](#)) has been developed, based on [FAANG/analysis-TAGADA](#) pipeline, to run all *feelnc* analyzing steps. The program has been applied on the final annotation, with same genome and reference annotation used for isoseq annotations.

```

nextflow run sguizard/nf-feelnc \
-r dev \
--ref_annotation Gallus_gallus_gca000002315v5.GRCg6a.108.gtf\
--new_annotation chicken_ensembl.gtf \
--genome Gallus_gallus.GRCg6a.dna.toplevel.fa \
--profile singularity,eddie

```



2.10 Merge FEELnc biotype to read support file

The FEELnc pipeline produces a GTF annotation with an additional field: `feelnc_biotype`. This information is added to the read support file.

```
library(tidyverse)

read_gtf <- function(file) {
  require(readr)
  read_tsv(
    file,
    col_names = c(
      "sequence", "source", "feature",
      "start",    "end",    "score",
      "strand",   "phase",  "attributes"),
    col_types = cols(
      sequence   = "c",
      source     = "c",
      feature    = "c",
      start      = "i",
      end        = "i",
      score      = "c",
      strand     = "c",
      phase      = "c",
      attributes = "c"),
    comment = "#") %>%
  mutate(
    gene_id          = str_match(attributes, 'gene_id "([^"]+)"')[,2],
    gene_version     = str_match(attributes, 'gene_version "([^"]+)"')[,2],
    transcript_id     = str_match(attributes, 'transcript_id "([^"]+)"')[,2],
    transcript_version = str_match(attributes, 'transcript_version "([^"]+)"')[,2],
    uniq_trans_id    = str_match(attributes, 'uniq_trans_id "([^"]+)"')[,2],
    gene_source       = str_match(attributes, 'gene_source "([^"]+)"')[,2],
    gene_biotype      = str_match(attributes, 'gene_biotype "([^"]+)"')[,2],
    feelnc_biotype    = str_match(attributes, 'feelnc_biotype "([^"]+)"')[,2],
    transcript_source  = str_match(attributes, 'transcript_source "([^"]+)"')[,2],
    transcript_biotype = str_match(attributes, 'transcript_biotype "([^"]+)"')[,2],
    exon_number       = str_match(attributes, 'exon_number "([^"]+)"')[,2],
    exon_id           = str_match(attributes, 'exon_id "([^"]+)"')[,2],
    exon_version      = str_match(attributes, 'exon_version "([^"]+)"')[,2])
}

gtf <-
  read_gtf("novel.gtf") %>%
  select(where(~!all(is.na(.x)))) %>%
  filter(feature == "transcript" & !is.na(feelnc_biotype)) %>%
  select(uniq_trans_id, feelnc_biotype) %>%
  rename(merge_trans_id = uniq_trans_id)

rs <-
  read_tsv(
    "chicken_ensembl_read_support.txt",
    col_types = c(
      "merge_gene_id" = "c",
      "merge_trans_id" = "c",
```



```

        "previous_trans_id" = "c",
        "trans_read_count"  = "i",
        "source_line"      = "c",
        "support_line"     = "c",
        "gene_read_count"  = "i"))

rs %>%
  left_join(gtf) %>%
  select(1:3, 7, 4, 8, 5, 6) %>%
  write_tsv("chicken_ensembl_read_support_and_feelnc.txt")

```

2.11 Add gene and transcripts sources to read support file

TAMA file the origin of an annotation (ISOseq, Ensembl, Ensembl & ISOseq) in two files gene_report and trans_report. This information is added to the read support file.

```

library(tidyverse)

src_g <-
  read_tsv(
    file = "chicken_ensembl_gene_report.txt",
    col_types = c(
      gene_id      = "c",
      num_clusters = "i",
      num_final_trans = "i",
      sources      = "c",
      chrom       = "c",
      start       = "i",
      end         = "i",
      source_genes = "c",
      source_summary = "c")) %>%
  rename(
    merge_gene_id = gene_id,
    sources_id_gene = sources) %>%
  select(merge_gene_id, sources_id_gene, source_genes, source_summary)

src_t <-
  read_tsv(
    file = "chicken_ensembl_trans_report.txt",
    col_types = c(
      transcript_id = "c",
      num_clusters = "i",
      sources      = "c",
      start_wobble_list = "c",
      end_wobble_list = "c",
      exon_start_support = "c",
      exon_end_support = "c",
      all_source_trans = "c")) %>%
  rename(
    merge_trans_id = transcript_id,
    sources_id_trans = sources) %>%
  mutate(merge_gene_id = str_match(merge_trans_id, "(G\\d+)\\.\\.\\.\\d+")[,2]) %>%
  select(merge_gene_id, merge_trans_id, sources_id_trans, all_source_trans)

rs <-
  read_tsv(
    file = "chicken_ensembl_read_support_and_feelnc.txt",
    col_types = c(
      merge_gene_id = "c",

