



EnsemblCompara GeneTrees: Analysis of complete, duplication aware phylogenetic trees in vertebrates.

Albert J Vilella, Jessica Severin, Abel Ureta-Vidal, et al.

Genome Res. published online November 24, 2008

Access the most recent version at doi:[10.1101/gr.073585.107](https://doi.org/10.1101/gr.073585.107)

P<P	Published online November 24, 2008 in advance of the print journal.
Accepted Manuscript	Peer-reviewed and accepted for publication but not copyedited or typeset; accepted manuscript is likely to differ from the final, published version.
Open Access	Freely available online through the <i>Genome Research</i> Open Access option.
License	This manuscript is Open Access.
Email Alerting Service	Receive free email alerts when new articles cite this article - sign up in the box at the top right corner of the article or click here .

Comprehensive immune receptor profiling.
Discover the **DriverMap™ AIR Assay** difference.

LEARN
MORE



To subscribe to *Genome Research* go to:
<https://genome.cshlp.org/subscriptions>

Copyright © 2008, Cold Spring Harbor Laboratory Press

EnsemblCompara GeneTrees: complete, duplication aware phylogenetic trees in vertebrates.

Albert J. Vilella¹, Jessica Severin¹⁺, Abel Ureta-Vidal^{1†}, Li Heng², Richard Durbin², Ewan Birney^{1*}

¹ EMBL-EBI, Wellcome Trust Genome Campus, Hinxton, Cambridge, CB10 1SD, United Kingdom

² Wellcome Trust Sanger Institute, Wellcome Trust Genome Campus, Hinxton, Cambridge, CB10 1HH, United Kingdom

* To whom correspondence should be addressed. Email: birney@ebi.ac.uk

⁺ Current address : RIKEN Yokohama Institute, Genomic Sciences Center (GSC), 1-7-22 Suehiro-cho, Tsurumi-ku, Yokohama, 230-0045 Japan

[†] Current address : Eagle Genomics, 19 Forge End, Stapleford, Cambridge CB22 5BN, United Kingdom

Abstract

We have developed a comprehensive gene orientated phylogenetic resource, EnsemblCompara GeneTrees based on a computational pipeline to handle clustering, multiple alignment and tree generation, including the handling of large gene families. We developed two novel non-sequence based metrics of gene tree correctness and benchmarked a number of tree methods. The TreeBeST method from TreeFam shows the best performance in our hands. We also compared this phylogenetic approach to clustering approaches for ortholog prediction, showing a large increase in coverage using the phylogenetic approach. All data is made available in a number of formats and will be kept up to date with the Ensembl project.

Introduction

The use of phylogenetic trees to describe the evolution of biological processes was established in the 1950s (Hennig 1952) and remains a fundamental approach to understand the evolution of individual genes through to complete genomes, for example in the mouse (Mouse Genome Sequencing Consortium, et al, 2002), rat (Gibbs, et al, 2004), chicken (International Chicken Genome Sequencing Consortium, 2004) and monodelphis (Mikkelsen, et al, 2007) genome papers, and numerous papers on individual sequences. Now routine, the determination of vertebrate genome sequences provides a rich data source to understand evolution, and using phylogenetic trees of the genes is one of the best ways to organise this data. However, the increased set of genomes makes the compute and engineering tasks to form all the gene trees progressively more complex and harder for individual groups to use. The Ensembl project provides an accurate and consistent protein coding gene set for all vertebrate genomes (Mouse Genome Sequencing Consortium, et al, 2002, Gibbs, et al, 2004, Mikkelsen, et al, 2007, Xie, et al, 2005, 'The International Human Genome Consortium', et al, 2001, Dehal, et al, 2002, Rhesus Macaque Genome Sequencing and Analysis Consortium, et al, 2007). Previously (until April 2006) Ensembl provided a basic method for tracing orthologs via the Best Reciprocal Blast method, similar to approaches used in other genome analyses such as *Drosophila melanogaster* (Adams, et al, 2000) or human ('The International Human Genome Consortium', et al, 2001). In June 2006 (Hubbard, et al, 2007), we replaced this system with a phylogenetically sound, genetree based approach, providing a complete set of phylogenetic trees spanning 91% of genes across vertebrates. In addition to the vertebrates we have included a few important non-vertebrate species (fly, worm and yeast) to act both as out groups and provide links to these model organisms. In this paper we provide the motivation, implementation and benchmarking of this method and document the

display and access methods for these trees.

There have been a number of methods proposed for routine generation of genomewide orthology descriptions, including InParanoid (Remm, et al, 2001), MSOAR (Fu, et al, 2007), OrthoMCL (Li, et al, 2003), Homologene (Wheeler, et al, 2008), TreeFam (Li, et al, 2006) and PhyOP (Goodstadt, et al, 2006) and PhiGs [REF]. The first four, InParanoid, MSOAR, OrthoMCL, and Homologene, focus on providing clusters (or linked clusters) of genes, without an explicit tree topology. PhyOP (Goodstadt, et al, 2006) uses a tree based method, but between pairs of closely related species, resolving paralogs accurately by using neutral substitution (as measured by dS, the synonymous substitution rate). TreeFam provides an explicit gene tree across multiple species, using both dS, dN (non-synonymous substitution rate), nucleotide and protein distance measures, and the standard species tree to balance duplications vs deletions to inform the tree construction, using the program TreeBeST (<http://treesoft.sourceforge.net/treebest.shtml>, Heng et al. submitted).

