



***De novo* search for non-coding RNA genes in the AT-rich genome of *Dictyostelium discoideum*: performance of Markov-dependent genome feature scoring**

Pontus Larsson, Andrea Hinas, David H Ardell, et al.

Genome Res. published online March 17, 2008

Access the most recent version at doi:[10.1101/gr.069104.107](https://doi.org/10.1101/gr.069104.107)

P<P Published online March 17, 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.

License

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](#).



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

Copyright © 2008, Cold Spring Harbor Laboratory Press

***De novo* search for non-coding RNA genes in the AT-rich genome of
Dictyostelium discoideum: performance of Markov-dependent genome
feature scoring**

**Pontus Larsson¹, Andrea Hinas^{2,4}, David H Ardell^{3,5*}, Leif A Kirsebom¹, Anders
Virtanen¹, and Fredrik Söderbom^{2,*}**

**¹Department of Cell and Molecular Biology, Biomedical Center, Uppsala University,
Box 596, S-75124 Uppsala, Sweden. ²Department of Molecular Biology, Biomedical
Center, Swedish University of Agricultural Sciences, Box 590, S-75124 Uppsala,
Sweden. ³Linnaeus Centre for Bioinformatics, Biomedical Center, Box 598, SE-751
24 Uppsala Sweden. ⁴Present address: Department of Molecular and Cellular
Biology, Harvard University, 16 Divinity Ave, Room 3050, Cambridge MA, 02138,
USA. ⁵Present address: School of Natural Sciences, University of California,
Merced, CA, 95344, USA.**

***Corresponding authors**

**David H Ardell: Dave.Ardell@lcb.uu.se; dardell@ucmerced.edu; phone 1 (209) 228-
2953; fax: 1 (209) 228-4060**

**Fredrik Söderbom: fredde@xray.bmc.uu.se; phone +46 (0)18 471 49 01; fax +46 (0)18
536971**

Running title: *de novo* identification of ncRNA genes

Key words: ncRNA, Dictyostelium, Karlin-Dembo statistics, Markov dependent partial sum, dinucleotide composition

Abstract

Genome data are increasingly important in the computational identification of novel regulatory non-coding RNAs (ncRNAs). However, most ncRNA gene-finders are either specialized to well-characterized ncRNA gene families or require comparisons of closely related genomes. We developed a method for *de novo* screening for ncRNA genes with a nucleotide composition that stands out against the background genome based on a partial sum process. We compared the performance when assuming independent and first-order Markov dependent nucleotides, respectively and used Karlin-Altschul and Karlin-Dembo statistics to evaluate significance of hits. We hypothesized that a first-order Markov-dependent process might have better power to detect ncRNA genes since nearest-neighbor models have shown to be successful in predicting RNA structures. A model based on a first-order partial sum process (analyzing overlapping dinucleotides) had better sensitivity and specificity than a zeroth-order model when applied to the AT-rich genome of the amoeba *Dictyostelium discoideum*. In this genome we detected 94 percent of previously known ncRNA genes (at this sensitivity, the false positive rate was estimated to 25% in a simulated background). The predictions were further refined by clustering candidate genes according to sequence similarity and/or searching for an ncRNA-associated upstream element. We experimentally verified six out of ten tested ncRNA gene predictions. We conclude that higher-order models, in combination with other information, are useful for identification of novel ncRNA gene families in single genome analysis of *D. discoideum*. Our generalizable approach extends the range of

genomic data that can be searched for novel ncRNA genes using well-grounded statistical methods.

Supplemental material is available online at <http://www.genome.org>. Sequence data from this study have been submitted to GenBank under accession nos. EF551319 and EF551320

Introduction

Non-coding RNAs (ncRNAs) have long been regarded as the exception to the rule that proteins are the major functional component of the cell. However, the recent discovery of a large number of functional ncRNAs, such as microRNAs (miRNAs), have led to a paradigm shift in which ncRNAs are seen to play a major role in the regulation of gene expression (Aravin and Tuschl 2005).

Computational gene prediction plays an increasingly important, but challenging, role in the identification of new ncRNA genes (Hüttenhofer and Vogel 2006). Although ncRNA gene homologs often contain sequence motifs that are conserved among distantly related organisms, these motifs are generally too short to be used as a sole search criterion. In addition, the structure of an ncRNA is often important for its function but difficult to confidently predict or use for genomic searches (Eddy 2002). Hence, methods for computational identification of ncRNA genes are usually specialized for a particular class of ncRNAs whose characteristics are already well-known e.g. tRNAs, snoRNAs, or RNase P RNA (Edvardsson et al. 2003; Griffiths-Jones et al. 2005; Kyriakopoulou et al. 2006; Laslett and Canbäck 2004; Lowe and Eddy 1997; Lowe and Eddy 1999; Piccinelli et al. 2005; Schattner et al. 2006). In contrast, *de novo* ncRNA gene prediction without *a priori* knowledge of RNA structure and function is difficult. Most progress in this area relies on the identification of conserved sequences and predicted structures through genome comparisons (Havgaard et al. 2005; Pedersen et al. 2006; Rivas and Eddy 2001;

Washietl et al. 2005). Therefore, there is room for improvement of methods to find new ncRNA genes in genomes for which related genomes have not yet been sequenced.

The organism in this study, the social amoeba *Dictyostelium discoideum*, has an AT-rich (78%) genome of 34 Mb and about 12,500 protein coding genes (Eichinger et al. 2005). Evolutionarily, *D. discoideum* is placed between the plant and animal-fungi lineages (Baldauf et al. 2000). *D. discoideum* lives on the forest floor as single cells where they feed on bacteria. Upon starvation, the cells embark upon a developmental program where up to 100,000 cells move together and differentiate, forming a multicellular-like organism. Recent experimental and computational analyses have identified a large number of expressed and predicted ncRNA genes in *D. discoideum*, including genes belonging to known ncRNA gene families as well as unique ncRNA classes that have not previously been described (Aspegren et al. 2004; Hinas et al. 2006; Hinas and Söderbom 2007; Kuhlmann et al. 2005).

Rivas et al. (Rivas and Eddy 2000) first applied nucleotide compositional contrasts in the average nucleotide composition to identify ncRNA genes within an AT-rich genome. Their work was extended by Klein *et al.* (Klein et al. 2002) and Schattner (Schattner 2002), who separately used the nucleotide contrast approach to identify ten and 19 ncRNA candidates, respectively in the AT-rich (~69%) genome of *Methanococcus jannaschii*. In total, 23 unique candidates were identified and Klein *et al.* showed expression from five out of their ten candidates. Another organism with an extremely AT-rich (80%) genome is the malarial parasite *Plasmodium falciparum* (Gardner et al.

2002). Its genome was recently searched for ncRNA genes using a compositional contrast approach combined with a search for conserved regions among three closely related *Plasmodium* species (Upadhyay et al. 2005). This study reported 18 candidates for new ncRNA genes, of which six could be detected by Northern blot analysis (Upadhyay et al. 2005).

The compositional contrast approach to predict ncRNA genes requires methods to locate compositionally deviating segments within a sequence, to measure the statistical significance of the compositional deviations, and to define the boundaries of the deviating segments. Schattner and Upadhyay *et al.* both employed a sliding window to segment genome sequences based on average GC-content (Schattner 2002; Upadhyay et al. 2005). The sliding window approach is attractive since it is relatively uncomplicated and large genomes can be rapidly searched. However, it suffers from complementary problems in the lack of a well-grounded theoretical basis for statistics and difficulty to achieve precision in defining boundaries of the segments. Instead of a sliding window approach, Klein and co-workers used a two-state hidden Markov model (HMM) to identify GC-rich segments. The HMM provides a much more flexible framework with natural solutions for boundary detection. However, a relatively simple and straightforward HMM approach, as used for instance by Klein *et al.*, carries implicit *a priori* assumptions about the length distributions and genomic dispersions of ncRNA genes. Also, the statistical significance of candidates obtained from the sliding window or HMM approaches cannot be calculated analytically but are rather computed by parametric simulation.

An alternative approach to segment an input sequence is to transform it into a sequence of scores according to a set of rules. Within this sequence, disjoint segments are determined so that the partial sum of the segment is maximized. Karlin and Altschul have proposed a number of suitable rules and have studied distributions of maximum segment scores, enabling the calculation of the statistical significance of such segments (Karlin and Altschul 1990; Karlin et al. 1990). This has applications to various problems, most notably sequence comparison, for which it has been implemented in BLAST (Altschul et al. 1990; Altschul et al. 1997). It has also been used for other purposes such as identification of trans-membrane domains and predicting replication origins in viral genomes. (Chew et al. 2007; Karlin and Altschul 1990). In a recent work, Csűrös developed an algorithm for optimally combining disjoint segments and used Karlin-Altschul theory for controlling the statistical significance of combined segments (Csűrös 2004).

It is well known that structural features of ncRNA often are important for the function of ncRNAs (Eddy 2002). The best algorithms in use today for RNA secondary structure prediction by free energy minimization rely on the empirical nearest-neighbor model (reviewed in (Mathews 2006)). Overlapping dinucleotide frequencies have proven important to statistically model RNA sequences (Clote et al. 2005; Workman and Krogh 1999), at least in part because the nearest neighbor model includes base-pair stacking increment parameters within stems (Xia et al. 1998). We therefore reasoned that the compositional contrast approach might be more powerful for the ncRNA gene-finding

application if it were generalized to consider overlapping dinucleotides. Fortunately, Karlin and Dembo derived results that generalize Karlin-Altschul statistics to Markov-dependent sequences (Karlin and Dembo 1992). However, to our knowledge their results have not been applied before to this or any other biological problem.

In this work we implemented a partial sum process using empirical log-likelihood ratio scoring schemes and Karlin-Altschul and Karlin-Dembo statistics and applied it to *de novo* ncRNA gene prediction in the AT-rich genome of the protist *D. discoideum*. With this method we recovered 94 percent of previously identified ncRNA genes. The predictions were subsequently filtered using biologically significant criteria such as the repeated occurrence of similar sequences within the *D. discoideum* genome or the nearby presence of a Dictyostelium upstream sequences element (DUSE) previously implicated to be associated with ncRNA genes (Aspegren et al. 2004; Hinas et al. 2006; Hinas and Söderbom 2007). Finally, we verified expression of six out of ten novel ncRNA genes or gene families by Northern blot analysis. We believe that our improved contrast approach to *de novo* single-genome prediction of ncRNA genes will be a useful search tool also for other genomes where a sufficient compositional contrast exists and will be particularly powerful in combination with other search criteria.

Results

Nucleotide and dinucleotide composition of ncRNA genes and the background search space in *D. discoideum*

Genes expressing ncRNAs often have higher GC-content than their surrounding DNA context, especially in organisms with AT-rich genomes (Rivas and Eddy 2000). This observation encouraged us to analyze the nucleotide compositions of ncRNA genes known or confidently predicted to be expressed in *D. discoideum* (Table 1 and (Hinas and Söderbom 2007)), including 379 tRNA genes that we predicted using the ARAGORN tRNA prediction software (Laslett and Canbäck 2004). The average GC-content of the ncRNA genes ranges from around 33% for the RNase P RNA gene up to approximately 51% for the tRNA genes (Table 1). The GC-content of the ncRNA genes was considerably higher than the genomic average of 22% (Eichinger et al. 2005). Furthermore, the vast majority of these were found in non-repetitive intergenic regions where the average GC-content was only about 14%. Thus, the contrast in average GC-content between the ncRNA genes and the non-repetitive, intergenic or intronic genome of *D. discoideum* (the “background genome”, see Methods) is approximately 2- to 3.5-fold.

In order to investigate the possibility to predict the orientation of the candidate ncRNA genes, we analyzed the strand-symmetry for the background genome and the known ncRNA genes. While the background genome is nearly strand-symmetric, there is some asymmetry in ncRNA strands (Supplemental Data S1, Supplemental Data S2). However,

when this was applied to determine the coding strand for the candidate ncRNA genes, the predictions were not reliable (data not shown). Hence, we did not discriminate between the strands during our search.

Considering dinucleotides, G-tests of independence (Zar 1999) strongly reject independence of consecutive nucleotides ($p \ll 0.001$, not surprisingly given the large sample sizes). Furthermore and in agreement with previous studies (Karlin and Burge 1995; Karlin et al. 1998), the background had stronger departures from those expected by independent mononucleotides with strong tendencies towards repeats (i.e. “CC” and “GG” are more than 2.4 times more prevalent than expected). For *D. discoideum* ncRNA genes, TC is most overrepresented among dinucleotides and exceeds the expected frequency 1.3-fold, while AC is most underrepresented at about 71% of expectation. Comparing conditional dinucleotide probabilities by ratios of those in target to background revealed substantial contrasts, such as a greater than six-fold higher conditional probability of G following an A in targets. This indicates that using higher-order Markov-dependent partial sum processes could yield better results in ncRNA gene-finding as detailed below. Based on these results, we were encouraged to pursue a screen for ncRNA genes based on contrasts of nucleotide and dinucleotide contents.