```



```

merge_trans_id    = "c",
previous_trans_id = "c",
gene_read_count   = "i",
trans_read_count  = "i",
feelnc_biotype    = "c",
source_line       = "c"))

left_join(src_g, src_t) %>%
  select(
    merge_gene_id,    merge_trans_id, sources_id_gene,
    sources_id_trans, source_summary, source_genes,
    all_source_trans) %>%
  left_join(rs) %>%
  write_tsv("chicken_ensembl_read_support_and_feelnc_and_sources.txt")

```

2.12 Add block information from bed file

The 12-column version of BED format carry the information about exon number and exon size (blockCount, blockSizes). This information is important is added to the read support file.

```

library(tidyverse)

read_bed12 <- function(file, skip) {
  skip = 0

  require(readr)
  read_tsv(
    file,
    col_names = c(
      "chrom",
      "chromStart",
      "chromEnd",
      "name",
      "score",
      "strand",
      "thickStart",
      "thickEnd",
      "itemRgb",
      "blockCount",
      "blockSizes",
      "blockStarts"),
    col_types = cols(
      chrom = col_character(),
      chromStart = col_integer(),
      chromEnd = col_integer(),
      name = col_character(),
      score = col_double(),
      strand = col_character(),
      thickStart = col_integer(),
      thickEnd = col_integer(),
      itemRgb = col_character(),
      blockCount = col_integer(),
      blockSizes = col_character(),
      blockStarts = col_character()),
    skip = skip
  )
}

bed <-

```



```

read_bed12("chicken_ensembl.bed") %>%
mutate(
  merge_gene_id = str_match(name, "(G\\d+);G\\d+\\.\\.\\d+")[,2],
  merge_trans_id = str_match(name, "G\\d+;(G\\d+\\.\\.\\d+")[,2]) %>%
  select(merge_gene_id, merge_trans_id, blockCount, blockSizes)

d <-
  read_tsv(
    file = "chicken_ensembl_read_support_and_feelnc_and_sources.txt",
    col_types = c(
      merge_gene_id = "c",
      merge_trans_id = "c",
      sources_id_gene = "c",
      sources_id_trans = "c",
      source_summary = "c",
      source_genes = "c",
      all_source_trans = "c",
      previous_trans_id = "c",
      gene_read_count = "i",
      trans_read_count = "i",
      feelnc_biotype = "c",
      source_line = "c",
      support_line = "c"))

d %>%
  left_join(bed) %>%
  write_tsv(file = "chicken_ensembl_read_support_and_feelnc_and_sources_and_blockInfo.txt")

```

2.13 Discard Ensembl annotations

The annotations specific to Ensembl reference are discarded in order to keep ISOseq specific and Ensembl & ISOseq ones.

```

library(tidyverse)

read_bed12 <- function(file, skip = 0) {require(readr); read_tsv(file, col_names =
c("chrom", "chromStart", "chromEnd", "name", "score", "strand", "thickStart",
"thickEnd", "itemRgb", "blockCount", "blockSizes", "blockStarts"), col_types =
cols(chrom = col_character(), chromStart = col_integer(), chromEnd = col_integer(),
name = col_character(), score = col_double(), strand = col_character(), thickStart
= col_integer(), thickEnd = col_integer(), itemRgb = col_character(), blockCount =
col_integer(), blockSizes = col_character(), blockStarts = col_character()), skip =
skip)}

bed <- read_bed12("chicken_ensembl.bed")

d <-
  read_tsv(
    "chicken_ensembl_read_support_and_feelnc_and_sources_and_blockInfo.txt",
    col_types = c(
      merge_gene_id = "c", merge_trans_id = "c", sources_id_gene = "c",
      sources_id_trans = "c", source_summary = "c", source_genes = "c",
      all_source_trans = "c", previous_trans_id = "c", gene_read_count = "i",
      trans_read_count = "i", feelnc_biotype = "c", source_line = "c",
      support_line = "c"))

d %>%
  filter(sources_id_trans != 'Ensembl') %>%

```



```
filter(sources_id_gene != 'Ensembl') %>%

write_tsv("chicken_ensembl_read_support_and_feelnc_and_sources_and_blockInfo_gs.txt")

gs_id <-
  d %>%
    filter(sources_id_trans != 'Ensembl') %>%
    filter(sources_id_gene != 'Ensembl') %>%
    select(merge_gene_id, merge_trans_id)

bed %>%
  mutate(
    merge_gene_id = str_match(name, "(G\\d+)\\.\\.\\d+")[,2],
    merge_trans_id = str_match(name, "(G\\d+\\.\\.\\d+")[,2]) %>%
    semi_join(gs_id) %>%
    select(-merge_gene_id, -merge_trans_id) %>%
    write_tsv("chicken_ensembl_gs.bed", col_names = FALSE)
```