The PhiGs method (Dehal, et al, 2006) is a leading phylogenetic based method which produced a comprehensive phylogenetic resource for the genomes at the time it was run, and the basic outline of its analysis, which was clustering of protein sequences followed by phylogenetic trees is similar to the method presented here. However, the PhiGs resource covered a smaller number of species (23 vs 45) and has been difficult to keep up to date with the advances in gene sets and genomes. Another major difference between PhiG based phylogenetic trees and the phylogenetic trees presented here is that the former was calculated using a single maximum likelihood method based on protein evolution. In contrast the Ensembl gene trees are calculated using a new method, TreeBeST which integrates multiple tree topologies, in particular both DNA level and protein level models and combines this with a species-

tree aware penalisation of topologies which are inconsistent with known species relationships. We show in this paper that this method produces trees which are more consistent with synteny relationships and less anomalous topologies than single protein based phylogenetic methods.

There are also many single phylogenetic tree building approaches, many of them based on maximum likelihood methods; one leading method is PhyML (Guindon, et al, 2003). It is unclear what is the best method to use, in particular in the context of genome-wide tree building with constraints on computational costs and the need to robustly handle many complex scenarios usually involving large families with heterogeneous phylogenetic depths. In this paper, we benchmark in vertebrates the tree programs TreeBeST and PhyML, and the resulting trees to basic Best Reciprocal Hit (BRH) methods, and cluster frameworks, in particular Inparanoid and Homologene. We also benchmark to a recent PhyOP dataset. The PhyOP pipeline has recently switched to use the same tree building program (TreeBeST) that we use, but differs in its input clusters. Although we adopted this same tree building method, we describe here considerable novel engineering in the deployment of these methods across all vertebrates. Similar to the PhiGs resource, we have used the dense coverage of genomes to provide topologically based timings (i.e., the standard use of outgroups vs subsequent lineages to bracket a duplication), in order to label duplication events.

Methods

Gene synteny metric

In order to assess the quality of our orthology predictions, we developed a synteny metric that provides a measure of gene order conservation. The main idea is that when

a predicted ortholog between two species is flanked (by a distance criteria) by orthologous pairs on each side of each genome in a ordered manner, the central orthologous link is considered to be consistent with synteny. This measure can be applied to both 1-to-1 orthologs and to 1-to-many orthologs where a recent tandem duplication in one species has duplicated a gene as the criteria for flanking orthologous genes is based on distance, not gene order. Considering 2 species (e.g. human and mouse) and one species as reference (e.g. human), we called a *perfect syntenic gene* (on the reference species) a gene that has an orthology relation for which both upstream and downstream orthologies exist at 250kb and are co-linear in both species. We called a *good syntenic gene* (on the reference species) a gene that has an orthology relation for which one orthology exists at 250kb, either an upstream or downstream, that is co-linear in both species. It is important to note that we are using this metric to assess the quality of resulting trees, and not directly as part of our tree building procedure.

Duplication consistency score

In order to assess the reliability of the duplication calls on internal nodes of our tree, we developed a simply measure of the consistency of lineages post a putative duplication node. This measure is based on the assumption that duplication followed by reciprocal complementary gene losses on the left and right branch of a duplication node is an unlikely scenario (see main text).

Duplication score =

Intersection of species between left and right branches / Union of species between left and right branches

Pipeline framework

We created a fault-tolerant pipeline using Object-Oriented Perl and a MySQL database. The EnsemblCompara schema API sits on top of the main Ensembl schema and API, and links to BioPerl (Stajich, et al, 2002) objects for the main data types. The EnsemblCompara GeneTrees are updated every two months which involves being built from scratch for every Ensembl release over a two-week period using a cluster of computers, generating about 50GB of data.

Datasets used for assessing our pipeline

In order to compare the various pipeline implementations, we have performed all analyses on an identical data set from Ensembl v41 (October 2006). Assessment comprised implementations of (a) BRH alone and the tree-based programs; (b) PHYML followed by tree reconciliation with RAP, and (c) TreeBeST.

We also compared our Ensembl release 45 (June 2007), using the TreeBeST approach dataset against other method of predictions or databases such as Homologene, Inparanoid, PhyOP and TreeFam.

The species tree for RAP is an adapted tree from the ENCODE Multiple Sequence Analysis paper and can be found in the Supplementary information. (Margulies, et al, 2007)The species tree provided to TreeBeST only requires topological constraints, so we have used an adapted topology from the NCBI taxonomic tree which can be found in the Supplementary information.

Results

A robust, computationally efficient pipeline for gene tree generation

We have built a fault-tolerant pipeline to run our orthology and paralogy gene prediction analysis using TreeFam methodology. The fault-tolerance works at two levels – firstly we use a robust compute scheduling engine (in our case, LSF, though other packages could substitute for this component) to schedule jobs, but even with the use of LSF's scheduling and job recovery, there can be periodic network or disk failures which result in apparent successful LSF completion without data being successfully stored. Our experience is that a second level of data tracking is required, in particular due to the complex interdependence on compute results in the pipeline, which is hard to express as single static LSF based set of dependencies. Finally the pipeline allows aggregation of multiple highly similar compute tasks (in our case, BLAST comparisons) into a single LSF task, which is important to allow the granularity of the LSF tracking component to be optimised. The pipeline can be divided into 8 main steps that are presented in the schema in Figure 1. These 8 steps are described as follow:

- 1- *Protein dataset* : for each species considered in the analysis, we only consider protein coding genes. For each gene, we only consider the longest protein translation.
- 2- *Blastp all versus all* : each protein is queried using wublastp against each

individual species protein database, including its self-species protein database.

