

Identification and analysis of functional elements in 1% of the human genome by the ENCODE pilot project

The ENCODE Project Consortium*

We report the generation and analysis of functional data from multiple, diverse experiments performed on a targeted 1% of the human genome as part of the pilot phase of the ENCODE Project. These data have been further integrated and augmented by a number of evolutionary and computational analyses. Together, our results advance the collective knowledge about human genome function in several major areas. First, our studies provide convincing evidence that the genome is pervasively transcribed, such that the majority of its bases can be found in primary transcripts, including non-protein-coding transcripts, and those that extensively overlap one another. Second, systematic examination of transcriptional regulation has yielded new understanding about transcription start sites, including their relationship to specific regulatory sequences and features of chromatin accessibility and histone modification. Third, a more sophisticated view of chromatin structure has emerged, including its inter-relationship with DNA replication and transcriptional regulation. Finally, integration of these new sources of information, in particular with respect to mammalian evolution based on inter- and intra-species sequence comparisons, has yielded new mechanistic and evolutionary insights concerning the functional landscape of the human genome. Together, these studies are defining a path for pursuit of a more comprehensive characterization of human genome function.

The human genome is an elegant but cryptic store of information. The roughly three billion bases encode, either directly or indirectly, the instructions for synthesizing nearly all the molecules that form each human cell, tissue and organ. Sequencing the human genome^{1–3} provided highly accurate DNA sequences for each of the 24 chromosomes. However, at present, we have an incomplete understanding of the protein-coding portions of the genome, and markedly less understanding of both non-protein-coding transcripts and genomic elements that temporally and spatially regulate gene expression. To understand the human genome, and by extension the biological processes it orchestrates and the ways in which its defects can give rise to disease, we need a more transparent view of the information it encodes.

The molecular mechanisms by which genomic information directs the synthesis of different biomolecules has been the focus of much of molecular biology research over the last three decades. Previous studies have typically concentrated on individual genes, with the resulting general principles then providing insights into transcription, chromatin remodelling, messenger RNA splicing, DNA replication and numerous other genomic processes. Although many such principles seem valid as additional genes are investigated, they generally have not provided genome-wide insights about biological function.

The first genome-wide analyses that shed light on human genome function made use of observing the actions of evolution. The ever-growing set of vertebrate genome sequences^{4–8} is providing increasing power to reveal the genomic regions that have been most and least acted on by the forces of evolution. However, although these studies convincingly indicate the presence of numerous genomic regions under strong evolutionary constraint, they have less power in identifying the precise bases that are constrained and provide little, if any, insight into why those bases are biologically important. Furthermore, although we have good models for how protein-coding regions

evolve, our present understanding about the evolution of other functional genomic regions is poorly developed. Experimental studies that augment what we learn from evolutionary analyses are key for solidifying our insights regarding genome function.

The Encyclopedia of DNA Elements (ENCODE) Project⁹ aims to provide a more biologically informative representation of the human genome by using high-throughput methods to identify and catalogue the functional elements encoded. In its pilot phase, 35 groups provided more than 200 experimental and computational data sets that examined in unprecedented detail a targeted 29,998 kilobases (kb) of the human genome. These roughly 30 Mb—equivalent to ~1% of the human genome—are sufficiently large and diverse to allow for rigorous pilot testing of multiple experimental and computational methods. These 30 Mb are divided among 44 genomic regions; approximately 15 Mb reside in 14 regions for which there is already substantial biological knowledge, whereas the other 15 Mb reside in 30 regions chosen by a stratified random-sampling method (see <http://www.genome.gov/10506161>). The highlights of our findings to date include:

- The human genome is pervasively transcribed, such that the majority of its bases are associated with at least one primary transcript and many transcripts link distal regions to established protein-coding loci.
- Many novel non-protein-coding transcripts have been identified, with many of these overlapping protein-coding loci and others located in regions of the genome previously thought to be transcriptionally silent.
- Numerous previously unrecognized transcription start sites have been identified, many of which show chromatin structure and sequence-specific protein-binding properties similar to well-understood promoters.

*A list of authors and their affiliations appears at the end of the paper.

- Regulatory sequences that surround transcription start sites are symmetrically distributed, with no bias towards upstream regions.

- Chromatin accessibility and histone modification patterns are highly predictive of both the presence and activity of transcription start sites.

- Distal DNaseI hypersensitive sites have characteristic histone modification patterns that reliably distinguish them from promoters; some of these distal sites show marks consistent with insulator function.