Strategy to search for ncRNA genes based on deviations in nucleotide and dinucleotide compositions

In order to search for new ncRNA genes in available genome sequences, we implemented a segmentation approach based on maximal partial sums in a sequence of scores. To

evaluate the statistical significance of segments, we used Karlin-Altschul statistics for independent nucleotides (“the M0 model”) and the Karlin-Dembo extension for nucleotides with first-order Markov dependence (“the M1 model”) (Karlin and Dembo 1992; Karlin et al. 1990). The only parameters required are the rules for scoring a nucleotide sequence and the probabilities of assigning each score. Karlin and Altschul proposed a scoring scheme using log-likelihood ratios of frequencies of word occurrence events in a “target” distribution over those in a “null” or “background” distribution (Karlin and Altschul 1990). We derived such a scheme by calculating target nucleotide word frequencies from a non-redundant set of known or inferred *D. discoideum* ncRNA genes (frequencies were calculated from both the forward and reverse complement sequence). Background frequencies were derived from the *D. discoideum* genome where complex repeats, exons and known or inferred ncRNA genes had been removed (see Methods). The relative entropies of the scoring matrices in the two models indicate that conditional dinucleotides have more than three times the information for detecting target ncRNA genes than mononucleotides (for M0, the entropy is 0.19 bits per nucleotide and for M1, the entropy is 0.68 bits per nucleotide). If nucleotides were independent, the expected information increase of the conditional dinucleotide matrix would be merely twice that of the mononucleotide matrix.

Performance evaluation

To compare the performances of the M0- and M1 models, we generated a simulated search space of the same sizes as the *D. discoideum* genome using a Markov chain of 5th order (i.e. the probability of generating each individual nucleotide depends on the

previous five nucleotides). We estimated parameters for the Markov chain from the repeat masked *D. discoideum* intergenic regions as detailed in Methods. The background was masked and known or confidently predicted ncRNA genes were inserted at the same locations as in the natural data. We could then compare the performance of the M0- and M1 models by analyzing the sensitivity and specificity in rediscovering the inserted ncRNA genes. We scored each strand independently but considered a prediction as a true positive if it overlapped a known ncRNA, regardless of the strand and the length of the overlap. Predictions that overlapped each other but on opposite strands were combined into one prediction using the outermost endpoints. To evaluate the performance we calculated the area under the curve (AUC) when sensitivity was plotted against the false positive rate (1-specificity) (ROC-plot, Fig. 1A). The sensitivity was weighted according to ncRNA family as described in Methods. With AUC 0.84 and 0.87 for the M0- and M1 model, respectively, we conclude that the M1 model performs better on the artificial data. We performed a two-fold cross-validation (see Methods). When using the artificial background data, the mean of the AUC for M0 and M1 was 0.76 and 0.80, respectively (standard deviations 0.04 and 0.02, respectively). When using background data sampled from the *D. discoideum* genome, the mean AUC was 0.21 and 0.22 for M0 and M1, respectively (standard deviation 0.02 and 0.005).

Search for novel ncRNA genes in *D. discoideum*

Approximately 12 Mb of the *D. discoideum* genome that remained after masking of exons and repeats were searched in both directions using both the M0- and M1 model (see Methods). Fig. 1B shows the sensitivity weighted by ncRNA family plotted against

the ranking of corresponding hits for both models. It can be seen that the models perform similarly but the M1 model has a higher specificity as the number of hits increase. This prompted us to proceed with the M1 model predictions to identify candidates that could potentially correspond to novel ncRNA genes and we found an E-value cutoff of 0.1 (dashed, vertical line in Fig. 1B) to be a good compromise between sensitivity and the number of hits. The search strategy is outlined in Fig. 2.

This search strategy resulted in 2296 candidate ncRNA gene-containing regions. We refer to these as drp for “*Dictyostelium* ncRNA predictions”. Of the 2296 predicted regions, 396 overlapped known or confidently predicted *D. discoideum* ncRNA genes. The shortest overlap between a prediction and a known ncRNA gene observed was 35 nt or 45% of the length of the ncRNA. Ten predictions each overlap two closely located tRNA genes and one prediction contains two C/D-box snoRNAs. Hence, in total 94% or 407 of the 435 known or confidently predicted ncRNA genes (not overlapped by repeats or protein coding genes) were successfully predicted by our method (Table 1), including 370 of 379 tRNAs, 17 of 18 snRNAs, the H/ACA-box snoRNA, RNase P and MRP RNAs, SRP RNAs and seven of 18 C/D-box snoRNAs. In the artificial background, the same sensitivity is reached at a false positive rate of approximately 25%, although this number is most likely an underestimation due to the ideal nature of this dataset.

In order to identify the predicted regions most likely to contain new ncRNA genes, we applied different filtering procedures (Fig. 2). The first filter was based on the observation that known ncRNA genes often occur in multigene families, i.e. similar or

identical RNA genes are often present in multiple copies throughout the genome. This organization also applies to the ncRNA genes in *D. discoideum* (Hinas and Söderbom 2007). Therefore, the candidates were clustered according to sequence similarity. From here on we will use “cluster” to describe two or more sequences that appear related by sequence similarity (not implying location in the genome). We use “family” to refer to clusters that we believe represent homologous ncRNA genes. Our analysis generated 131 distinct clusters with two to 159 members (Supplemental Fig. S1). In total, 892 of the 2296 putative ncRNAs could be clustered by sequence similarity. More than half of the clusters contained only two members and roughly one third of the two-member clusters (25 of 73) can be explained by the 750 kb genomic duplication previously described (Eichinger et al. 2005). Known *D. discoideum* ncRNAs were present in 21 of the identified clusters (Table 1 and Supplemental Fig. S1).

As a second filter, we searched for the presence of a ncRNA-gene-associated upstream sequence element within 150 nucleotides flanking the 2296 predicted gene-containing regions. This eight nucleotide putative promoter element, DUSE, is located approximately 63 nucleotides upstream of the transcription start site of the majority of the recently identified and expressed ncRNA genes in *D. discoideum* (Aspegren et al. 2004; Hinas et al. 2006; Hinas and Söderbom 2007). The DUSE sequence was found in association with 116 predicted ncRNA genes. Of these, 28 correspond to genes already known or confidently predicted to express ncRNAs. We also combined the two filters to search for sequences within clusters that are associated with DUSE, which resulted in 54

members dispersed in 25 clusters. Nine of these clusters contain known or confidently predicted ncRNA genes.

Genomic distribution

The predictions are in general evenly distributed across the genome, with on average one prediction every 15 kb. The known or confidently predicted ncRNAs occur on average once every 80 kb but this number varies substantially between chromosomes with chromosome six having the highest ncRNA density of one every 40 kb and chromosomes four and five the lowest with one every 120 kb. We and others have previously reported that ncRNA genes in *D. discoideum* frequently occur in pairs or triplets of highly similar copies (Aspegren et al. 2004; Eichinger et al. 2005; Hinas et al. 2006). These copies are located in all possible orientations, i.e. tandem, divergent, and convergent, usually within 25 kb of each other but often even closer, e.g. some of the snRNA genes are separated only by a few hundred base pairs. Even though the mechanism is unclear, this organization suggests recent duplication events (Eichinger et al. 2005). We found that 34 of our identified clusters of candidate ncRNA genes show a similar organization in the genome, i.e. two members are closer than 25 kb apart. In addition, eight of the 12 members of cluster *drf9* (drf for “Dictyostelium ncRNA family”) seem to co-localize with ncRNA genes belonging to the Class I RNA genes, a family of short, abundantly expressed ncRNAs only known in *D. discoideum* (see Supplemental Table S1) (Aspegren et al. 2004; Hinas and Söderbom 2007). Despite the close association, it seems unlikely that these genes are co-transcribed since both the Class I genes and the *drf9* genes carry

their own upstream DUSE sequence and in many cases they are located on opposite strands.

Experimental validation

In order to validate the search algorithm employed to screen for ncRNA genes, we investigated gene expression by Northern blot analyses on total RNA isolated from growing *D. discoideum* cells. For each tested candidate ncRNA gene containing region, we assayed both strands using strand- and target-specific oligonucleotide probes. We chose to analyze the expression of a set of randomly selected candidate RNA genes that could be clustered together by sequence similarity. The probes were designed to recognize the maximum number of members within the same cluster, i.e. they were complementary to the most conserved regions.

For the expressed family *drf17* (see below), the hybridization oligonucleotide matched five additional sites in the genome. However, upon closer inspection these sequences were found to be similar to *drf17* but did not meet the significance threshold of the search algorithm. Similarly, the probe for *drf9* (see below) matched one additional site in the genome, which was positioned within a predicted candidate. Due to the stringency of the clustering algorithm, the candidate failed to cluster together with *drf9*.

The search generated families containing sequences of variable length. In the majority of cases, the member sequences within each family have a core motif where most of the nucleotides align (target for the hybridization probes). However, other motifs were

sometimes shifted in relation to each other due to insertions and gaps of non-similar sequences.

Eight of the 131 identified clusters were analyzed for expression and four of these were transcribed into small RNAs (*drf115*, *drf17*, *drf22*, and *drf9*) (Fig. 3 and Table 2). The putative multigene family *drf27* yielded signals indicating RNAs substantially longer than the predictions and expressed from both strands while *drf15* showed a hybridization signal of about 160 nt from one strand. However, the signal disappeared at high stringency washes. Hence, these two families were considered to be false positives. The additional two predicted multigene families, *drf32* and *drf11* failed to give rise to any detectable hybridization signals.

Properties of novel expressed ncRNA multigene families

The family *drf17* contains eight similar member sequences, approximately 220 nucleotides in length according to Northern blot results (Fig. 3A). One of the predictions spans two highly similar copies located on opposite strands. Hence, each copy was treated as individual candidates, (*drf17_5* and *drf17_6*) (Table 2). We could detect only a faint signal by Northern blot analysis on RNA extracted from growing *D. discoideum* cells. However, we readily detected hybridization signals when we investigated expression during development (i.e. 16h slugs, and 24h fruiting bodies) (Fig. 3A). The developmental expression was only analyzed for this family. The multigene family *drf22* is made up of six sequences and the major signal detected on Northern blot appears to be about 160 nucleotides (Fig. 3A). In addition, a heterogeneous population of longer RNAs

possibly derived from the different members of this family was visible. An alternative explanation is that the transcripts from this family are of the same size but subjected to polyadenylation as reported for *D. discoideum* spliceosomal RNAs (snRNAs) (Hinas et al. 2006).

The family *drf115* consists of two sequences located merely 80 bp apart. The member recognized by the hybridization probe has a DUSE motif located 65 nt upstream of the predicted start site. Northern blot detects a major ~100 nucleotide RNA transcript. In addition, a less strong hybridization signal is visible corresponding to an RNA sized ~130 nucleotides (Fig. 3A).

The family *drf9* contains 12 member sequences (Fig. 3C and Table 2) and was subjected to a more detailed analysis. Two slightly different oligonucleotide probes were used to verify expression (9-1 and 9-2 in Fig. 3A). Hybridization signals indicated expressed RNA of approximately 240-270 nucleotides. In order to determine the 5' and 3' ends of the RNA, *drf9* RNA was subjected to Rapid Amplification of cDNA Ends (RACE) analyses (data not shown). The sequence of the 5'RACE product (only one sequenced) could be derived from three members of *drf9* with identical sequences to the determined 5' end (Fig. 3C). In addition, two sequenced 3' RACE products determined the 3' ends as well as most of the full-length RNA sequence of two additional member RNAs (Fig. 3C). Hence, three different RNAs from within this family have been directly verified. The sizes of these RNAs are estimated to 265 and 267 nucleotides by 3'end determination and 265-268 nucleotides by the 5'RACE. These estimations are based on the following

observations: i) the RNA sequences are similar; ii) the 5' and 3' ends are positioned in short stretches of fully conserved regions; iii) the major size as determined by Northern blot analysis is ~270 nucleotides (one of the hybridization primers was also used for 5'RACE); iv) the presence of a conserved sequence element (DUSE) predicted/verified to be located 63 nucleotides in front of nine of the 12 members and preceding all the RNAs where the 5' and 3' ends have been confirmed and v) the genomic region downstream of the experimentally verified 3' ends consists of a stretch of thymidine residues. This run of T's is conserved among the members of *drf9* and may constitute a termination signal. The RNA polymerase III transcribed U6 snRNA genes in other organisms have been shown to use stretches of T's as terminators (Will and Lührmann) and the same organization is present in *D. discoideum* (Hinas et al. 2006).