- 3- *Graph construction* : connections (edges) between the nodes (proteins) are retained when they satisfy either a Best Reciprocal Hit (BRH) or a Blast Score Ratio (BSR) over 0.33.

A BSR for two proteins, P1 and P2, is defined as $\text{score}_{P1P2} / \max(\text{self-score}_{P1} \text{ or } \text{self-score}_{P2})$.

- 4- *Clusters* : we extract from the graph the connected components (i.e. single linkage clusters). Each connected component represents a cluster, i.e. a gene family. If the cluster has greater than 750 members, steps 3 and 4 are repeated at higher stringency (see below).
- 5- *Multiple alignment* : proteins in the same cluster are aligned using MUSCLE (Edgar, 2004) to obtain a multiple alignment.
- 6- *Gene tree and reconciliation* : the CDS back-translated protein-based multiple alignment is used as an input to the tree program, TreeBeST, as well as the multifurcated species tree necessary for the reconciliation and the duplication calls on internal nodes.
- 7- *Inference of orthologs and paralogs* : As many users like to use ortholog look-up tables, we flatten the resulting trees into ortholog and paralog tables of pairwise relationships between genes. In the case of paralogs, this flattening also records the timing of the duplication due to the presence of extant species past the duplication, and thus implicitly outgroup lineages before the duplication (See Supplementary Figures 1A and 1B for a detailed explanation).

- 8- Pairwise dN/dS (*non-synonymous substitutions/synonymous substitutions*) : we calculate the pairwise dN/dS between pairs of genes for closely-related species using *codeml* from the PAML package (Yang, 2007).

For the Ensembl v41 assessment, step 6 was divided in step 6a using PHYML (Guindon, et al, 2003) to build the tree, and step 6b using RAP (Dufayard, et al, 2005) for the tree reconciliation.

At the end of step 4, if the cluster is large (currently parametrised as containing more than 750 genes), the genes in this cluster are then re-injected into step 3 (Figure 1, dashed lines), with only the BSR threshold condition to satisfy. If more iterations are necessary, the BSR threshold is increased by 0.1 at each iteration. The same kind of iterations are applied at the end of steps 5 and 6 when MUSCLE or the tree building program failed to process a cluster. This iteration procedure is effectively a hierarchical breakdown of the initial clustering to get more fine grained set of clusters that can easily be processed. This iterative approach is critical to generate sensible trees for complex large families, such as those of zinc finger proteins or olfactory receptors. Although it would be desirable to place all genes from these gene families into a comprehensive single tree, there are numerous engineering, algorithmic and display problems associated with large trees, and the hierarchical breakdown provides a pragmatic solution for such families.

The TreeBeST method has two new components. Firstly it runs a number of independent phylogenetic methods, in particular DNA, codon and protein maximum likelihood models are created on the same data. The TreeBeST method then creates a

combined tree using a stochastic context free grammar approach to integrate the different tree information with a model to penalise duplications and deletions relative to a known species tree. The result is that TreeBeST will tend to use DNA or codon based methods in the parts of the phylogeny which do not have saturated DNA mutation rates (for example, intra-mammalian comparisons) but utilises protein information at longer distances (for example, comparisons between mammals and fish).

We developed two metrics to assess the different methods.

Duplication consistency score

We developed a consistency score for proposed duplications, where we measure the intersection of the number of species post duplication over the union; one expects that most duplications should have the gene persisting at least in an equally likely manner in subsequent lineages. In contrast, incorrect topologies will often have simply reordered a deep node leading to usually a few species in the topologically incorrect positions; reconciliation to the species tree then forces the prediction of duplication followed by extensive loss in a precisely correlated manner across the two daughter lineages. The duplication consistency score captures this unbalanced nature of poor topologies as the intersection in subsequent lineages is low. Figure 2 shows clearly that the PHYML/RAP approach made many more duplication nodes compared to TreeBeST, and the vast majority of the additional duplications from PHYML/RAP have a low duplication consistency score. This result is unsurprising as TreeBeST takes the species tree as input and explicitly penalises both duplication and deletion of genes, in other words, the TreeBeST program tends to produce duplication nodes

when the gene tree has extensive extant members on each side of the duplication. Although this metric fundamentally reflects the difference in methodology between PhyML, a pure sequence based tree, and TreeBeST, which uses the species tree as input, it is clear that the TreeBeST results are more biologically consistent, given the assumption that gene duplication and deletion rates are rare.

Gene synteny metric

We also developed an alternative metric which was not confounded by the tree methodology using the fact that gene order and orientation (informally called synteny) are conserved across species. None of the tree approaches used synteny information in the tree construction, though the old Best Reciprocal Hit method extended its range using synteny information. Supplementary Figure 2 shows the difference in results between a strict BRH approach with no syntenic information used, PHYML/RAP and TreeBeST. In both cases, for perfect and good syntenic genes, the TreeBeST pipeline shows better results. PHYML/RAP gave poorer results than BRH. We believe this was mainly due to a large number of wrong gene tree topologies and hence difficult tree reconciliations that overestimated duplication events. Such overestimation led to missed orthology predictions.