- DNA replication timing is correlated with chromatin structure.

- A total of 5% of the bases in the genome can be confidently identified as being under evolutionary constraint in mammals; for approximately 60% of these constrained bases, there is evidence of function on the basis of the results of the experimental assays performed to date.

- Although there is general overlap between genomic regions identified as functional by experimental assays and those under evolutionary constraint, not all bases within these experimentally defined regions show evidence of constraint.

- Different functional elements vary greatly in their sequence variability across the human population and in their likelihood of residing within a structurally variable region of the genome.

- Surprisingly, many functional elements are seemingly unconstrained across mammalian evolution. This suggests the possibility of a large pool of neutral elements that are biochemically active but provide no specific benefit to the organism. This pool may serve as a 'warehouse' for natural selection, potentially acting as the source of lineage-specific elements and functionally conserved but non-orthologous elements between species.

Below, we first provide an overview of the experimental techniques used for our studies, after which we describe the insights gained from analysing and integrating the generated data sets. We conclude with a perspective of what we have learned to date about this 1% of the

human genome and what we believe the prospects are for a broader and deeper investigation of the functional elements in the human genome. To aid the reader, Box 1 provides a glossary for many of the abbreviations used throughout this paper.

Experimental techniques

Table 1 (expanded in Supplementary Information section 1.1) lists the major experimental techniques used for the studies reported here, relevant acronyms, and references reporting the generated data sets. These data sets reflect over 400 million experimental data points (603 million data points if one includes comparative sequencing bases). In describing the major results and initial conclusions, we seek to distinguish 'biochemical function' from 'biological role'. Biochemical function reflects the direct behaviour of a molecule(s), whereas biological role is used to describe the consequence(s) of this function for the organism. Genome-analysis techniques nearly always focus on biochemical function but not necessarily on biological role. This is because the former is more amenable to large-scale data-generation methods, whereas the latter is more difficult to assay on a large scale.

The ENCODE pilot project aimed to establish redundancy with respect to the findings represented by different data sets. In some instances, this involved the intentional use of different assays that were based on a similar technique, whereas in other situations, different techniques assayed the same biochemical function. Such redundancy has allowed methods to be compared and consensus data sets to be generated, much of which is discussed in companion papers, such as the ChIP-chip platform comparison^{10,11}. All ENCODE data have been released after verification but before this publication, as befits a 'community resource' project (see http://www.wellcome.ac.uk/doc_wtd003208.html). Verification is defined as when the experiment is reproducibly confirmed (see Supplementary Information section 1.2). The main portal for ENCODE data is provided by the UCSC Genome Browser (<http://genome.ucsc.edu/ENCODE/>); this is

Box 1 | Frequently used abbreviations in this paper

AR Ancient repeat: a repeat that was inserted into the early mammalian lineage and has since become dormant; the majority of ancient repeats are thought to be neutrally evolving.

CAGE tag A short sequence from the 5' end of a transcript

CDS Coding sequence: a region of a cDNA or genome that encodes proteins

ChIP-chip Chromatin immunoprecipitation followed by detection of the products using a genomic tiling array

CNV Copy number variants: regions of the genome that have large duplications in some individuals in the human population

CS Constrained sequence: a genomic region associated with evidence of negative selection (that is, rejection of mutations relative to neutral regions)

DHS DNaseI hypersensitive site: a region of the genome showing a sharply different sensitivity to DNaseI compared with its immediate locale

EST Expressed sequence tag: a short sequence of a cDNA indicative of expression at this point

FAIRE Formaldehyde-assisted isolation of regulatory elements: a method to assay open chromatin using formaldehyde crosslinking followed by detection of the products using a genomic tiling array

FDR False discovery rate: a statistical method for setting thresholds on statistical tests to correct for multiple testing

GENCODE Integrated annotation of existing cDNA and protein resources to define transcripts with both manual review and experimental testing procedures

GSC Genome structure correction: a method to adapt statistical tests to make fewer assumptions about the distribution of features on the genome sequence. This provides a conservative correction to standard tests

HMM Hidden Markov model: a machine-learning technique that can establish optimal parameters for a given model to explain the observed data

Indel An insertion or deletion; two sequences often show a length difference within alignments, but it is not always clear whether this reflects a previous insertion or a deletion