The secondary structure of *drf9* was predicted by the RNAalifold software (Hofacker et al. 2002) using the five sequences which correspond to the RACE-determined members (Fig. 3C). The predicted consensus secondary structure contains several stem structures where compensatory base pair changes are present, corroborating the structure (Fig. 3B).

In summary, four out of eight of the computationally predicted clusters of small RNA genes were expressed. One of the expressed ncRNA families showed very low expression in growing cells, which increased during development. This indicates that transcription of some of the predicted genes may have escaped detection since the majority of the analyses were performed on RNA from growing cells only. One predicted ncRNA gene

family, *drf9*, was investigated in more detail and the presence and partial sequences of at least three member RNAs was verified.

Increased stringency by searching for conserved upstream elements

As an alternative approach to improve the quality of the predictions we searched for the conserved upstream sequence element, DUSE (see above), in association with the ncRNA candidates. Since the exact site of transcription initiation could not reliably be predicted for the majority of the RNAs, a region of 150 nucleotides flanking the predicted genes was investigated. The DUSE element was found to be associated with 116 candidate RNA genes. Of these, two candidate ncRNA genes (notably not belonging to any apparent clusters) were randomly chosen for further analysis and both genes were demonstrated to be transcribed. The transcripts correspond to the DNA strand containing the DUSE (Fig. 4) and are located to the cytoplasm (data not shown). *drd38* (drd for “Dictyostelium ncRNA with DUSE”) was detected as a major transcript of ~90 nucleotides which was found to correspond to the recently identified selenocysteine tRNA (Fig. 4A) (Shrimali et al. 2005). The major RNA derived from *drd3* is ~155 nucleotides while a less abundant species of ~180 nucleotides was also detected (Fig. 4C). The secondary structures of the RNAs were predicted by RNAfold (Hofacker et al. 1994) based on the sizes estimated by Northern blot analysis and the assumption that the start of transcription is located 63 nucleotides from the DUSE sequence (Fig. 4). In addition, the majority of the members of *drf9* (whose expression was verified – see above) as well as the verified member of *drf115* also have an upstream DUSE sequence.

These results show that combining our algorithm for ncRNA gene prediction with searches for the conserved DUSE element generates an additional level of confidence. All six of the analyzed genes (including one member of *drf115* and the three individual genes within *drf9* for which expression was verified by RACE) that were preceded by a DUSE were transcribed.

Discussion

The options for *de novo* ncRNA gene prediction using a single genome are limited (Meyer 2007) and there is a need for improved methods. In the protist *D. discoideum*, we observed that the known ncRNA genes have a substantially higher GC-content and markedly different dinucleotide composition than the background genome sequence. This prompted us to develop a computational search strategy to identify compositionally deviating regions that could contain novel ncRNA genes. For this purpose we implemented a search algorithm based on maximal segment scores among partial sums and the Karlin and Dembo extension for Markov dependence to the Karlin and Altschul theory on the statistical significance. We used a scoring scheme based on nucleotide word frequencies derived from background and target sequences. There are several major advantages with this method, e.g. statistical significance of scores of ncRNA gene candidates can be readily calculated and there are few nuisance parameters that need to be estimated.

First-order Markov dependence, i.e. a model where the probability of observing a particular nucleotide depends on the identity of the nucleotide at the previous position, is motivated from a biological point of view. Nearest-neighbor interactions are known to be important for RNA structural prediction as a consequence of e.g. stacking interactions (for review, see (Turner and Sugimoto 1988)). We implemented algorithms based on this theory using both a model that assumes independent nucleotide positions (the M0 model) and a model that assumes nucleotide positions with first-order Markov dependence (the M1 model). Mononucleotide frequencies as well as conditional dinucleotide frequencies

in *D. discoideum* are nearly strand-symmetric for the background genome but show asymmetries to some degree for ncRNA genes, suggesting that a compositional contrast approach could predict the orientation of a potential ncRNA gene. However, when this possibility was explored we did not find the predictions accurate and we chose not to take this factor into account when we searched for novel candidates.

When we used the M1 model to search the genome sequence of *D. discoideum* with an expect threshold of 0.1, 2296 regions were identified that had a deviating nucleotide composition compared to the surrounding genome. The great majority of the verified and confidently predicted ncRNA genes were included in this set. The reason for failure to identify all the C/D-box snoRNAs and the Class I and Class II RNAs may be due to their relatively short length and their lack of extensive intramolecular secondary structure. Although our implementation does not explicitly impose length requirements on candidate regions, short sequences may not accumulate an aggregate score large enough to meet the significance threshold. It should be noted that the snRNA and nine tRNAs that escaped detection by the search method, were masked due to overlap with exons or repeat sequences.

It is well known that many ncRNAs occur in families with several copies of similar RNA genes spread throughout the genome. This also applies to *D. discoideum* ncRNA genes (Aspegren et al. 2004; Eichinger et al. 2005; Hinas et al. 2006). This observation was used to refine the search method by identifying possible families of new ncRNA genes. Of these, we confidently verified expression of four out of eight tested. To further

improve the stringency of the method, we searched the flanking sequences of all the predicted ncRNA genes for a conserved sequence element, DUSE, known to be associated with *D. discoideum* ncRNA genes (Aspegren et al. 2004; Hinas et al. 2006; Hinas and Söderbom 2007). Two of these, notably where the sequences did not belong to a cluster but were unique in the genome, were tested for expression and transcripts for both genes were detected. Interestingly, one of these, *drd38*, was found to correspond to the recently reported selenocysteine tRNA gene (Shrimali et al. 2005) that neither tRNAscan-SE nor ARAGORN identified. The observation that this unusual tRNA gene is preceded by the upstream sequence element is intriguing since this motif is not normally associated with tRNA genes. In fact, besides the selenocysteine tRNA gene, only a leucine tRNA gene seems to have the element at the expected upstream distance (our unpublished observation). It is tempting to speculate that the presence of DUSE in front of these tRNA genes could be linked to a specific regulatory mechanism, which is distinct from the canonical tRNA genes. In addition, nine out of 12 members of one of the predicted ncRNA clusters have DUSE in close proximity and RACE experiments identified three separate sequences within this family of which all were derived from sequences preceded by DUSE.

We have assessed the source of the performance improvement of the M1 model in *D. discoideum* and the potential of this extension for this purpose in other organisms, i.e. *Caenorhabditis elegans*, *Giardia lamblia*, *P. falciparum* and *M. jannaschii*, with a variety of background genomic nucleotide compositions. When we searched the AT-rich genome of *M. jannaschii*, previously studied by Klein *et al.* and Schattner (Klein et al. 2002;

Schattner 2002), we were able to detect all known as well as all predicted previously reported ncRNAs (Supplemental Data S3 and Supplemental Fig. S2).

It appears that the contrast in mononucleotide content is the largest contributor to the detection of contrasting regions even when the M1 model is used. In addition, the marginal advantage of M1 in *D. discoideum* lies apparently in an unusual conditional dinucleotide signature in the background rather than in the targets, reflecting the fact that if we are detecting any RNA-specific dinucleotide signature it is swamped by much more distinct signatures of the background. This was largely consistent when compared across the organisms we studied, however we noticed an interesting consistency in certain dinucleotides, *i.e.* AA/TT and AG/CT. These have roughly the same observed/expected dinucleotide ratio for RNAs from almost all organisms we examined but lack the corresponding consistency in background frequency ratios (Supplemental Fig. S3). Nonetheless, our investigation clearly shows that higher-order models are potentially beneficial for detecting compositionally contrasting regions of genomes in order to detect non-coding RNA genes.

In conclusion, we present a computational search for new ncRNA genes based on a screen for genomic regions with unusual dinucleotide composition. The extracted sequences are subsequently filtered with clustering of similar predicted regions and/or search for a conserved upstream sequence element. We applied this approach on the *D. discoideum* genome and discovered, with high confidence, new ncRNA gene candidates. Out of a total of ten tested ncRNA candidate genes, the expression of six was verified

using stringent criteria. Furthermore, when the presence of the upstream motif was used as the sole search criterion, all of the predicted genes were found to be expressed. Hence, our combined approach of using partial sum processes together with our filtering criteria gives a useful method to detect new ncRNA genes. In addition, other ways of filtering the candidates, such as clustering according to secondary structure (Will et al. 2007), could also prove important. We anticipate that our search method will be useful for the analysis of other organisms, particularly those for which genomes of closely related species have not yet been sequenced and the use of higher-order relationships can give a marginal but valuable improvement.

Methods

Sequence data

The 2.5 release of the *D. discoideum* genome (Eichinger et al. 2005) along with annotations of sequence features was downloaded from the dictyBase web site, <http://dictybase.org/> (Chisholm et al. 2006) in June 2006. RepeatMasker (Smit et al. 1996-2004) was used to mask abundant complex repeats commonly found in the intergenic sequences of *D. discoideum* using a custom repeat library, (Glöckner et al. 2001) obtained from the Genome Sequencing Centre Jena website at <http://genome.imb-jena.de/dictyostelium/> in June 2006. In addition, a 250 kb sequence at the beginning of chromosome 1, which is believed to function as a centromeric region (Eichinger et al. 2005), was also masked from the analyses due to its abundance of repeat sequences. Finally, exons in the dictyBase annotation were masked. Masking means that these

nucleotides were changed to the letter X, leaving chromosome sequences otherwise intact. The approximately 12 Mb of sequence from the six *D. discoideum* chromosomes remaining after the masking procedure is called the “Background genome” in which contrasting regions were searched and from which, with additional masking of known or inferred ncRNA genes, background frequencies were estimated.

Search algorithm

A nucleotide sequence L of length n is transformed into a score sequence A according to a set of scoring rules. For the independent nucleotide case (the M0 model), each nucleotide is scored according to the set of scores $\{s_\alpha\}$, $\alpha \in \{A, C, G, T\}$ and similarly, for the Markov-dependent nucleotide case (the M1 model), each overlapping dinucleotide is scored based on a set of dinucleotide scores $\{s_{\alpha\beta}\}$, $\alpha\beta \in \{AA, AC, \dots, TT\}$. A partial sum

S_m is defined as $S_0 = 0$, $S_m = \sum_{j=1}^m \Lambda_j$, $m \leq n$ where Λ_j corresponds to the score at

position j in the score sequence. Then the greatest aggregate score

$M(n) = \max_{0 \leq k \leq l \leq n} (S_l - S_k)$, is the aggregate score for the subsequence of L starting at the

nucleotide with index k and ending at the nucleotide with index l . We consider all subsequences having an aggregate score exceeding a chosen significance level as ncRNA gene candidates. For the M0 model, it is assumed that at least one of the scores s_α is positive and the probability of observing that score, p_α , is non-zero. Furthermore, the expected score, $\sum_\alpha p_\alpha s_\alpha$, is assumed to be negative. For the M1 model, the corresponding assumptions are that for each nucleotide α there exists nucleotides β and γ such that

$s_{\alpha\beta} > 0$ and $s_{\alpha\gamma} < 0$ and the associated probabilities $p_{\alpha\beta}$, $p_{\alpha\gamma}$ are non-zero. In addition, the expected score $\sum_{\alpha\beta} \pi_{\alpha} p_{\alpha\beta} s_{\alpha\beta}$, where π is the stationary frequency vector for the dinucleotide probabilities $p_{\alpha\beta}$, is assumed to be negative. Provided that the conditions above hold, it has been shown that for large n , the probability of obtaining an aggregate score greater than $x + \frac{\ln n}{\lambda^*}$ can be approximated by $1 - \exp(-K^* e^{-\lambda^* x})$ (Karlin and Altschul 1990; Karlin and Brendel 1992). K^* and λ^* are parameters that depend on the probabilities p_{α} and scores s_{α} (or $p_{\alpha\beta}$ and $s_{\alpha\beta}$ for the M1 case) and $\frac{\ln n}{\lambda^*}$ is a scaling parameter. Specifically, for the M0 case, λ^* can be found as the unique positive solution to the equation $\sum_{\alpha} p_{\alpha} * \exp(\lambda s_{\alpha}) = 1$. Formulas for calculating K^* are detailed in the appendix of (Karlin and Altschul 1990) and step-by-step instructions on calculating λ^* and K^* for the M1 case can be found in (Karlin and Dembo 1992) (p. 137-139). We have implemented routines in Java for calculating these parameters. The routines include solving eigen systems and estimation of infinite sums using iteration by matrix recursions, for which we used the JAMA matrix package (v1.0.2) developed by MathWorks and NIST (<http://math.nist.gov/javanumerics/jama/>).

Hence, for a desired significance level, the corresponding score threshold can be determined and conversely, the statistical significance of a high scoring sequence can be calculated. Specifically, it has been shown that the number of subsequences having a

score greater than or equal to a score S can be approximated from a Poisson distribution according to $E = nK^* \exp(-\lambda^* S)$ (Karlin and Altschul 1990).