Comparison to bootstrap metrics

We compared the duplication consistency measure to the bootstrap support of the duplication nodes from TreeBeST. As expected, there is a strong correlation between the bootstrap support and the duplication consistency measure, with most of the high

duplication consistency measure scores also having high bootstrap support (Figure 3), and low duplication consistency measures having a variety of bootstrap scores, but nearly always below the 80% support level. Interestingly, there was a set of low bootstrap but high duplication-consistent set of duplication nodes (bottom right in Figure 3), but not the inverse set of high bootstrap low duplication-consistent nodes. These duplication nodes were not obviously correlated with either internal aspects of the multiple alignment or tree (for example, length or average distance from the duplication node to extant species) nor external properties of the genes (for example, GO term distribution, Pfam domain sets or the position of the duplication node with respect to the vertebrate tree). This set of genes might reflect the fact that the bootstrap statistic is about the consistency of the tree across the columns of the multiple alignment, and this consistency measure does not necessarily have to apply to every gene equally. In contrast, the duplication consistency measure, which is a property of the behaviour of genes post duplication, may be more consistent across different genes.

The conclusion of this investigation is that TreeBeST was the best of these extensively tested methods on the criteria of duplication consistency and synteny consistency criteria. It is hard to have entirely objective measures of the accuracy of trees (see discussion below). We also briefly investigated further tree programs and further multiple alignment programs (e.g., ClustalW), but many of these were not robust enough to work in a large scale compute environment with the complex gene families present across vertebrates. In the future these metrics will permit the testing of both other tree construction programs and other multiple alignment programs, and we will continue to test and assess new robustly engineered programs with a good

chance of improving the trees.

External Benchmarking to other orthology sets

Overlap of ortholog sets

Table 1 shows the overlap of ortholog sets between EnsemblCompara GeneTrees v45 to InParanoid, Homologene, PhyOP or TreeFamCurated for certain pairs of species. In all our comparisons we have taken genes as reference, and have counted for each gene its best homology prediction. The ranking from best to worst favoured one-to-one orthologs over one-to-many orthologs and both over paralogs. When a gene was not involved in any homology relation, it has been labelled as unclassified.

In all the data shown in the tables, EnsemblCompara always shows better or similar coverage to any other method. This is clearly visible in Homologene, where 1/3 more human genes and 2-fold more mouse genes are lost in Homologene as compared with EnsemblCompara GeneTrees v45. Part of this large difference is the absence of RefSeq (Pruitt, et al, 2007) entries for particular human genes, i.e., a problem with gene prediction sets or coordination between Ensembl and RefSeq IDs in the genomes rather than the inability to create an orthology relationship. As Homologene is a database, and not a method that can be applied to a new dataset, one cannot perform a perfectly matched set. We then restricted the 22,568 human protein coding genes and 24,496 mouse protein coding genes present in the Ensembl database to the common RefSeq set used as an input to Homologene to compare the homology types as fairly as possible between the two datasets. For this set there were 838 Homologene associations which could have been made in EnsemblCompara, compared to 1,519

EnsemblCompara cases between genes with RefSeq IDs but no Homologene association. Manual inspection of these cases show some complex tree topologies but also clear cases of one-to-one orthology that had been missed in Homologene (e.g. the *MAGIX* gene), whereas the majority of the missing EnsemblCompara cases came from complex scenarios with unclear correct tree topologies, such as Ig locus genes.

When comparing our results with InParanoid (Remm, et al, 2001) we used a matched protein set (Ensembl v45) for both methods. We observed that although they return very similar results, EnsemblCompara has increased coverage with marginally increased specificity (see below for specificity measure). The gain in gene coverage in favour of EnsemblCompara v45 becomes clearer when looking at more distant species such as human/medaka or human/drosophila.

The PhyOP pipeline has recently moved to using the same tree building program (TreeBeST) as TreeFam and EnsemblCompara. This means any difference is due to the input clusters. The PhyOP pipeline shows marginally less unclassified genes than the EnsemblCompara pipeline, i.e., having orthologous genes predicted which were not present in EnsemblCompara. Examination of these cases showed many genes involved in large families. Currently, EnsemblCompara handles 35 species compared to the more restricted set of 6 species in the PhyOP run, and it seems that in breaking down the large families into appropriate clusters, some genes in these large families can become orphaned. This is clearly an area which can be improved in the future.

The TreeFamCurated entry corresponds to the comparison of our dataset against the curated set of TreeFam, with 1,247 such cases only. The curated trees in TreeFam incorporate expert knowledge to change the topology of trees, for example, by using

information on the conservation of function. In Table 1 we show that the concordance between our automated prediction set and the manually curated TreeFam dataset is very high. The only exception seems to be that our method tends to miss orthology relationships in favour of within-species paralogs. We believe this is mainly due to wrong tree topologies involving mispredicted (merged/split/partial) genes for which automatic tree building has difficulties to place the genes correctly. The manual curation in TreeFam then corrects this problem and results in a better tree topology. In the long term, the incorporation of more manual curation into the human and mouse gene sets, coupled with more improvements in the gene prediction methodology in Ensembl should progressively remove these errors.

Comparison using the Synteny metric

We were interested in looking at the differences in the synteny metric, as defined above, between the different methods as indicative of their specificity. The plot in Figure 4 shows the number of human genes involved in homology relations as a function of the number of human syntenic genes. EnsemblCompara and PhyOP always perform better in terms of the number of covered human genes, but there is a remarkably similar level of syntenous predictions between all the different methods, with in some cases Inparanoid showing a marginally higher rate of syntenous predictions. In the case of a distant vertebrate species such as medaka, EnsemblCompara is best on both coverage and specificity measures. The teleost genomes represent a particular challenge for the clustering mechanism due to the ancient duplication at the root of the teleost lineage, leading to proportionally more

“ancient” paralog relationships which are hard to capture using the clustering methods.

Display and access of orthologs

We provide different ways to access and visualize the orthology/paralogy data and have used it in a number of ways in house.