PET A short sequence that contains both the 5' and 3' ends of a transcript

RACE Rapid amplification of cDNA ends: a technique for amplifying cDNA sequences between a known internal position in a transcript and its 5' end

RFBR Regulatory factor binding region: a genomic region found by a ChIP-chip assay to be bound by a protein factor

RFBR-Seqsp Regulatory factor binding regions that are from sequence-specific binding factors

RT-PCR Reverse transcriptase polymerase chain reaction: a technique for amplifying a specific region of a transcript

RxFrag Fragment of a RACE reaction: a genomic region found to be present in a RACE product by an unbiased tiling-array assay

SNP Single nucleotide polymorphism: a single base pair change between two individuals in the human population

STAGE Sequence tag analysis of genomic enrichment: a method similar to ChIP-chip for detecting protein factor binding regions but using extensive short sequence determination rather than genomic tiling arrays

SVM Support vector machine: a machine-learning technique that can establish an optimal classifier on the basis of labelled training data

TR50 A measure of replication timing corresponding to the time in the cell cycle when 50% of the cells have replicated their DNA at a specific genomic position

TSS Transcription start site

TxFrag Fragment of a transcript: a genomic region found to be present in a transcript by an unbiased tiling-array assay

Un.TxFrag A TxFrag that is not associated with any other functional annotation

UTR Untranslated region: part of a cDNA either at the 5' or 3' end that does not encode a protein sequence

augmented by multiple other websites (see Supplementary Information section 1.1).

A common feature of genomic analyses is the need to assess the significance of the co-occurrence of features or of other statistical tests. One confounding factor is the heterogeneity of the genome, which can produce uninteresting correlations of variables distributed across the genome. We have developed and used a statistical framework that mitigates many of these hidden correlations by adjusting the appropriate null distribution of the test statistics. We term this correction procedure genome structure correction (GSC) (see Supplementary Information section 1.3).

In the next five sections, we detail the various biological insights of the pilot phase of the ENCODE Project.

Transcription

Overview. RNA transcripts are involved in many cellular functions, either directly as biologically active molecules or indirectly by encoding other active molecules. In the conventional view of genome organization, sets of RNA transcripts (for example, messenger RNAs) are encoded by distinct loci, with each usually dedicated to a single biological role (for example, encoding a specific protein). However, this picture has substantially grown in complexity in recent years¹². Other forms of RNA molecules (such as small nucleolar RNAs and micro (mi)RNAs) are known to exist, and often these are encoded by regions that intercalate with protein-coding genes. These observations are consistent with the well-known discrepancy between the levels of observable mRNAs and large structural RNAs

compared with the total RNA in a cell, suggesting that there are numerous RNA species yet to be classified^{13–15}. In addition, studies of specific loci have indicated the presence of RNA transcripts that have a role in chromatin maintenance and other regulatory control. We sought to assay and analyse transcription comprehensively across the 44 ENCODE regions in an effort to understand the repertoire of encoded RNA molecules.

Transcript maps. We used three methods to identify transcripts emanating from the ENCODE regions: hybridization of RNA (either total or polyA-selected) to unbiased tiling arrays (see Supplementary Information section 2.1), tag sequencing of cap-selected RNA at the 5' or joint 5'/3' ends (see Supplementary Information sections 2.2 and S2.3), and integrated annotation of available complementary DNA and EST sequences involving computational, manual, and experimental approaches¹⁶ (see Supplementary Information section 2.4). We abbreviate the regions identified by unbiased tiling arrays as TxFrag, the cap-selected RNAs as CAGE or PET tags (see Box 1), and the integrated annotation as GENCODE transcripts. When a TxFrag does not overlap a GENCODE annotation, we call it an Un.TxFrag. Validation of these various studies is described in papers reporting these data sets¹⁷ (see Supplementary Information sections 2.1.4 and 2.1.5).

These methods recapitulate previous findings, but provide enhanced resolution owing to the larger number of tissues sampled and the integration of results across the three approaches (see Table 2). To begin with, our studies show that 14.7% of the bases represented in the unbiased tiling arrays are transcribed in at least one tissue sample. Consistent with previous work^{14,15}, many (63%) TxFrag reside outside of GENCODE annotations, both in intronic (40.9%) and intergenic (22.6%) regions. GENCODE annotations are richer than the more-conservative RefSeq or Ensembl annotations, with 2,608 transcripts clustered into 487 loci, leading to an average of 5.4 transcripts per locus. Finally, extensive testing of predicted protein-coding sequences outside of GENCODE annotations was positive in only 2% of cases¹⁶, suggesting that GENCODE annotations cover nearly all protein-coding sequences. The GENCODE annotations are categorized both by likely function (mainly, the presence of an open reading frame) and by classification evidence (for example, transcripts based solely on ESTs are distinguished from other scenarios); this classification is not strongly correlated with expression levels (see Supplementary Information sections 2.4.2 and 2.4.3).