Both strands of the background genome data were searched at an expect value of 0.1, but for most analyses in this paper we took a region as predicted if either strand for that region was significant at the given threshold.

Scoring rules

Scores are calculated as log-odds ratios between target and background nucleotide frequencies. We derived target frequencies from inferred genome sequence parts (“genes”) corresponding to a non-redundant subset of the known ncRNA genes (Table 1).

The non-redundant subset was created by removing sequences sharing more than 90% identity with any other sequence in the set, using the *weight* program from the *squid* v1.9g package, written by Sean Eddy and available at

<http://selab.janelia.org/software.html>. Let Ω denote the set of different RNA families present in the training set, such that $\Omega = \{tRNA, rRNA, snRNA, \textit{Class I RNA}, \textit{Class II RNA}, \textit{RNA D2}, \textit{MRP RNA}, \textit{RNaseP RNA}, \textit{C/D-box snoRNA}, \textit{H/ACA snoRNA}, \textit{SRP RNA}\}$. Let

n_{α}^i and $n_{\alpha\beta}^i$ be the mononucleotide and dinucleotide counts of nucleotides

$\alpha \in \{A, C, G, T\}$ and dinucleotides $\alpha\beta \in \{AA, AC, \dots, TT\}$ in the i^{th} family, $i \in \Omega$,

respectively and let M^i and D^i denote the total number of mononucleotides and

dinucleotides in the i^{th} family when only nucleotides and dinucleotides involving A, C, G

or T are counted. Then the mononucleotide target frequency equals $f_{\alpha}^t = \frac{1}{\|\Omega\|} \sum_{i \in \Omega} \frac{n_{\alpha}^i}{M^i}$ and

the dinucleotide target frequency equals $f_{\alpha\beta}^t = \frac{1}{\|\Omega\|} \sum_{i \in \Omega} \frac{n_{\alpha\beta}^i}{D^i}$.

In calculations of the background frequencies we used the previously defined background genome in which all known or confidently predicted ncRNA genes have additionally been masked. Using the notation above, let n_{α} and $n_{\alpha\beta}$ be the count of nucleotide α and dinucleotide $\alpha\beta$, respectively and let M and D be the total number of mononucleotides and dinucleotides, respectively. The background mononucleotide and dinucleotide

background frequencies were calculated according to $f_{\alpha}^b = \frac{n_{\alpha}}{M}$ and $f_{\alpha\beta}^b = \frac{n_{\alpha\beta}}{D}$,

respectively. The score associated with the nucleotide α was then calculated as the log-

likelihood ratio $s_{\alpha} = c \log_2 \frac{f_{\alpha}^t}{f_{\alpha}^b}$, where c is a scaling constant introduced for practical

purposes (we used a value of $c=10$) and the scores were rounded to the closest integer.

Similarly, the score for the dinucleotide $\alpha\beta$ was calculated as $s_{\alpha\beta} = c \log_2 \frac{q_{\alpha\beta}}{p_{\alpha\beta}}$, where

$p_{\alpha\beta}$ and $q_{\alpha\beta}$ are the probabilities of making a transition from state α to state β in the

Markov chain and can be calculated according to $p_{\alpha\beta} = \frac{f_{\alpha\beta}^b}{f_{\alpha}^b}$ for the background and

$q_{\alpha\beta} = \frac{f_{\alpha\beta}^t}{f_{\alpha}^t}$ for the target. In addition, a set of custom scores were introduced to

accommodate masked segments according to $s_{\gamma} = s_{\alpha\gamma} = s_{\gamma\beta} = s_{\gamma\delta} = -\infty$, $\alpha, \beta \neq X$,

$\gamma, \delta = X$ (the letter X indicates a masked nucleotide). Ambiguous or unrecognized nucleotides encountered were scored according to $s_\gamma = \min_\alpha s_{\alpha\gamma}$, $s_{\alpha\gamma} = \min_\beta s_{\alpha\beta}$, $s_{\gamma\beta} = \min_\alpha s_{\alpha\beta}$, $s_{\gamma\delta} = \min_{\alpha\beta} s_{\alpha\beta}$, $\alpha, \beta \in \{A, C, G, T\}$, $\gamma, \delta \in \{R, Y, M, K, S, W, H, B, V, D, N\}$ (the set of IUPAC ambiguous nucleotide symbols).

Model Validation and Selection

ROC-plots were calculated by running the M0 and M1 models on simulated datasets for which known or confidently predicted ncRNA genes were inserted at their natural positions into background genome data generated according to a 5th order Markov model of the background genome data (estimated from the same data as the background model for the score matrices). The same pattern of feature masking as the natural data was retained in the simulated data. A true-positive was counted if one or more significant regions overlapped it with at least one nucleotide. A true negative was counted if no significant region overlapped it. If a significant region overlapped a true negative and simultaneously overlapped a true positive on one side, the true negative was still counted as such. Sensitivity was calculated as $Sn = \frac{TP}{TP + FN}$, where TP and FN denotes number

of true positives and false negatives, respectively. Specificity was calculated as

$Sp = \frac{TN}{TN + FP}$, where TN and FP denotes number of true negatives and false positives,

respectively. Since we only had a limited amount of non-redundant target sequences, we chose to do a cross-validation analysis with two folds to have true positive examples representing diverse classes of ncRNAs in each fold. We divided the non-redundant

target set of known ncRNAs in two halves and the background set was divided in the same way. The first half of the ncRNA set was used together with one half of the background set for calculating the scores and the other half was randomly inserted into the other half of the background set for validation. This procedure was repeated for the two halves. We used both the artificial background described above and real data consisting of the non-repetitive, intergenic regions from *D. discoideum*, assuming that no unknown true ncRNA genes overlapped those. ROC-plot analyses were performed and the mean and standard deviation of the AUC was calculated.

Upstream sequence elements

In order to determine if a DUSE was associated with a candidate gene, we searched for the presence of sequence elements ACCCATAA or TCCCAHAA (H = A, C or T) within the 150 nucleotides flanking the candidate.

Sequence clustering

Sequence clusters were built by single-linkage clustering using the blastclust v2.2.9 software. Neighbors in a cluster were required to share at least 80% sequence identity covering a minimum of 50% of one of the sequence's length. The word size was set to 11 and the *dust* low complexity filter was used.

Bioinformatic tools

Version 1.1 of the tRNA gene prediction software ARAGORN was used for prediction of tRNA genes (Laslett and Canbäck 2004). The RepeatMasker software, version open-3.1.2 (Smit et al. 1996-2004), was used for masking of interspersed repeats. Clustering of

candidates was performed using the blastclust v2.2.9 software, available in the blast package at <http://www.ncbi.nlm.nih.gov/BLAST/download.shtml>. Secondary structure predictions were performed using RNAfold and RNAalifold in version 1.6.2 of the Vienna RNA Package, available at <http://www.tbi.univie.ac.at/~ivo/RNA/> (Hofacker et al. 2002; Hofacker et al. 1994). The predicted secondary structures were manually edited using the RNAviz v2.0 editor, available at <http://rnaviz.sourceforge.net/> (De Rijk et al. 2003). Multiple sequence alignments were constructed using ClustalW v1.83 (Thompson et al. 1994).

Java routines and source code

We have bundled the java routines used for calculation of score matrices and statistical parameters as well as identification of high scoring regions of contrasting composition into a package named SIGRS (Single Genome ncRna Search). Class files and source code are available for download at http://www.icm.uu.se/molbio/virtanen/publications/larsson_2008/supplemental/S4

Sequence data

Experimentally verified ncRNAs have been submitted to dictyBase (<http://dictybase.org>). An annotation file compatible with the dictyBase genome browser that can be used to visualize all predictions in their genomic context has been compiled (Supplemental Data S5) and a fasta file with nucleotide sequences of the predictions is available (Supplemental Data S6).

Oligonucleotides

Sequences of the DNA oligonucleotides (Invitrogen) used in this study can be found in Supplemental Table S2. Sequences of the RNA oligonucleotides used for 5' and 3' RACE (Dharmacon) were described in (Aspegren et al. 2004).

Growth conditions

D. discoideum strain AX4 (Knecht et al. 1986) was grown axenically in HL5 medium and developed on nitrocellulose membranes (Sussman 1987).

Expression analysis

Total RNA was extracted from axenically growing cells and cells developed for 16 and 24 hours, respectively, by the TRIzol method (Invitrogen). For Northern blot analysis, 20 µg RNA was separated on an 8% polyacrylamide/7M urea/1xTBE gel and electroblotted to a Hybond N⁺ membrane (GE Healthcare) together with radiolabeled markers pUC19/MspI (Fermentas) and RNA decade marker (Ambion). The membranes were hybridized over night with P³²-end-labeled DNA oligonucleotides at 42°C in Church buffer (7% SDS, 0.5 M NaPO₄ pH 7.2, 1 mM EDTA and 1% BSA) and subsequently washed twice for 5 min in 2xSSC/0.1%SDS, twice for 10 min in 1xSSC/0.1% SDS and twice for 5 min in 0.5xSSC/0.1%SDS at 42°C, followed by washing twice for 5 min in 0.5xSSC/0.1%SDS at 50°C. Hybridization signals were detected by a PhosphorImager (Molecular Dynamics). For determination of RNA termini, 5' and 3'RACE were carried out as described previously (Aspegren et al. 2004). Only probes that were fully

complementary to candidate RNAs were regarded to be recognized by the oligonucleotide.

Deposited sequences

Primary sequences for ncRNAs *drf9_3* and *drf9_7* determined by 3'RACE analysis have been deposited -in the GenBank database. Accession numbers are [GenBank: EF551319] for *drf9_3* and [GenBank: EF551320] for *drf9_7*.

Acknowledgements

We thank Lotta Avesson and Björn Garpefjord for their contributions on the experimental validation of candidate ncRNAs and Johan Reimegård for valuable suggestions.

dictyBase (<http://dictybase.org/>) is also gratefully acknowledged. This work was supported by grants from the European Community (FOSRAK, EC005120) to F.S., from Wallenberg Consortium North Foundation, The Swedish Research Council to Uppsala RNA Research Center in the form of Linnéstöd and Swedish Foundation for Strategic Research to L.K.

Figure legends

Figure 1. Performance of the M0- and M1 models. Solid red line and red circles indicate performance of M0 model. Dotted black line and black circles indicate performance of M1 model. The dashed lines indicate the $E=0.1$ threshold for the M0 (red line) and M1 (black line) models, respectively. **(A)** An artificial background was generated (see Methods). The performance of the M0 and M1 models was measured in terms of sensitivity and specificity. The sensitivity was weighted so that each ncRNA family has an equal contribution to the total sensitivity. The weighted sensitivity is plotted against 1-specificity. **(B)** Performance on real data. The weighted sensitivity in detecting known or confidently predicted ncRNAs in the *D. discoideum* genome is plotted against the rank of the corresponding hits.

Figure 2. Flow chart and clustering of sequences. **(A)** The six chromosomes of the *D. discoideum* where exons and complex repeats had been masked were searched for high scoring segments with an expect value of at most 0.1. Overlapping high scoring segments residing on opposite strands were combined. The predictions were filtered for presence in multi-copy families and/or the occurrence of an upstream sequence element (DUSE) within 150 nucleotides flanking the prediction. Pie charts represent the number and character of sequences that were used for the operation. The size of the pie chart is proportional to the total number of sequences. Yellow and the associated numbers indicate novel predictions and blue and associated numbers denotes known or confidently predicted ncRNA genes. After combining the strands, the number of known or

confidently predicted ncRNA genes increased since genes with an overlapping prediction on the template strand counts as detected after the combination step.

Figure 3. Analysis of expressed ncRNA families. **(A)** Northern blot analysis of families hybridized with sense and antisense (as) probes. *drf17* was analyzed for developmental regulation on RNA extracted at 0h (growing cells), 16h and 24h. Two slightly different probes were used in the analysis of *drf9* (9-1 and 9-2). **(B)** Consensus secondary structure, predicted by RNAalifold, of the five members matching sequences determined by RACE analysis. Compensatory base pair changes are circled. **(C)** Alignment of the *drf9* member sequences including upstream and downstream flanking sequences. Green box beneath sequences indicates location of the DUSE sequence present in nine sequences. Red and blue boxes beneath sequences depict sequences determined by 5' and 3'RACE analysis, respectively. Candidate gene names marked by red and blue boxes indicate members matching the sequences determined by RACE analysis. Nucleotides complementary to the hybridization probes are indicated in bold letters between the vertical bars. Start and end point nucleotides of predictions are marked by black boxes in the alignment. The predictions for *drf9_3*, *drf9_4* and *drf9_8* extend beyond the alignment for additional 124, 226 and 111 nucleotides, respectively. The alifold consensus structure is indicated in dot-bracket notation above the alignment. The notation has been slightly adopted to accommodate the full family alignment. Compensatory base pair changes in the RACE-determined sequences are indicated by green, blue or yellow boxes for a change to a GC, AT or GT base pair, respectively. A red box indicates a change that disrupts base pairing.