Web display

The main entry points are GeneView

(http://jun2007.archive.ensembl.org/Homo_sapiens/geneview?gene=ENSG00000129965)

and GeneTreeView

(http://jun2007.archive.ensembl.org/Homo_sapiens/genetreeview?db=core;gene=ENSG00000129965 and Figure 5).

In GeneView, we list the orthologous and within-species paralogous gene predictions.

In each case, the user has access to MultiContigView, a display that shows the ortholog or paralog relation in the genomic the context of both species. The user can also see the alignment between the 2 ortholog/paralog protein sequences via AlignView.

GeneTreeView (Figure 5) displays the gene tree and shows the considered gene highlighted in red in context of all its homologous relations. Duplication nodes are coloured red, whereas speciation nodes are coloured blue. The user can dump

multiple alignment of this gene tree with the “Export” menu, as well as a picture of the tree in different formats (PDF, PS and SVG). Future development will include zoom in/out at specific internal nodes to display subtrees. We will also include the ability to dump the gene list of the whole tree and a subtree, as well as the multiple alignment of the protein/CDS in a subtree.

Projection of GO terms via orthology links

One of the benefits of extensive and accurate prediction of orthologs is that one can infer that they have (usually) retained the same function in extant species. Using this methodology we have automatically projected GO terms from the main two mammalian sources, human and mouse, out across other vertebrate species. When we project GO terms, we tag the evidence as “Inferred from Electronic Annotation” (IEA), consistent with other GO annotations, to prevent confusion with directly assigned GO terms, and we only project from experimentally referenced GO evidence in the source organism. After discussion with the GO community we have only projected via one-to-one ortholog links, though it is worth considering in the future a more flexible approach of projection through duplications for some terms (for example, molecular function terms may rarely be changed by recent duplication structure, while biological process terms may change more frequently). Table 2 shows the set of species for which we have projected GO terms and the comparison with existing sets. Even in the well annotated human and mouse genomes, this projection provides a small increase in the overall number of genes and a marked increase when not considering genes already IEA annotated. Currently the bulk of IEA assignments come via domain matching (eg, Interpro2Go), and thus have to use quite broad specificity terms, whereas our ortholog annotation can provide far more detailed GO

terms. Of course in less intensively studied genomes this creates a large set (e.g. ~5,000 previous annotated genes for dog) of GO mappings.

Datamining using BioMart

BioMart is a flexible data mining application which can be addressed using a user-friendly web page, programmatic access, webservice access and the biomaRt package in the R statistical environment. BioMart enables the user to do bulk dumps of ortholog or paralog gene pair lists given a species pair and to restrict this by any valid BioMart query. It can also to dump the peptide/cDNA sequence of the gene in question.

Raw dump accessible via ftp

We also provide dumps of the gene tree multiple alignments and the trees themselves in “emf” format, described in more detail at

[ftp://ftp.ensembl.org/pub/current_multi_species/data/emf/protein_trees/README](http://ftp.ensembl.org/pub/current_multi_species/data/emf/protein_trees/README).

The tree written down in newick format, embedded in the emf format itself, such there is only one file representing the entire dataset. We are developing format readers in the BioPerl libraries to ensure easy integration of this flat file data into other pipelines in a standalone manner.

Programmatic access using the Perl API

The data stored in a Ensembl Compara database is finally also accessible in programmatic way using a Perl API. More detailed documents and tutorials on how to install and use the API can be found here <http://www.ensembl.org/info/software/index.html>. Supplement 1 shows three

examples of Perl scripts using the API.

Discussion

The orthology pipeline presented here is robust and provides a framework in which we can assess different components in tree generation. We have used three key metrics; coverage in homology relationships, duplication consistency score and consistency with genome synteny, to assess both different internal components of our pipeline and to other orthology sets. In our assessments the Muscle+TreeBeST system, which is the set of methods used in TreeFam, performs best according to these metrics. One problem in phylogenetic method development is that it is hard to have access to objectively correct trees to assess methods. Simulation based assessment can explore the potential source of errors, and TreeBeST performs well, and critically better than sequence-only methods, with simulated data (Heng et al., submitted). Of the three metrics used in this paper, the first two (coverage and duplication consistency) are somewhat arbitrary choices which nevertheless correspond to observations about “poor” trees from biological experts who use other information (such as conservation of function) to infer orthology. Such observations are necessarily anecdotal, but having methods which produce trees with higher coverage and less duplication followed by mirroring loss in the daughter lineages is more consistent with biological expertise. This is shown by the comparison to the curated TreeFam trees, which attempt to capture systematically this expert knowledge for a subset of trees. The third metric, the conservation of synteny for orthologous genes in mammals, is more principled. However, new methods integrating this information into phylogenetic methods, such as MSOAR (Fu, et al, 2007) method

and (Jiang, et al, 2007), could provide more accurate trees at the expense of not being able to use synteny to assess accuracy.

In comparisons to other genome-wide frameworks, mainly cluster based, these methods performed better in terms of coverage with at least as good specificity, as measured by the synteny metric. In particular much improvement is seen in the teleost lineage, where the complex ancient duplication structure which has been differentially lost in extant species leads to more complex phylogenies. In addition, this phylogenetic method provides a far richer dataset including the topological timings of duplications and the ability to implement other tree-dependent methods, such as global dN/dS methods. Obviously, one can expect improvements in both alignments and tree methods in the future, and this framework is flexible enough both to assess and replace the components we are using currently.