Analyses of more biological samples have allowed a richer description of the transcription specificity (see Fig. 1 and Supplementary Information section 2.5). We found that 40% of TxFrag are present in only one sample, whereas only 2% are present in all samples. Although exon-containing TxFrag are more likely (74%) to be expressed in more than one sample, 45% of unannotated TxFrag are also expressed in multiple samples. GENCODE annotations of separate loci often (42%) overlap with respect to their genomic coordinates, in particular on opposite strands (33% of loci). Further analysis of GENCODE-annotated sequences with respect to the positions of open reading frames revealed that some component exons do not have the expected synonymous versus non-synonymous substitution patterns of protein-coding sequence (see Supplement Information section 2.6) and some have deletions incompatible with

Table 1 | Summary of types of experimental techniques used in ENCODE

Feature class	Experimental technique(s)	Abbreviations	References	Number of experimental data points
Transcription	Tiling array, integrated annotation	TxFrag, RxFrag, GENCODE	117 118 19 119	63,348,656
5' ends of transcripts*	Tag sequencing	PET, CAGE	121 13	864,964
Histone modifications	Tiling array	Histone nomenclature†, RFBR	46	4,401,291
Chromatin‡ structure	QT-PCR, tiling array	DHS, FAIRE	42 43 44 122	15,318,324
Sequence-specific factors	Tiling array, tag sequencing, promoter assays	STAGE, ChIP-Chip, ChIP-PET, RFBR	41,52 11,120 123 81 34,51 124 49 33 40	324,846,018
Replication	Tiling array	TR50	59 75	14,735,740
Computational analysis	Computational methods	CCI, RFBR cluster	80 125 10 16 126 127	NA
Comparative sequence analysis*	Genomic sequencing, multi-sequence alignments, computational analyses	CS	87 86 26	NA
Polymorphisms*	Resequencing, copy number variation	CNV	103 128	NA

* Not all data generated by the ENCODE Project.

† Histone code nomenclature follows the Brno nomenclature as described in ref. 129.

‡ Also contains histone modification.

Table 2 | Bases detected in processed transcripts either as a GENCODE exon, a TxFrag, or as either a GENCODE exon or a TxFrag

	GENCODE exon	TxFrag	Either GENCODE exon or TxFrag
Total detectable transcripts (bases)	1,776,157 (5.9%)	1,369,611 (4.6%)	2,519,280 (8.4%)
Transcripts detected in tiled regions of arrays (bases)	1,447,192 (9.8%)	1,369,611 (9.3%)	2,163,303 (14.7%)

Percentages are of total bases in ENCODE in the first row and bases tiled in arrays in the second row.

protein structure¹⁸. Such exons are on average less expressed (25% versus 87% by RT-PCR; see Supplementary Information section 2.7) than exons involved in more than one transcript (see Supplementary Information section 2.4.3), but when expressed have a tissue distribution comparable to well-established genes.

Critical questions are raised by the presence of a large amount of unannotated transcription with respect to how the corresponding sequences are organized in the genome—do these reflect longer transcripts that include known loci, do they link known loci, or are they completely separate from known loci? We further investigated these issues using both computational and new experimental techniques. **Unannotated transcription.** Consistent with previous findings, the Un.TxFrags did not show evidence of encoding proteins (see Supplementary Information section 2.8). One might expect Un.TxFrags to be linked within transcripts that exhibit coordinated expression and have similar conservation profiles across species. To test this, we clustered Un.TxFrags using two methods. The first method¹⁹ used expression levels in 11 cell lines or conditions, dinucleotide composition, location relative to annotated genes, and evolutionary conservation profiles to cluster TxFrags (both unannotated and annotated). By this method, 14% of Un.TxFrags could be assigned to annotated loci, and 21% could be clustered into 200 novel loci (with an average of ~7 TxFrags per locus). We experimentally examined these novel loci to study the connectivity of transcripts amongst Un.TxFrags and between Un.TxFrags and known exons. Overall, about 40% of the connections (18 out of 46) were validated by RT-PCR. The second clustering method involved analysing a time course (0, 2, 8 and 32 h) of expression changes in human HL60 cells following retinoic-acid stimulation. There is a coordinated program of expression changes from annotated loci, which can be shown by plotting Pearson correlation values of the expression levels of exons inside annotated loci versus unrelated exons (see Supplementary Information section 2.8.2). Similarly, there is coordinated expression of nearby Un.TxFrags, albeit lower, though still significantly different from randomized sets. Both clustering methods indicate that there is coordinated behaviour of many Un.TxFrags, consistent with them residing in connected transcripts.