Figure 4. Analysis of two candidates with an upstream DUSE sequence. Northern blot analyses were performed for each candidate in sense and anti-sense (as) direction as indicated. RNA secondary structure predictions were obtained from RNAfold using the sequence length deduced from the Northern blot analyses and assuming a transcription start site located 63 nucleotides downstream of the DUSE. **(A)** and **(B)** show Northern blots and secondary structure predictions for the *drd38* and *drd3* candidates, respectively.

Tables

Table 1. Known *D. discoideum* ncRNAs

RNA	Average			Detected ^a				Reference
	N	GC%	Length	C	D	C+D	None	
tRNA ^b	379	51	79	362	1	4	3	(Eichinger et al. 2005) and this study
C/D-box								
snoRNA ^c	18	35	81		1		6	(Aspegren et al. 2004)
Class I ^c	17	37	58		1			(Aspegren et al. 2004)
Class II	2	40	61					(Aspegren et al. 2004)
H/ACA	1	35	146		1			(Aspegren et al. 2004)
MRP RNA	1	36	366				1	(Piccinelli et al. 2005)
RNase P RNA	1	33	369				1	(Piccinelli et al. 2005),(Marquez et al. 2005)
rRNA ^c	4	43	1348		Not searched			(Sugang et al. 2003)
snRNA ^b	18	37	176	6	2	9		(Aspegren et al. 2004), (Hinas et al. 2006)
SRP RNA	2	46	281			2		(Aspegren et al. 2004)
D2/U3	7	34	210			7		(Wise and Weiner 1981)

Number of genes, average GC-content and length for each RNA are shown. ^aNumber of genes identified by the M1 model screen and fulfilling either of the filtering criteria: clustering (C), DUSE (D), combined clustering and DUSE (C+D) or not identified by any filter (None). ^bNine tRNA genes and one snRNA gene are masked from the search since they overlap e.g. exons. ^cIn the natural data, the ribosomal RNA genes are located on an extrachromosomal DNA element not included in the search and were consequently never detected. The search in the artificial data readily identified the ribosomal RNA genes. Similarly, one C/D-box snoRNA gene and one Class I RNA gene are located on unassembled contigs that were not included in the search.

Table 2. Candidate ncRNAs analyzed for expression.

Name	Northern blot ^a			DUSE	M1 prediction ^b			Rank ^c	
	Detected	Length	Strand		Length	GC%	Expect	M0	M1
<i>drd38</i>	+	90	→	x	90	50	5.2E-09	933	734
<i>drd3</i>	+	154	→	x	198	39	7.0E-09	532	751
<i>drf115 1</i>	+	100	→	x	119	41	3.1E-03	864	1777
<i>drf115 2</i>	ND				143	38	1.0E-03	864	1655
<i>drf17 1</i>	+	220	←		142	44	3.9E-03	661	1811
<i>drf17 2</i>			←		75	48	3.9E-02	1445	2203
<i>drf17 5</i>			→		544	36	6.8E-39	46	88
<i>drf17 6</i>			←		216	46	6.8E-39	46	88
<i>drf17 7</i>			←		75	48	4.5E-02	1475	2225
<i>drf17 3</i>	ND				104	42	2.2E-02	1236	1994
<i>drf17 4</i>					337	45	3.8E-25	82	130
<i>drf17 8</i>					360	33	5.2E-13	618	431
<i>drf27 1</i>	-				150	47	1.6E-10	398	587
<i>drf27 2</i>					88	48	1.7E-03	1088	1611
<i>drf27 4</i>					136	48	2.8E-08	450	839
<i>drf27 3</i>	ND				150	47	4.3E-09	373	723
<i>drf27 5</i>					176	41	2.8E-08	578	839
<i>drf22 1</i>	+	160	→		179	44	2.7E-14	379	358
<i>drf22 2</i>			→		184	44	4.1E-14	331	370
<i>drf22 3</i>			←		302	36	1.3E-13	331	397
<i>drf22 6</i>			←		296	37	3.3E-15	331	314
<i>drf22 4</i>	ND				84	46	1.1E-05	1347	1234
<i>drf22 5</i>					85	40	8.5E-02	2353	2349
<i>drf9 3</i>	+	240-270	←	x	385	35	2.1E-16	245	278
<i>drf9 4</i>			→	x	504	32	6.4E-13	204	439
<i>drf9 5</i>			→	x	273	35	2.1E-12	649	461
<i>drf9 7</i>			→	x	265	34	1.7E-09	910	691
<i>drf9 8</i>			→	x	384	32	3.6E-10	774	624
<i>drf9 9</i>			→	x	243	33	1.1E-08	1214	774
<i>drf9 10</i>			→	x	268	35	5.7E-12	694	487
<i>drf9 11</i>			→	x	268	33	2.6E-12	976	468
<i>drf9 12</i>			←		162	38	1.3E-09	1106	680
<i>drf9 13</i>			→	x	152	42	1.2E-10	649	581
<i>drf9 1</i>	ND				106	34	6.0E-02	3882	2184
<i>drf9 6</i>					263	32	1.6E-07	1155	939
<i>drf15 2</i>	-				107	36	1.1E-02	2038	1871
<i>drf15 3</i>					137	36	4.5E-03	1410	1834
<i>drf15 4</i>					137	36	4.5E-03	1599	1834
<i>drf15 5</i>					137	36	5.5E-03	1475	1859
<i>drf15 6</i>					127	39	4.6E-05	1309	1337
<i>drf15 7</i>					108	41	1.0E-04	1395	1333
<i>drf15 1</i>					124	34	3.4E-02	2610	2170
<i>drf15 8</i>					131	33	8.4E-03	2799	1918
<i>drf32 1</i>	-			x	146	36	9.7E-03	1712	1940
<i>drf32 2</i>					311	32	7.0E-05	955	1371
<i>drf32 3</i>					352	40	3.8E-25	155	153
<i>drf32 4</i>					231	33	1.4E-06	1177	1079
<i>drf11 7</i>	-				174	38	1.7E-08	910	797
<i>drf11 8</i>					315	37	5.4E-15	313	321
<i>drf11 11</i>					420	35	1.9E-24	224	160
<i>drf11 1</i>					140	38	1.0E-06	1236	986
<i>drf11 2</i>					725	34	6.9E-35	106	98
<i>drf11 3</i>					87	43	6.8E-03	1796	1881
<i>drf11 4</i>					712	34	3.0E-34	108	101
<i>drf11 5</i>					111	34	3.0E-02	3483	2142
<i>drf11 6</i>					94	36	6.9E-02	3197	2201
<i>drf11 9</i>					219	33	6.6E-08	1475	895
<i>drf11 10</i>					92	39	2.5E-03	2185	1659

^aExpression analysis by Northern blot. (+) and (-) indicate presence or absence,

respectively, of relevant hybridization signals. ND indicates that the hybridization probe

was not expected to recognize the candidate. Length and strands are estimated from Northern blot. Strand is relative to the official v2.5 *D. discoideum* genome sequence.

^bCharacteristics of the predicted candidates. Presence of DUSE within 150 nt of flanking sequences is indicated by x. ^cRank measured by standard competition ranking for the M0 and M1 search, respectively.

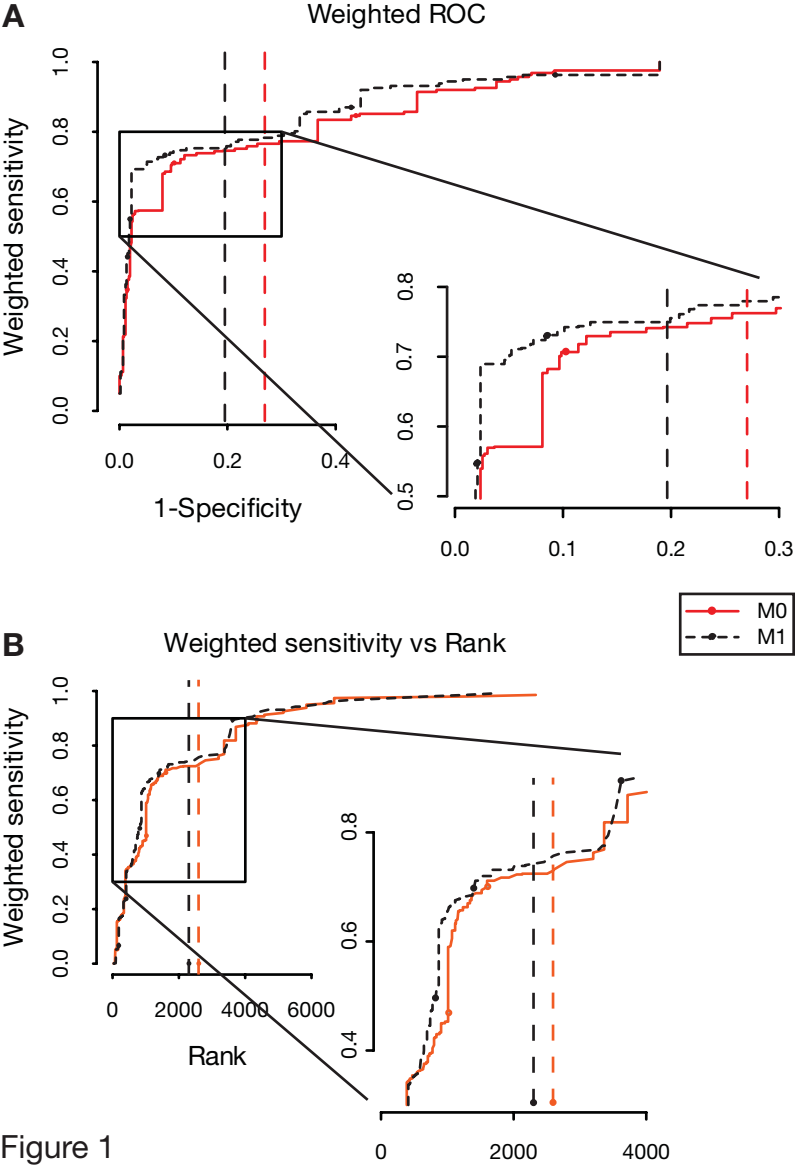
References

- Altschul, S.F., W. Gish, W. Miller, E.W. Myers, and D.J. Lipman. 1990. Basic local alignment search tool. *J Mol Biol* **215**: 403-410.
- Altschul, S.F., T.L. Madden, A.A. Schaffer, J. Zhang, Z. Zhang, W. Miller, and D.J. Lipman. 1997. Gapped BLAST and PSI-BLAST: a new generation of protein database search programs. *Nucleic Acids Res* **25**: 3389-3402.
- Aravin, A. and T. Tuschl. 2005. Identification and characterization of small RNAs involved in RNA silencing. *FEBS Lett* **579**: 5830-5840.
- Aspegren, A., A. Hinas, P. Larsson, A. Larsson, and F. Söderbom. 2004. Novel non-coding RNAs in Dictyostelium discoideum and their expression during development. *Nucl. Acids Res.* **32**: 4646-4656.
- Baldauf, S.L., A.J. Roger, I. Wenk-Siefert, and W.F. Doolittle. 2000. A kingdom-level phylogeny of eukaryotes based on combined protein data. *Science* **290**: 972-977.
- Chew, D.S., M.Y. Leung, and K.P. Choi. 2007. AT Excursion: a new approach to predict replication origins in viral genomes by locating AT-rich regions. *BMC Bioinformatics* **8**: 163.
- Chisholm, R.L., P. Gaudet, E.M. Just, K.E. Pilcher, P. Fey, S.N. Merchant, and W.A. Kibbe. 2006. dictyBase, the model organism database for Dictyostelium discoideum. *Nucl. Acids Res.* **34**: D423-427.
- Clote, P., F. Ferre, E. Kranakis, and D. Krizanc. 2005. Structural RNA has lower folding energy than random RNA of the same dinucleotide frequency. *Rna* **11**: 578-591.
- Csűrös, M. 2004. Maximum-Scoring Segment Sets. *IEEE/ACM Trans. Comput. Biol. Bioinformatics* **1**: 139-150.
- De Rijk, P., J. Wuyts, and R. De Wachter. 2003. RnaViz 2: an improved representation of RNA secondary structure. *Bioinformatics* **19**: 299-300.
- Eddy, S.R. 2002. Computational genomics of noncoding RNA genes. *Cell* **109**: 137-140.
- Edvardsson, S., P.P. Gardner, A.M. Poole, M.D. Hendy, D. Penny, and V. Moulton. 2003. A search for H/ACA snoRNAs in yeast using MFE secondary structure prediction. *Bioinformatics* **19**: 865-873.
- Eichinger, L., J.A. Pachebat, G. Glockner, M.A. Rajandream, R. Sucgang, M. Berriman, J. Song, R. Olsen, K. Szafranski, Q. Xu et al. 2005. The genome of the social amoeba Dictyostelium discoideum. *Nature* **435**: 43.