The phylogenetic information presented here is now a standard part of the Ensembl system, and will be present in all future releases, as well as available through the Ensembl archives. This provides both an individual gene-specific biologist the opportunity to explore the evolution of a gene family, discovering potentially unappreciated ancestral duplications, or draws his or her attention to other lineages where a gene has been duplicated. As the presence of recent lineage-specific duplications is often associated with positive selection, this could lead a biologist to look into the specific biology of a previously unappreciated species to understand the functional role of a gene. More mundanely, the presence of these accurate orthology links allows other groups in Ensembl to provide appropriate projection of information across the vertebrate lineages, using the concentration of information on human and mouse to inform all of the species in the vertebrate tree. This is a great boon when coupled with the GO annotation dictionary, and also allows us to project the HGNC

symbols across species confidently to provide a useful visual tag for genes in different species. Ensembl also includes the MCL (Enright, et al, 2002) generated Ensembl families resource. This is a clustering based method designed to work at a far deeper phylogenetic distance (incorporating events in protein families that occurred during early eukaryotic evolution) than the ortholog prediction framework presented here. There are both conceptual problems due to large scale domain changes over this depth of evolution, which in some cases involve genuine gene split and merge events, and engineering problems due to the considerable increase in family membership when working at this scale. We are currently investigating ways both to deepen our gene family clusters and to reconcile these deeper families to broader protein family representation, such as Tribe-MCL, but currently both methods show complementary aspects of protein evolution.

Acknowledgements

We would like to acknowledge Chris Ponting, Leo Goodstadt and Andreas Heger for providing the recent PhyOP run and many useful discussions. We also want to thank the rest of the TreeFam and the Ensembl teams for many insightful discussions. We are grateful to Michael Hoffman for his assistance in using R. Also, we thank the Sanger Institute computer services for the reliable running of the compute farm. A.V, A. U-V, J. S were supported by the Wellcome Trust. E.B was supported by EMBL. R.D. and L.H. were supported by the Wellcome Trust via the Sanger Institute.

Figure Legends

Figure 1. A figure showing the computational pipeline for the EnsemblCompara process.

Figure 2. A diagram of the duplication consistency score on an example tree showing unlikely coordinated deletions on the subsequent lineages. The histogram shows the distribution of consistency scores for both PhyML/RAP and TreeBeST methods. PhyML/RAP has both a higher absolute number of duplications and far more at low consistency values

Figure 3. A scatter plot of the duplication consistency score (x-axis) compared to the bootstrap value of duplication nodes (y-axis). Because of the large number of values, the density of points is shown using the smoothScatter kernel based density function in R.

Figure 4. A plot showing different methods in terms of coverage in human genes (x-axis) vs number of genes in syntenic relationships (y-axis).

Figure 5. A screen shot of the gene tree page at Ensembl for the INS (insulin peptide) gene. This shows two independent duplications in rodents (giving rise to Ins1 and Ins2 genes) and Teleost fish. Duplication nodes are shown as red squares whereas speciation nodes are in blue. The green bars to the right provide a graphical view of the multiple alignment, showing partial gene structures in Hamster (*Cavia p.*), Cat (*Felis c.*) and Rabbit (*Oryctolagus c.*), due to their low coverage status.

Tables

Homologene vs EnsemblCompara v45

human/mouse

	ortholog one2one	app ortholog one2one	ortholog one2many	ortholog many2many	within species paralog	unclassified	TOTAL
ortholog_one2one	9284/9284	349/351	577/592	116/110	200/196	97/90	10623/10623
ortholog_one2many	2055/2048	92/90	228/287	70/83	89/107	18/24	2552/2639
ortholog_many2many	283/284	13/12	25/48	20/36	76/123	19/33	436/536
within_species_paralog	0/1	0/0	39/82	43/67	119/247	1/0	202/397
unclassified	55/60	8/9	372/874	223/329	2589/3668	5508/5361	8755/10301
TOTAL	11677/11677	462/462	1241/1883	472/625	3073/4341	5643/5508	22568/24496

Phyop vs EnsemblCompara v45

human/mouse

	ortholog one2one	app ortholog one2one	ortholog one2many	ortholog many2many	within species paralog	unclassified	TOTAL
ortholog one2one	12622/12597	981/970	180/175	19/12	77/60	447/512	14326/14326
app ortholog one2one	494/493	60/60	15/13	1/2	7/8	41/42	618/618
ortholog one2many	164/164	12/11	1081/1556	41/87	83/135	229/546	1610/2499
ortholog many2many	8/7	2/3	68/68	331/390	43/41	144/309	596/818
within species paralog	128/143	19/31	333/438	186/253	497/838	1239/1704	2402/3407
unclassified	163/175	33/32	81/58	10/6	105/88	2624/2469	3016/2828
TOTAL	13579/13579	1107/1107	1758/2308	588/750	812/1170	4724/5582	22568/24496

**Inparanoid vs
EnsemblCompara v45****human/mouse**

	ortholog one2one	app ortholog one2one	ortholog one2many	ortholog many2many	within species paralog	unclassified	TOTAL
ortholog one2one	14001/13984	545/546	489/492	25/26	137/153	133/129	15330/15330
ortholog one2many	47/37	4/8	771/1134	126/141	284/270	31/20	1263/1610
ortholog many2many	7/7	2/2	20/33	249/314	181/163	2/0	461/519
unclassified	271/298	67/62	330/840	196/337	1800/2821	2850/2679	5514/7037
TOTAL	14326/14326	618/618	1610/2499	596/818	2402/3407	3016/2828	22568/24496