Transcript connectivity. We used a combination of RACE and tiling arrays²⁰ to investigate the diversity of transcripts emanating from protein-coding loci. Analogous to TxFrags, we refer to transcripts

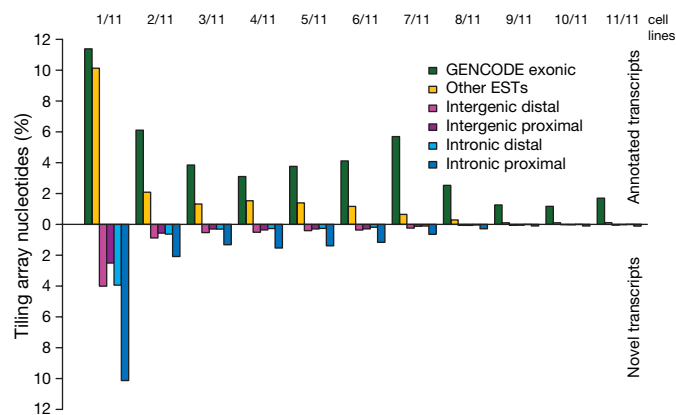


Figure 1 | Annotated and unannotated TxFrags detected in different cell lines. The proportion of different types of transcripts detected in the indicated number of cell lines (from 1/11 at the far left to 11/11 at the far right) is shown. The data for annotated and unannotated TxFrags are indicated separately, and also split into different categories based on GENCODE classification: exonic, intergenic (proximal being within 5 kb of a gene and distal being otherwise), intronic (proximal being within 5 kb of an intron and distal being otherwise), and matching other ESTs not used in the GENCODE annotation (principally because they were unspliced). The y axis indicates the per cent of tiling array nucleotides present in that class for that number of samples (combination of cell lines and tissues).

detected using RACE followed by hybridization to tiling arrays as RxFrags. We performed RACE to examine 399 protein-coding loci (those loci found entirely in ENCODE regions) using RNA derived from 12 tissues, and were able to unambiguously detect 4,573 RxFrags for 359 loci (see Supplementary Information section 2.9). Almost half of these RxFrags (2,324) do not overlap a GENCODE exon, and most (90%) loci have at least one novel RxFrags, which often extends a considerable distance beyond the 5' end of the locus. Figure 2 shows the distribution of distances between these new RACE-detected ends and the previously annotated TSS of each locus. The average distance of the extensions is between 50 kb and 100 kb, with many extensions (>20%) being more than 200 kb. Consistent with the known presence of overlapping genes in the human genome, our findings reveal evidence for an overlapping gene at 224 loci, with transcripts from 180 of these loci (~50% of the RACE-positive loci) appearing to have incorporated at least one exon from an upstream gene.

To characterize further the 5' RxFrags extensions, we performed RT-PCR followed by cloning and sequencing for 550 of the 5' RxFrags (including the 261 longest extensions identified for each locus). The approach of mapping RACE products using microarrays is a combination method previously described and validated in several studies^{14,17,20}. Hybridization of the RT-PCR products to tiling arrays confirmed connectivity in almost 60% of the cases. Sequenced clones confirmed transcript extensions. Longer extensions were harder to clone and sequence, but 5 out of 18 RT-PCR-positive extensions over 100 kb were verified by sequencing (see Supplementary Information section 2.9.7 and ref. 17). The detection of numerous RxFrags extensions coupled with evidence of considerable intronic transcription indicates that protein-coding loci are more transcriptionally complex than previously thought. Instead of the traditional view that many genes have one or more alternative transcripts that code for alternative proteins, our data suggest that a given gene may both encode multiple protein products and produce other transcripts that include sequences from both strands and from neighbouring loci (often without encoding a different protein). Figure 3 illustrates such a case, in which a new fusion transcript is expressed in the small intestine, and consists of at least three coding exons from the *ATP5O* gene and at least two coding exons from the *DONSON*