- Gardner, M.J., N. Hall, E. Fung, O. White, M. Berriman, R.W. Hyman, J.M. Carlton, A. Pain, K.E. Nelson, S. Bowman et al. 2002. Genome sequence of the human malaria parasite *Plasmodium falciparum*. *Nature* **419**: 498-511.
- Glöckner, G., K. Szafranski, T. Winckler, T. Dingermann, M.A. Quail, E. Cox, L. Eichinger, A.A. Noegel, and A. Rosenthal. 2001. The Complex Repeats of *Dictyostelium discoideum*. *Genome Res.* **11**: 585-594.
- Griffiths-Jones, S., S. Moxon, M. Marshall, A. Khanna, S.R. Eddy, and A. Bateman. 2005. Rfam: annotating non-coding RNAs in complete genomes. *Nucl. Acids Res.* **33**: D121-124.
- Havgaard, J.H., R.B. Lyngso, G.D. Stormo, and J. Gorodkin. 2005. Pairwise local structural alignment of RNA sequences with sequence similarity less than 40%. *Bioinformatics* **21**: 1815-1824.
- Hinas, A., P. Larsson, L. Avesson, L.A. Kirsebom, A. Virtanen, and F. Söderbom. 2006. Identification of the Major Spliceosomal RNAs in *Dictyostelium discoideum* Reveals Developmentally Regulated U2 Variants and Polyadenylated snRNAs. *Eukaryotic Cell* **5**: 924-934.
- Hinas, A. and F. Söderbom. 2007. Treasure hunt in an amoeba: non-coding RNAs in *Dictyostelium discoideum*. *Curr Genet* **51**: 141-159.
- Hofacker, I.L., M. Fekete, and P.F. Stadler. 2002. Secondary structure prediction for aligned RNA sequences. *J Mol Biol* **319**: 1059-1066.
- Hofacker, I.L., W. Fontana, P.F. Stadler, L.S. Bonhoeffer, M. Tacker, and P. Schuster. 1994. Fast folding and comparison of RNA secondary structures. *Monatshefte für Chemie* **125**: 167-188.
- Hüttenhofer, A. and J. Vogel. 2006. Experimental approaches to identify non-coding RNAs. *Nucleic Acids Res* **34**: 635-646.
- Karlin, S. and S.F. Altschul. 1990. Methods for Assessing the Statistical Significance of Molecular Sequence Features by Using General Scoring Schemes. *PNAS* **87**: 2264-2268.
- Karlin, S. and V. Brendel. 1992. Chance and statistical significance in protein and DNA sequence analysis. *Science* **257**: 39-49.
- Karlin, S. and C. Burge. 1995. Dinucleotide relative abundance extremes: a genomic signature. *Trends in Genetics* **11**: 283-290.
- Karlin, S., A.M. Campbell, and J. Mrazek. 1998. COMPARATIVE DNA ANALYSIS ACROSS DIVERSE GENOMES. *Annual Review of Genetics* **32**: 185-225.
- Karlin, S. and A. Dembo. 1992. Limit Distributions of Maximal Segmental Score among Markov-Dependent Partial Sums. *Advances in Applied Probability* **24**: 113.
- Karlin, S., A. Dembo, and T. Kawabata. 1990. Statistical Composition of High-Scoring Segments from Molecular Sequences. *The Annals of Statistics* **18**: 571.
- Klein, R.J., Z. Misulovin, and S.R. Eddy. 2002. Noncoding RNA genes identified in AT-rich hyperthermophiles. *PNAS* **99**: 7542-7547.
- Knecht, D.A., S.M. Cohen, W.F. Loomis, and H.F. Lodish. 1986. Developmental regulation of *Dictyostelium discoideum* actin gene fusions carried on low-copy and high-copy transformation vectors. *Mol Cell Biol* **6**: 3973-3983.
- Kuhlmann, M., B.E. Borisova, M. Kaller, P. Larsson, D. Stach, J. Na, L. Eichinger, F. Lyko, V. Ambros, F. Söderbom et al. 2005. Silencing of retrotransposons in *Dictyostelium* by DNA methylation and RNAi. *Nucleic Acids Res* **33**: 6405-6417.

- Kyriakopoulou, C., P. Larsson, L. Liu, J. Schuster, F. Söderbom, L.A. Kirsebom, and A. Virtanen. 2006. U1-like snRNAs lacking complementarity to canonical 5' splice sites. *Rna* **12**: 1603-1611.
- Laslett, D. and B. Canbäck. 2004. ARAGORN, a program to detect tRNA genes and tmRNA genes in nucleotide sequences. *Nucl. Acids Res.* **32**: 11-16.
- Lowe, T.M. and S.R. Eddy. 1997. tRNAscan-SE: a program for improved detection of transfer RNA genes in genomic sequence. *Nucl. Acids Res.* **25**: 955-964.
- Lowe, T.M. and S.R. Eddy. 1999. A Computational Screen for Methylation Guide snoRNAs in Yeast. *Science* **283**: 1168-1171.
- Marquez, S.M., J.K. Harris, S.T. Kelley, J.W. Brown, S.C. Dawson, E.C. Roberts, and N.R. Pace. 2005. Structural implications of novel diversity in eucaryal RNase P RNA. *RNA* **11**: 739-751.
- Mathews, D.H. 2006. Revolutions in RNA secondary structure prediction. *J Mol Biol* **359**: 526-532.
- Meyer, I.M. 2007. A practical guide to the art of RNA gene prediction. *Brief Bioinform*: bbm011.
- Pedersen, J.S., G. Bejerano, A. Siepel, K. Rosenbloom, K. Lindblad-Toh, E.S. Lander, J. Kent, W. Miller, and D. Haussler. 2006. Identification and Classification of Conserved RNA Secondary Structures in the Human Genome. *PLoS Computational Biology* **2**: e33.
- Piccinelli, P., M.A. Rosenblad, and T. Samuelsson. 2005. Identification and analysis of ribonuclease P and MRP RNA in a broad range of eukaryotes. *Nucl. Acids Res.* **33**: 4485-4495.
- Rivas, E. and S. Eddy. 2001. Noncoding RNA gene detection using comparative sequence analysis. *BMC Bioinformatics* **2**: 8.
- Rivas, E. and S.R. Eddy. 2000. Secondary structure alone is generally not statistically significant for the detection of noncoding RNAs. *Bioinformatics* **16**: 583-605.
- Schattner, P. 2002. Searching for RNA genes using base-composition statistics. *Nucl. Acids Res.* **30**: 2076-2082.
- Schattner, P., S. Barberan-Soler, and T.M. Lowe. 2006. A computational screen for mammalian pseudouridylation guide H/ACA RNAs. *RNA* **12**: 15-25.
- Shrimali, R.K., A.V. Lobanov, X.-M. Xu, M. Rao, B.A. Carlson, D.C. Mahadeo, C.A. Parent, V.N. Gladyshev, and D.L. Hatfield. 2005. Selenocysteine tRNA identification in the model organisms Dictyostelium discoideum and Tetrahymena thermophila. *Biochemical and Biophysical Research Communications* **329**: 147.
- Smit, A., R. Hubley, and P. Green. 1996-2004. RepeatMasker Open-3.0. In *RepeatMasker*.
- Sucgang, R., G. Chen, W. Liu, R. Lindsay, J. Lu, D. Muzny, G. Shaulsky, W. Loomis, R. Gibbs, and A. Kuspa. 2003. Sequence and structure of the extrachromosomal palindrome encoding the ribosomal RNA genes in Dictyostelium. *Nucleic Acids Res* **31**: 2361-2368.
- Sussman, M. 1987. Cultivation and synchronous morphogenesis of Dictyostelium under controlled experimental conditions. *Methods Cell Biol* **28**: 9-29.
- Thompson, J.D., D.G. Higgins, and T.J. Gibson. 1994. CLUSTAL W: improving the sensitivity of progressive multiple sequence alignment through sequence

- weighting, position-specific gap penalties and weight matrix choice. *Nucl. Acids Res.* **22**: 4673-4680.
- Turner, D.H. and N. Sugimoto. 1988. RNA structure prediction. *Annu Rev Biophys Chem* **17**: 167-192.
- Upadhyay, R., P. Bawankar, D. Malhotra, and S. Patankar. 2005. A screen for conserved sequences with biased base composition identifies noncoding RNAs in the A-T rich genome of *Plasmodium falciparum*. *Molecular and Biochemical Parasitology* **144**: 149.
- Washietl, S., I.L. Hofacker, and P.F. Stadler. 2005. From The Cover: Fast and reliable prediction of noncoding RNAs. *PNAS* **102**: 2454-2459.
- Will, C.L. and R. Lührmann. 2001. Spliceosomal UsnRNP biogenesis, structure and function. *Curr Opin Cell Biol* **13**: 290-301.
- Will, S., K. Reiche, I.L. Hofacker, P.F. Stadler, and R. Backofen. 2007. Inferring Noncoding RNA Families and Classes by Means of Genome-Scale Structure-Based Clustering. *PLoS Computational Biology* **3**: e65.
- Wise, J.A. and A.M. Weiner. 1981. The small nuclear RNAs of the cellular slime mold *Dictyostelium discoideum*. Isolation and characterization. *J Biol Chem* **256**: 956-963.
- Workman, C. and A. Krogh. 1999. No evidence that mRNAs have lower folding free energies than random sequences with the same dinucleotide distribution. *Nucleic Acids Res* **27**: 4816-4822.
- Xia, T., J. SantaLucia, Jr., M.E. Burkard, R. Kierzek, S.J. Schroeder, X. Jiao, C. Cox, and D.H. Turner. 1998. Thermodynamic parameters for an expanded nearest-neighbor model for formation of RNA duplexes with Watson-Crick base pairs. *Biochemistry* **37**: 14719-14735.
- Zar, J.H. 1999. *Biostatistical Analysis*. Prentice-Hall, New Jersey.