human/medaka

	ortholog one2one	app ortholog one2one	ortholog one2many	ortholog many2many	within species paralog	unclassified	TOTAL
ortholog one2one	7563/7551	169/167	1514/1519	50/48	188/177	523/545	10007/10007
ortholog one2many	100/81	9/9	1148/1194	137/97	392/207	139/91	1925/1679
ortholog many2many	8/6	2/3	26/20	159/126	480/223	48/41	723/419
unclassified	854/887	65/66	848/1768	234/185	3374/2822	4538/2298	9913/8026
TOTAL	8525/8525	245/245	3536/4501	580/456	4434/3429	5248/2975	22568/20131

human/drosophila

	ortholog one2one	app ortholog one2one	ortholog one2many	ortholog many2many	within species paralog	unclassified	TOTAL
ortholog one2one	2673/2674	14/14	354/353	33/30	69/26	236/282	3379/3379
ortholog one2many	37/38	2/2	3094/1563	252/167	848/85	540/433	4773/2288
ortholog many2many	3/3	0/0	56/38	349/308	219/203	86/82	713/634
unclassified	253/251	0/0	1321/368	299/358	5714/1454	6116/5307	13703/7738
TOTAL	2966/2966	16/16	4825/2322	933/863	6850/1768	6978/6104	22568/14039

**TreeFamCurated vs
EnsemblCompara v45****human/mouse**

	ortholog one2one	app ortholog one2one	ortholog one2many	ortholog many2many	within species paralog	unclassified	TOTAL
ortholog one2one	1542/1539	82/83	40/38	0/1	11/19	17/16	1692/1696
app ortholog one2one	1/1	0/0	0/0	0/0	0/0	0/0	1/1
ortholog one2many	24/21	6/6	104/138	4/2	31/36	4/2	173/205
ortholog many2many	4/3	1/1	2/6	23/26	14/16	3/1	47/53
within species paralog	14/8	1/1	0/0	0/0	6/8	2/3	23/20
unclassified	12741/12754	528/527	1464/2317	569/789	2340/3328	2990/2806	20632/22521
TOTAL	14326/14326	618/618	1610/2499	596/818	2402/3407	3016/2828	22568/24496

Table 1. Table showing the comparison of EnsemblCompara vs other methods. In each case, the different Ensembl categories are listed in the columns, whereas the comparison database is listed in rows. Each cell shows the number of gene IDs for the two species for the intersection of the two categories. As well as one-to-many and many-to-many relationships making the between species numbers different, each homology program can make a different pairing of genes leading to different numbers in the one2one category. The unclassified column shows genes not captured at all by that method.

Species	Genes with all GO terms	Genes with Non-IEA GO terms	Additional Genes	Additional genes, discounting IEA	Total protein coding genes
Human	16758	10314	143	6587	22680
Mouse	17664	10454	108	7318	24118
Rat	11434	2006	1779	11207	22993
Dog	1262	74	5207	6395	19305
Cow	5662	330	3220	8552	21755
Chicken	3636	264	2712	6084	16736

Table 2. GO Term projection. Table 2 shows the number of genes associated with GO terms in different species. The first column shows the number of genes with GO terms in total, whereas the second column shows the number of genes with non “inferred by electronic annotation” (IEA) terms. The third column shows the number of additional genes with a GO term added by projection, and the fourth column shows the number of genes with a GO term added including cases which previously only had IEA terms.

References

- Adams, M. D., Celniker, S. E., Holt, R. A., Evans, C. A., Gocayne, J. D., Amanatides, P. G., Scherer, S. E., Li, P. W., Hoskins, R. A., Galle, R. F., et al. 2000. The genome sequence of *Drosophila melanogaster*. *Science* **287**: 2185-2195.
- Dehal, P., Satou, Y., Campbell, R. K., Chapman, J., Degnan, B., De Tomaso, A., Davidson, B., Di Gregorio, A., Gelpke, M., Goodstein, D. M., et al. 2002. The draft genome of *Ciona intestinalis*: insights into chordate and vertebrate origins. *Science* **298**: 2157-2167.
- Dehal, P. S. and Boore, J. L. 2006. A phylogenomic gene cluster resource: the Phylogenetically Inferred Groups (PhIGs) database. *BMC Bioinformatics* **7**: 201.
- Dufayard, J. F., Duret, L., Penel, S., Gouy, M., Rechenmann, F., and Perriere, G. 2005. Tree pattern matching in phylogenetic trees: automatic search for orthologs or

- paralogs in homologous gene sequence databases. *Bioinformatics* **21**: 2596-2603.
- Edgar, R. C. 2004. MUSCLE: a multiple sequence alignment method with reduced time and space complexity. *BMC Bioinformatics* **5**: 113.
- Enright, A. J., Van Dongen, S., and Ouzounis, C. A. 2002. An efficient algorithm for large-scale detection of protein families. *Nucleic Acids Res.* **30**: 1575-1584.
- Fu, Z., Chen, X., Vacic, V., Nan, P., Zhong, Y., and Jiang, T. 2007. MSOAR: a high-throughput ortholog assignment system based on genome rearrangement. *J. Comput. Biol.* **14**: 1160-1175.
- Gibbs, R. A., Weinstock, G. M., Metzker, M. L., Muzny, D. M., Sodergren, E. J., Scherer, S., Scott, G., Steffen, D., Worley, K. C., Burch, P. E., et al. 2004. Genome sequence of the Brown Norway rat yields insights into mammalian evolution. *Nature* **428**: 493-521.
- Goodstadt, L. and Ponting, C. P. 2006. Phylogenetic reconstruction of orthology, paralogy, and conserved synteny for dog and human. *PLoS Comput. Biol.* **2**: e133.
- Guindon, S. and Gascuel, O. 2003. A simple, fast, and accurate algorithm to estimate large phylogenies by maximum likelihood. *Syst. Biol.* **52**: 696-704.
- Hennig, W. 1952 *Grundzüge der Theorie der Phylogenetischen Systematik*. Deutscher Zentralverlag, Berlin.
- Hubbard, T. J., Aken, B. L., Beal, K., Ballester, B., Caccamo, M., Chen, Y., Clarke, L., Coates, G., Cunningham, F., Cutts, T., et al. 2007. Ensembl 2007. *Nucleic Acids Res.* **35**: D610-7.