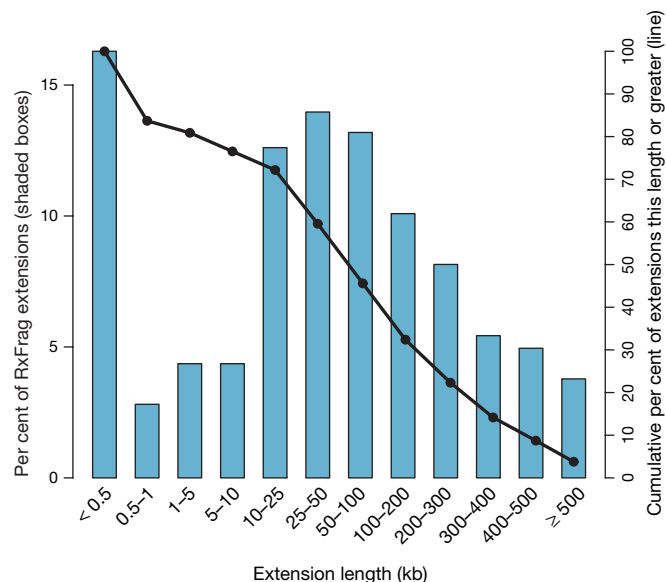


Figure 2 | Length of genomic extensions to GENCODE-annotated genes on the basis of RACE experiments followed by array hybridizations (RxFrags). The indicated bars reflect the frequency of extension lengths among different length classes. The solid line shows the cumulative frequency of extensions of that length or greater. Most of the extensions are greater than 50 kb from the annotated gene (see text for details).

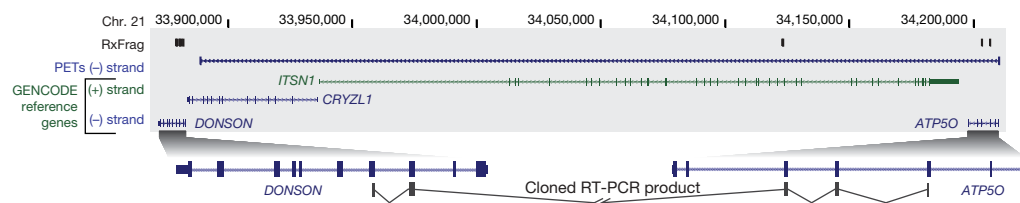


Figure 3 | Overview of RACE experiments showing a gene fusion.

Transcripts emanating from the region between the *DONSON* and *ATP5O* genes. A 330-kb interval of human chromosome 21 (within ENM005) is shown, which contains four annotated genes: *DONSON*, *CRYZL1*, *ITSN1* and *ATP5O*. The 5' RACE products generated from small intestine RNA and detected by

tiling-array analyses (RxFrag) are shown along the top. Along the bottom is shown the placement of a cloned and sequenced RT-PCR product that has two exons from the *DONSON* gene followed by three exons from the *ATP5O* gene; these sequences are separated by a 300 kb intron in the genome. A PET tag shows the termini of a transcript consistent with this RT-PCR product.

gene, with no evidence of sequences from two intervening protein-coding genes (*ITSN1* and *CRYZL1*).

Pseudogenes. Pseudogenes, reviewed in refs 21 and 22, are generally considered non-functional copies of genes, are sometimes transcribed and often complicate analysis of transcription owing to close sequence similarity to functional genes. We used various computational methods to identify 201 pseudogenes (124 processed and 77 non-processed) in the ENCODE regions (see Supplementary Information section 2.10 and ref. 23). Tiling-array analysis of 189 of these revealed that 56% overlapped at least one TxRag. However, possible cross-hybridization between the pseudogenes and their corresponding parent genes may have confounded such analyses. To assess better the extent of pseudogene transcription, 160 pseudogenes (111 processed and 49 non-processed) were examined for expression using RACE/tiling-array analysis (see Supplementary Information section 2.9.2). Transcripts were detected for 14 pseudogenes (8 processed and 6 non-processed) in at least one of the 12 tested RNA sources, the majority (9) being in testis (see ref. 23). Additionally, there was evidence for the transcription of 25 pseudogenes on the basis of their proximity (within 100 bp of a pseudogene end) to CAGE tags (8), PETs (2), or cDNAs/ESTs (21). Overall, we estimate that at least 19% of the pseudogenes in the ENCODE regions are transcribed, which is consistent with previous estimates^{24,25}.