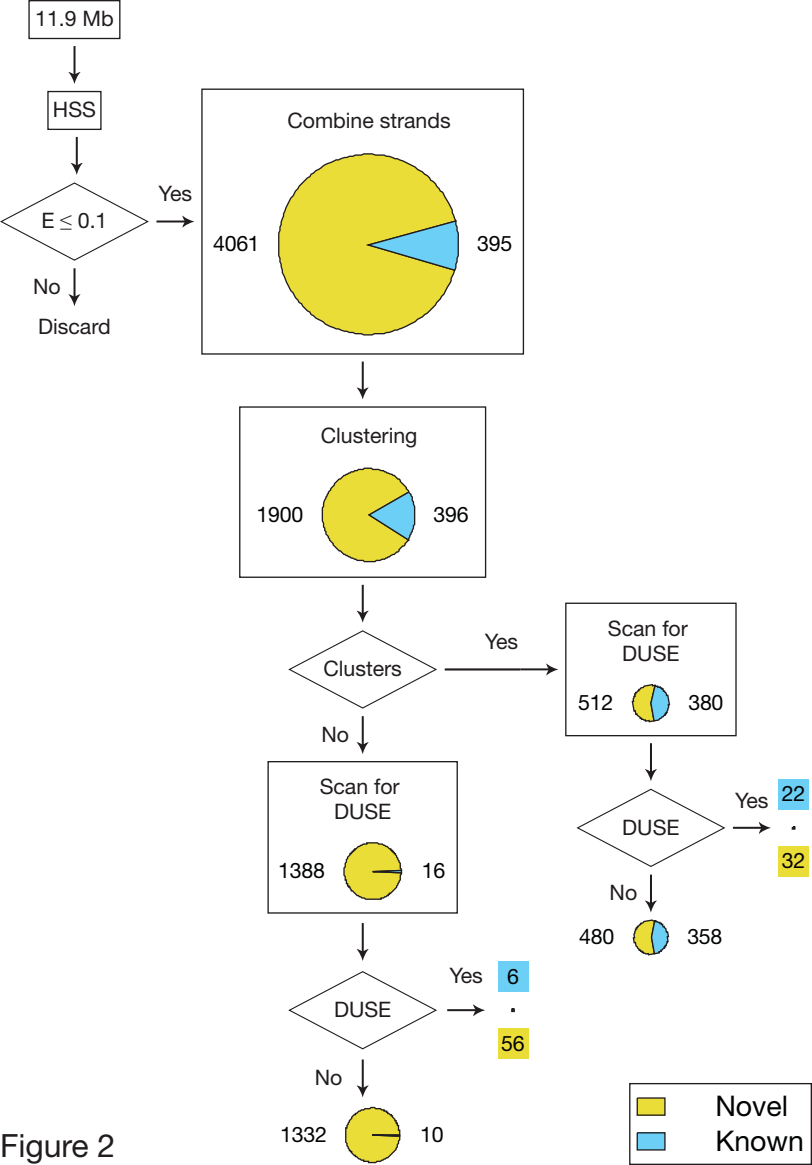


Figure 2

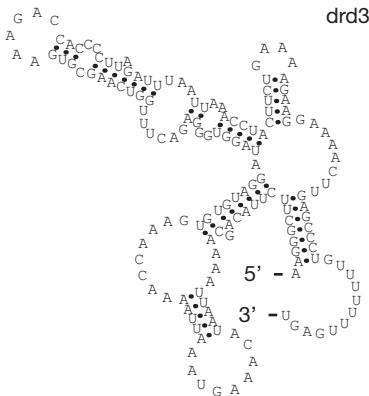
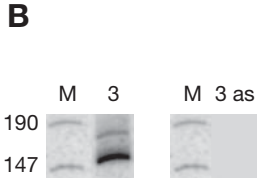
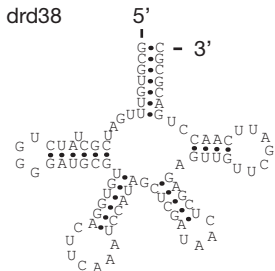
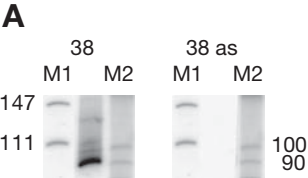
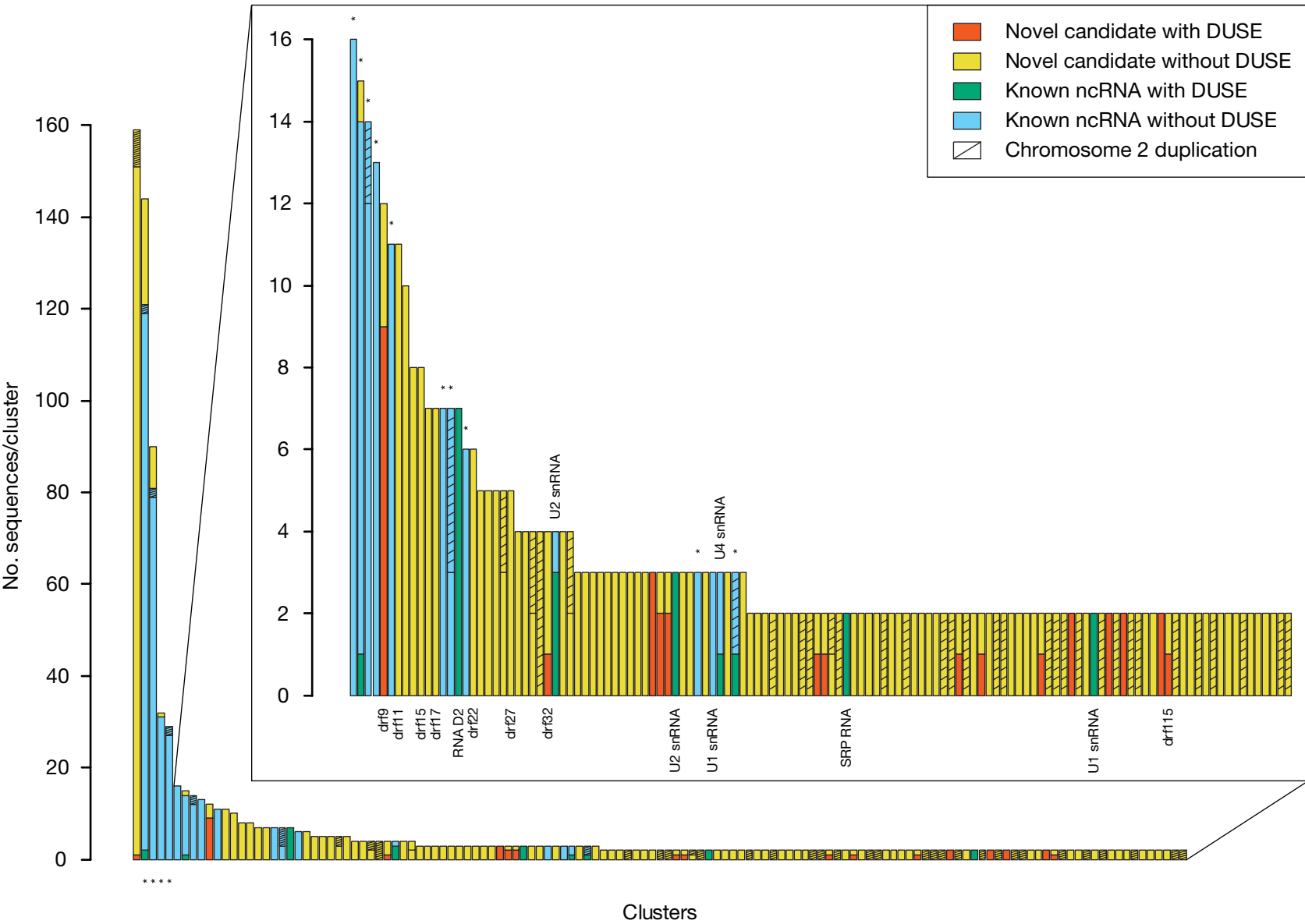


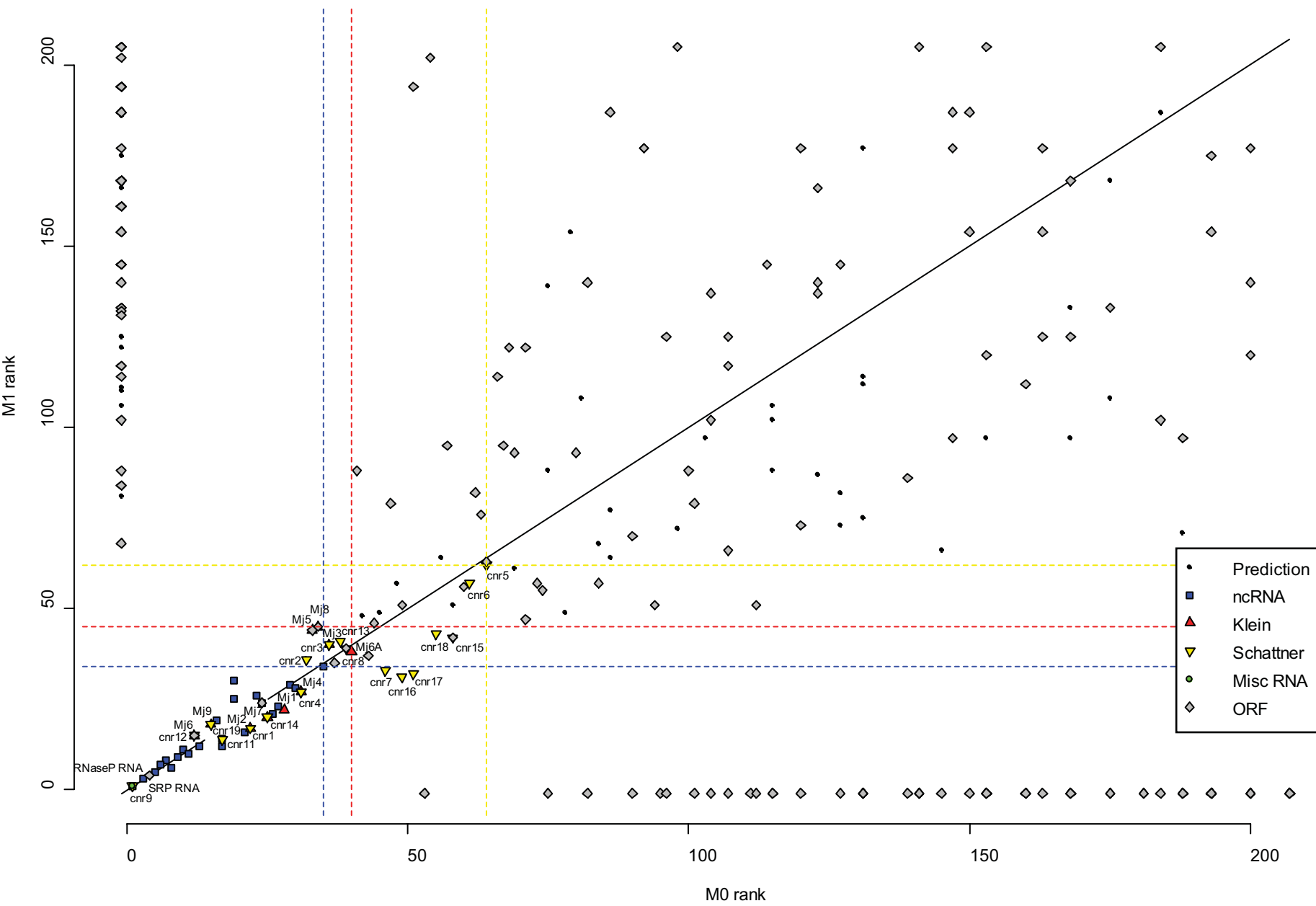
Figure 4

Supplemental Figure S1. Bar graph showing the size distribution and contents of clusters. Red and yellow colors indicate novel predictions with or without an upstream DUSE sequence, respectively. Green and blue bars represent predictions overlapping known or confidently predicted ncRNA genes with or without upstream DUSE sequence, respectively. Hatched bars indicate a sequence residing on the D. discoideum chromosome two duplication. Known or experimentally verified ncRNAs belonging to a cluster are indicated above or below bars. A star above or below a bar indicates presence of tRNAs within that cluster.

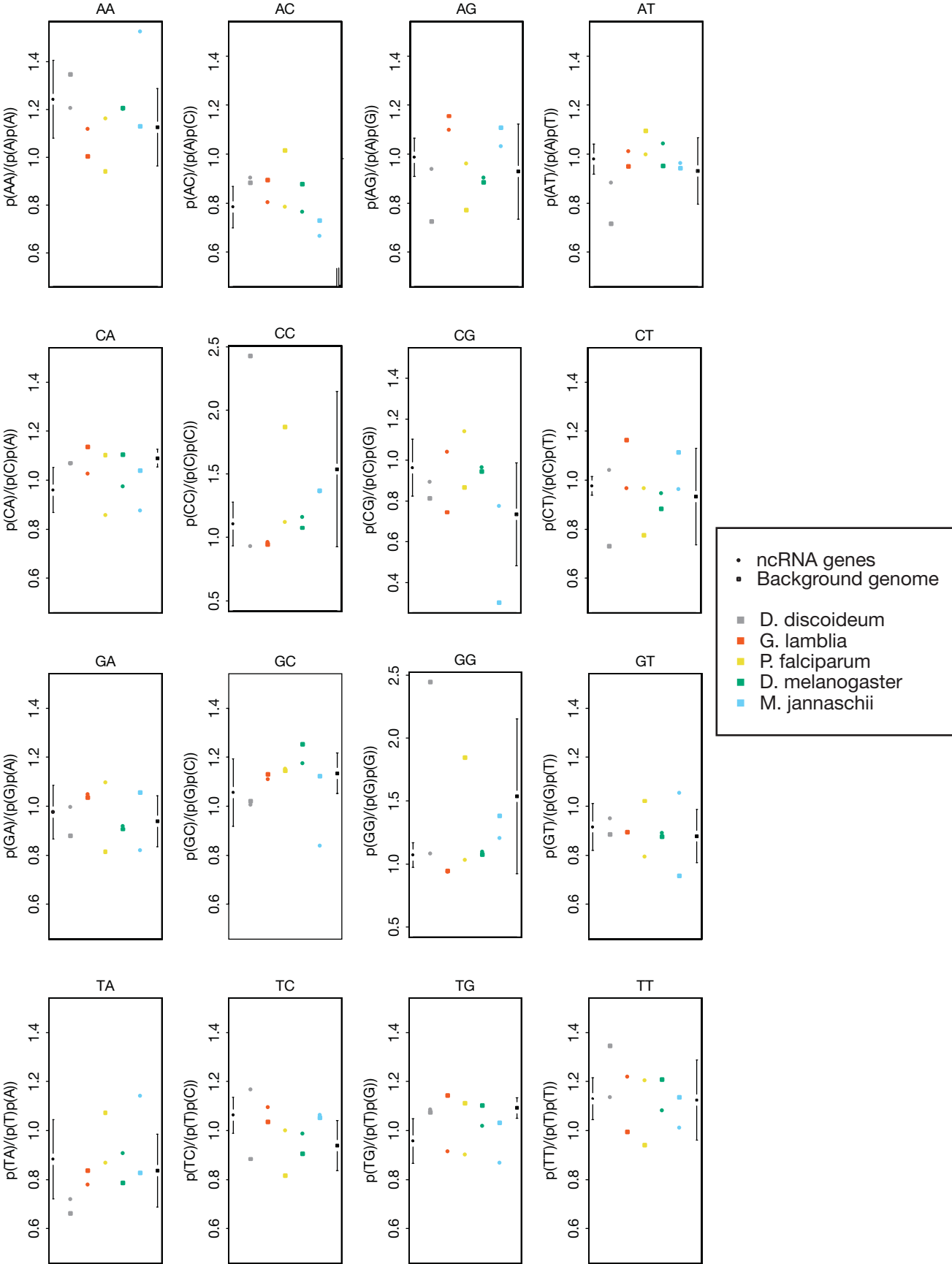


Supplemental Figure S2. Ranks of *M. jannaschii* M1 ncRNA predictions plotted against M0 ncRNA prediction. A solid black dot indicates a prediction. Gray diamonds indicates predictions that overlap known ORFs. Blue squares indicate predictions overlapping a known tRNA or rRNA gene. Red and yellow triangles represent predictions that overlap sequences reported in the works by Klein et al. and Schattner (Klein et al. 2002; Schattner 2002), respectively. Green circles represent SRP RNA and RNase P RNA. Identities of some predictions are indicated. The yellow and red dotted horizontal and vertical lines mark the ranks where all predictions from Schattner and Klein et al. are included. The blue dotted line indicates the ranks where all tRNA and rRNA genes are successfully predicted. A rank of -1 indicates that the region was not predicted by the corresponding model.