International Chicken Genome Sequencing Consortium. 2004. Sequence and comparative analysis of the chicken genome provide unique perspectives on vertebrate evolution. *Nature* **432**: 695-716.

Jiang, Z., Tang, H., Ventura, M., Cardone, M. F., Marques-Bonet, T., She, X., Pevzner, P. A., and Eichler, E. E. 2007. Ancestral reconstruction of segmental duplications reveals punctuated cores of human genome evolution. *Nat. Genet.* **39**: 1361-1368.

Li, H., Coghlan, A., Ruan, J., Coin, L. J., Heriche, J. K., Osmotherly, L., Li, R., Liu, T., Zhang, Z., Bolund, L., et al. 2006. TreeFam: a curated database of phylogenetic trees of animal gene families. *Nucleic Acids Res.* **34**: D572-80.

Li, L., Stoeckert, C. J., Jr, and Roos, D. S. 2003. OrthoMCL: identification of ortholog groups for eukaryotic genomes. *Genome Res.* **13**: 2178-2189.

Margulies, E. H., Cooper, G. M., Asimenos, G., Thomas, D. J., Dewey, C. N., Siepel, A., Birney, E., Keefe, D., Schwartz, A. S., Hou, M., et al. 2007. Analyses of deep mammalian sequence alignments and constraint predictions for 1% of the human genome. *Genome Res.* **17**: 760-774.

Mikkelsen, T. S., Wakefield, M. J., Aken, B., Amemiya, C. T., Chang, J. L., Duke, S., Garber, M., Gentles, A. J., Goodstadt, L., Heger, A., et al. 2007. Genome of the marsupial *Monodelphis domestica* reveals innovation in non-coding sequences. *Nature* **447**: 167-177.

Mouse Genome Sequencing Consortium, Waterston, R. H., Lindblad-Toh, K., Birney, E., Rogers, J., Abril, J. F., Agarwal, P., Agarwala, R., Ainscough, R., Alexandersson, M., et al. 2002. Initial sequencing and comparative analysis of the mouse genome.

Nature **420**: 520-562.

Pruitt, K. D., Tatusova, T., and Maglott, D. R. 2007. NCBI reference sequences (RefSeq): a curated non-redundant sequence database of genomes, transcripts and proteins. *Nucleic Acids Res.* **35**: D61-5.

Remm, M., Storm, C. E., and Sonnhammer, E. L. 2001. Automatic clustering of orthologs and in-paralogs from pairwise species comparisons. *J. Mol. Biol.* **314**: 1041-1052.

Rhesus Macaque Genome Sequencing and Analysis Consortium, Gibbs, R. A., Rogers, J., Katze, M. G., Bumgarner, R., Weinstock, G. M., Mardis, E. R., Remington, K. A., Strausberg, R. L., Venter, J. C., et al. 2007. Evolutionary and biomedical insights from the rhesus macaque genome. *Science* **316**: 222-234.

Stajich, J. E., Block, D., Boulez, K., Brenner, S. E., Chervitz, S. A., Dagdigian, C., Fuellen, G., Gilbert, J. G., Korf, I., Lapp, H., et al. 2002. The Bioperl toolkit: Perl modules for the life sciences. *Genome Res.* **12**: 1611-1618.

'The International Human Genome Consortium', Lander, E. S., Linton, L. M., Birren, B., Nusbaum, C., Zody, M. C., Baldwin, J., Devon, K., Dewar, K., Doyle, M., et al. 2001. Initial sequencing and analysis of the human genome. *Nature* **409**: 860-921.

Wheeler, D. L., Barrett, T., Benson, D. A., Bryant, S. H., Canese, K., Chetvernin, V., Church, D. M., Dicuccio, M., Edgar, R., Federhen, S., et al. 2008. Database resources of the National Center for Biotechnology Information. *Nucleic Acids Res.* **36**: D13-21.

Xie, X., Lu, J., Kulbokas, E. J., Golub, T. R., Mootha, V., Lindblad-Toh, K., Lander,

E. S., and Kellis, M. 2005. Systematic discovery of regulatory motifs in human promoters and 3' UTRs by comparison of several mammals. *Nature* **434**: 338-345.

Yang, Z. 2007. PAML 4: phylogenetic analysis by maximum likelihood. *Mol. Biol. Evol.* **24**: 1586-1591.

1 Load genes and longest translations
for all species in Ensembl

Downloaded from genome.cshlp.org on July 17, 2020. Published by Cold Spring Harbor Laboratory Press

2 blastp each longest translation against every
other species protein dataset, incl. self-species

3 Build a graph of protein relations based on
Blast Reciprocal Hits or Blast Score Ratio

4 Extract the connected components using single
linkage clustering with the groups of peptides

5 Generate a protein alignment for each cluster

6 Build a gene tree and reconciling it
with the species tree

7 Infer the orthology and paralogy relationships
for every pair of genes in the gene tree

8 Calculate the pairwise dN/dS ratio for
closely-related pairs of genes in the gene tree