Non-protein-coding RNA. Non-protein-coding RNAs (ncRNAs) include structural RNAs (for example, transfer RNAs, ribosomal RNAs, and small nuclear RNAs) and more recently discovered regulatory RNAs (for example, miRNAs). There are only 8 well-characterized ncRNA genes within the ENCODE regions (*U70*, *ACA36*, *ACA56*, *mir-192*, *mir-194-2*, *mir-196*, *mir-483* and *H19*), whereas representatives of other classes, (for example, box C/D snoRNAs, tRNAs, and functional snRNAs) seem to be completely absent in the ENCODE regions. Tiling-array data provided evidence for transcription in at least one of the assayed RNA samples for all of these ncRNAs, with the exception of *mir-483* (expression of *mir-483* might be specific to fetal liver, which was not tested). There is also evidence for the transcription of 6 out of 8 pseudogenes of ncRNAs (mainly snoRNA-derived). Similar to the analysis of protein-pseudogenes, the hybridization results could also originate from the known snoRNA gene elsewhere in the genome.

Many known ncRNAs are characterized by a well-defined RNA secondary structure. We applied two *de novo* ncRNA prediction algorithms—EvoFold and RNAz—to predict structured ncRNAs (as well as functional structures in mRNAs) using the multi-species sequence alignments (see below, Supplementary Information section 2.11 and ref. 26). Using a sensitivity threshold capable of detecting all known miRNAs and snoRNAs, we identified 4,986 and 3,707 candidate ncRNA loci with EvoFold and RNAz, respectively. Only 268 loci (5% and 7%, respectively) were found with both programs, representing a 1.6-fold enrichment over that expected by chance; the lack of more extensive overlap is due to the two programs having optimal sensitivity at different levels of GC content and conservation. We experimentally examined 50 of these targets using RACE/tiling-array analysis for brain and testis tissues (see Supplementary

Information sections 2.11 and 2.9.3); the predictions were validated at a 56%, 65%, and 63% rate for EvoFold, RNAz and dual predictions, respectively.

Primary transcripts. The detection of numerous unannotated transcripts coupled with increasing knowledge of the general complexity of transcription prompted us to examine the extent of primary (that is, unspliced) transcripts across the ENCODE regions. Three data sources provide insight about these primary transcripts: the GENCORE annotation, PETs, and RxRag extensions. Figure 4 summarizes the fraction of bases in the ENCODE regions that overlap transcripts identified by these technologies. Remarkably, 93% of bases are represented in a primary transcript identified by at least two independent observations (but potentially using the same technology); this figure is reduced to 74% in the case of primary transcripts detected by at least two different technologies. These increased spans are not mainly due to cell line rearrangements because they were present in multiple tissue experiments that confirmed the spans (see Supplementary Information section 2.12). These estimates assume that the presence of PETs or RxRags defining the terminal ends of a transcript imply that the entire intervening DNA is transcribed and then processed. Other mechanisms, thought to be unlikely in the human genome, such as *trans*-splicing or polymerase jumping would also produce these long termini and potentially should be reconsidered in more detail.

Previous studies have suggested a similar broad amount of transcription across the human^{14,15} and mouse²⁷ genomes. Our studies confirm these results, and have investigated the genesis of these transcripts in greater detail, confirming the presence of substantial intragenic and intergenic transcription. At the same time, many of the resulting transcripts are neither traditional protein-coding

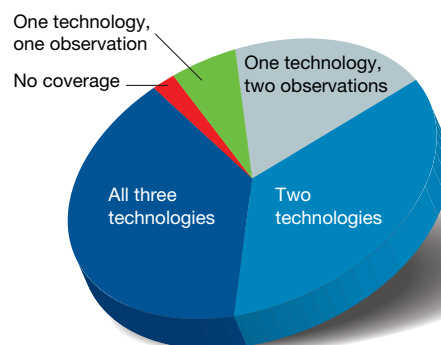


Figure 4 | Coverage of primary transcripts across ENCODE regions. Three different technologies (integrated annotation from GENCORE, RACE-array experiments (RxRags) and PET tags) were used to assess the presence of a nucleotide in a primary transcript. Use of these technologies provided the opportunity to have multiple observations of each finding. The proportion of genomic bases detected in the ENCODE regions associated with each of the following scenarios is depicted: detected by all three technologies, by two of the three technologies, by one technology but with multiple observations, and by one technology with only one observation. Also indicated are genomic bases without any detectable coverage of primary transcripts.