M. jannaschii M1 vs M0 prediction ranks



Supplemental Figure S3. Observed/expected dinucleotide ratio for ncRNAs and background genome for *D. discoideum* (Dd), *G. lamblia* (Gl), *P. falciparum* (Pf), *D. melanogaster* (Dm) and *M. jannaschii* (Mj). Filled circles represent observed/expected dinucleotide ratios for ncRNA genes and filled squares for background genome. Mean and standard deviation is indicated. Background genomes were constructed by masking annotated ncRNAs and exons except for *M. jannaschii* where only ncRNAs were masked. Sequence data and annotations were downloaded from the Generic Model Organism Database (<http://gmod.mbl.edu>), PlasmoDB (<http://www.plasmodb.org>), FlyBase (<http://www.flybase.org>) for Gl, Pf and Dm, respectively. Data for Mj was obtained from GenBank (<http://www.ncbi.nlm.nih.gov>).



Supplemental Data S1. Strand asymmetry

Interestingly, mononucleotide composition of the background genome is very nearly strand-symmetric, i.e. $f(A) = f(T) = 0.43$ and $f(G) = f(C) = 0.069$ within a strand (Sueoka's Parity Rule 2 (PR2) (Sueoka 1995)) while ncRNA copy strands (the DNA strand that corresponds to the RNA) are about 14% richer than their template strands in G. This suggests that the compositional contrast approach could in principle not only discriminate ncRNA gene-containing regions but also the template from the copy strand. Like mononucleotides, conditional dinucleotides exhibit substantial strand-asymmetry in the target but not the background, again suggesting a potential for strand-detection via contrast methods. However, when searching the genome using scores that were derived in a strand-dependent manner, the predictions of the coding strand did not turn out to be reliable.

Analyses of background and target ncRNA gene data. Perl script for calculating frequencies is available as Supplemental Data S2.

Observed-expected ratios of the target

	A	C	G	T
A	1.1588	0.7075	1.1121	0.9881
C	0.9748	1.0545	0.9614	1.0159
G	1.0004	0.9572	1.0296	1.0045
T	0.9170	1.2607	0.8330	1.0273

The G-test statistic G for the target = 542.0854 with 9 degrees of freedom.
The probability that sites are independent in the target is 0.0000

Observed-expected ratios of the background

	A	C	G	T
A	1.3476	0.8832	0.7252	0.7153
C	1.0695	2.4276	0.8127	0.7312
G	0.8799	1.0208	2.4461	0.8849
T	0.6617	0.8838	1.0733	1.3451

The G-test statistic G for the background = 1217561.4755 with 9 degrees of freedom.
The probability that sites are independent in the target is 0.0000

Conditional dinucleotide ratios of target over background

	A	C	G	T
A	0.4646	2.5359	6.0323	0.8858
C	0.4925	1.3752	4.6532	0.8910
G	0.6143	2.9687	1.6556	0.7280
T	0.7488	4.5161	3.0527	0.4898

Supplemental Data S3. Search for ncRNA genes in *M. jannaschii*

We compared our search method to the works of Klein *et al.* and Schattner (Klein *et al.* 2002; Schattner 2002) by searching for ncRNA genes in the AT-rich genome of *M. jannaschii*. Sequence data and annotations were downloaded from GenBank. We used the set of 43 tRNA and rRNA genes as target set and the genome with these ncRNA genes masked as background for calculating the scores. We subsequently searched the genome with both the M0 and the M1 model. The predictions are presented in Supplemental Fig. S2 where the ranks of the predictions for the different models are plotted against each other. Using both M0 and M1 we detect all of the tRNA and rRNA genes as well as all the five predicted ncRNAs that Klein *et al.* verified experimentally (some predictions overlap more than one gene). At this rank, the SRP RNA and RNase P RNA reported by Schattner along with 14 of his 19 predictions are also detected. All of the Schattner predictions are included along with 19 and 17 extra predictions at rank 64 and 62 for M0 and M1, respectively. It is interesting to note that the two predictions made by Klein *et al.* that overlap ORFs are the ones that score worst with the M1-model. Specifically, out of 207 and 205 predictions for M0 and M1, 59 predictions do not overlap neither annotated ORFs nor known ncRNAs or predictions by Klein *et al.* or Schattner. Of these, the M1 prediction has a lower rank in 46 cases. Hence, we readily reproduce the results of Klein *et al.* and Schattner using our approach.

Supplemental Table S1. Genomic coordinates and neighboring non-coding RNA genes of experimentally tested candidates. Strand indications are relative to the official v2.5 *D. discoideum* genome release.

Name	M1 prediction				Upstream			Downstream		
	Chr	Start	End	Strand	Gene	Strand	Distance	Gene	Strand	Distance
<i>drd38</i>	4	2475145	2475234	→						
<i>drd3</i>	1	1893172	1893369	→						
<i>drf115_1</i>	5	693177	693295	→				<i>drf115_2</i>		83
<i>drf115_2</i>	5	693378	693520		<i>drf115_1</i>	→	83			
<i>drf17_1</i>	1	1755846	1755987	←						
<i>drf17_2</i>	1	1756394	1756468	←						
<i>drf17_3</i>	1	4776055	4776158							
<i>drf17_4</i>	2	4773544	4773880							
<i>drf17_5</i>	5	849074	849617	→				<i>drf17_6</i>	←	1
<i>drf17_6</i>	5	849618	849833	←	<i>drf17_5</i>	→	1	<i>drf17_7</i>	←	482
<i>drf17_7</i>	5	850315	850389	←	<i>drf17_6</i>	←	482			
<i>drf17_8</i>	5	1065031	1065390							
<i>drf27_1</i>	2	87872	88021							
<i>drf27_2</i>	2	93620	93707							
<i>drf27_3</i>	2	291848	291997							
<i>drf27_4</i>	2	328690	328825							
<i>drf27_5</i>	4	3250352	3250527							
<i>drf22_1</i>	1	1651963	1652141	→						
<i>drf22_2</i>	3	774812	774995	→						
<i>drf22_3</i>	5	2268228	2268529	←						
<i>drf22_4</i>	5	2632109	2632192					<i>drf22_5</i>		1420
<i>drf22_5</i>	5	2633612	2633696		<i>drf22_4</i>		1420			
<i>drf22_6</i>	6	1904491	1904786	←						
<i>drf9_1</i>	2	4428918	4429023							
<i>drf9_3</i>	2	6866495	6866879							
<i>drf9_4</i>	3	1098356	1098859	←				ClassI Predicted	←	1534
<i>drf9_5</i>	3	1101912	1102184	→	ClassI Predicted	←	265			
<i>drf9_6</i>	4	897136	897398	→						
<i>drf9_7</i>	4	957509	957773		DdR-21	←	632	ClassI Predicted	←	471
<i>drf9_8</i>	4	960593	960976	→	DdR-22	←	255	ClassI Predicted	→	616
<i>drf9_9</i>	4	1662995	1663237	→	DdR-23C	←	249	DdR-24A	→	647
<i>drf9_10</i>	4	1667271	1667538	→	DdR-23B	←	270	ClassI Predicted	→	336
<i>drf9_11</i>	4	2792499	2792766	→	ClassI Predicted	←	261	<i>drf9_12</i>	←	399
<i>drf9_12</i>	4	2793165	2793326	→	<i>drf9_11</i>	→	399	ClassI Predicted	→	203
<i>drf9_13</i>	5	673919	674070	←						
<i>drf15_1</i>	1	1789139	1789262	→						
<i>drf15_2</i>	1	2135534	2135640							
<i>drf15_3</i>	2	89370	89506							
<i>drf15_4</i>	2	95426	95562							
<i>drf15_5</i>	2	326232	326368							
<i>drf15_6</i>	2	6241777	6241903							
<i>drf15_7</i>	3	2833138	2833245							
<i>drf15_8</i>	5	312466	312596							
<i>drf32_1</i>	1	2014925	2015070							
<i>drf32_2</i>	1	2858071	2858381							
<i>drf32_3</i>	2	8182027	8182378							
<i>drf32_4</i>	6	718936	719166							
<i>drf11_1</i>	2	294846	294985							
<i>drf11_2</i>	2	316677	317401							
<i>drf11_3</i>	2	323800	323886							
<i>drf11_4</i>	3	2830876	2831587							
<i>drf11_5</i>	3	2863427	2863537							
<i>drf11_6</i>	3	3665682	3665775							
<i>drf11_7</i>	3	4921473	4921646							
<i>drf11_8</i>	4	4276362	4276676							
<i>drf11_9</i>	5	2611116	2611334							
<i>drf11_10</i>	5	3909598	3909689							
<i>drf11_11</i>	5	4281170	4281589							

Supplemental Table S2. Oligonucleotides used for Northern blots, 5'RACE, and 3'RACE. Column "RNA": The RNA or group of RNA for which the oligos were used. 1) indicates hybridization signal > 500 nucleotides. 2 and 3 depict oligonucleotides also used for 3' and 5'RACE, respectively. For further details, see Material and Methods.

Name	Sequence (5'-3')	RNA	Hybridization
189	GGATTTGAAGTCCACACGCATC	drd38	+
206	GATGCGTGTGGACTTCAAATCC	drd38	-
190	GAACCAACTGTATGGATTTCTCC	drf115	+
284	GGAGAAATCCATACAGTTGGTTC	drf115	-
191	AGTCTCCCACCTATCCTACACAC	drd3	+
285	GTGTGTAGGATAGGTGGGAGACT	drd3	-
223	TACCCACCCATCCAACCCCAT	drf17	+
224	ATGGGGTTGGATGGGTGGGTA	drf17	-
227	AAATCCGGCTATTCGTCTTACAG	drf27	1) -
228	CTGTAAGACGAATAGCCGGATTT	drf27	1) -
229	CGACCGACCACTCCTCATTTAC	drf22	+
230	GTAAATGAGGAGTGGTCGGTCG	drf22	-
193	TTGTATTCTGGAGTCTGAATAGGT	drf9	+
207	ACCTATTCAGACTCCAGAATACAA	drf9	-
231 ²⁾	ACCTATTCAGACACCAGAATACAA	drf9	-
232 ³⁾	TTGTATTCTGGTGTCTGAATAGGT	drf9	+
233	CCCTTCCCATGGACTTCT	drf15	+
234	AGAAGTCCATGGGAAGGG	drf15	-
235	GGGAATCTATCCTGGTCAAATA	drf32	-
236	TATTTGACCAGGATAGATTCCC	drf32	-
239	TATACCCTGCAAGACATCCATTG	drf11	-
240	CAATGGATGTCTTGCAGGGTATA	drf11	-

Supplemental References

- Klein, R.J., Z. Misulovin, and S.R. Eddy. 2002. Noncoding RNA genes identified in AT-rich hyperthermophiles. *PNAS* **99**: 7542-7547.
- Schattner, P. 2002. Searching for RNA genes using base-composition statistics. *Nucl. Acids Res.* **30**: 2076-2082.
- Sueoka, N. 1995. Intrastrand parity rules of DNA base composition and usage biases of synonymous codons. *J Mol Evol* **40**: 318-325.