transcripts nor easily explained as structural non-coding RNAs. Other studies have noted complex transcription around specific loci or chimaeric-gene structures (for example refs 28–30), but these have often been considered exceptions; our data show that complex intercalated transcription is common at many loci. The results presented in the next section show extensive amounts of regulatory factors around novel TSSs, which is consistent with this extensive transcription. The biological relevance of these unannotated transcripts remains unanswered by these studies. Evolutionary information (detailed below) is mixed in this regard; for example, it indicates that unannotated transcripts show weaker evolutionary conservation than many other annotated features. As with other ENCODE-detected elements, it is difficult to identify clear biological roles for the majority of these transcripts; such experiments are challenging to perform on a large scale and, furthermore, it seems likely that many of the corresponding biochemical events may be evolutionarily neutral (see below).

Regulation of transcription

Overview. A significant challenge in biology is to identify the transcriptional regulatory elements that control the expression of each transcript and to understand how the function of these elements is coordinated to execute complex cellular processes. A simple, commonplace view of transcriptional regulation involves five types of *cis*-acting regulatory sequences—promoters, enhancers, silencers, insulators and locus control regions³¹. Overall, transcriptional regulation involves the interplay of multiple components, whereby the availability of specific transcription factors and the accessibility of specific genomic regions determine whether a transcript is generated³¹. However, the current view of transcriptional regulation is known to be overly simplified, with many details remaining to be established. For example, the consensus sequences of transcription factor binding sites (typically 6 to 10 bases) have relatively little information content and are present numerous times in the genome, with the great majority of these not participating in transcriptional regulation. Does chromatin structure then determine whether such a sequence has a regulatory role? Are there complex inter-factor interactions that integrate the signals from multiple sites? How are signals from different distal regulatory elements coupled without affecting all neighbouring genes? Meanwhile, our understanding of the repertoire of transcriptional events is becoming more complex, with an increasing appreciation of alternative TSSs^{32,33} and the presence of non-coding^{27,34} and anti-sense transcripts^{35,36}.

To better understand transcriptional regulation, we sought to begin cataloguing the regulatory elements residing within the 44 ENCODE regions. For this pilot project, we mainly focused on the binding of regulatory proteins and chromatin structure involved in transcriptional regulation. We analysed over 150 data sets, mainly from ChIP-chip^{37–39}, ChIP-PET and STAGE^{40,41} studies (see Supplementary Information section 3.1 and 3.2). These methods use chromatin immunoprecipitation with specific antibodies to enrich for DNA in physical contact with the targeted epitope. This enriched DNA can then be analysed using either microarrays (ChIP-chip) or high-throughput sequencing (ChIP-PET and STAGE). The assays included 18 sequence-specific transcription factors and components of the general transcription machinery (for example, RNA polymerase II (Pol II), TAF1 and TFIIB/GTF2B). In addition, we tested more than 600 potential promoter fragments for transcriptional activity by transient-transfection reporter assays that used 16 human cell lines³³. We also examined chromatin structure by studying the ENCODE regions for DNaseI sensitivity (by quantitative PCR⁴² and tiling arrays^{43,44}, see Supplementary Information section 3.3), histone composition⁴⁵, histone modifications (using ChIP-chip assays)^{37,46}, and histone displacement (using FAIRE, see Supplementary Information section 3.4). Below, we detail these analyses, starting with the efforts to define and classify the 5' ends of transcripts with respect to their associated regulatory signals. Following that are summaries of

generated data about sequence-specific transcription factor binding and clusters of regulatory elements. Finally, we describe how this information can be integrated to make predictions about transcriptional regulation.

Transcription start site catalogue. We analysed two data sets to catalogue TSSs in the ENCODE regions: the 5' ends of GENCODE-annotated transcripts and the combined results of two 5'-end-capture technologies—CAGE and PET-tagging. The initial results suggested the potential presence of 16,051 unique TSSs. However, in many cases, multiple TSSs resided within a single small segment (up to ~200 bases); this was due to some promoters containing TSSs with many very close precise initiation sites⁴⁷. To normalize for this effect, we grouped TSSs that were 60 or fewer bases apart into a single cluster, and in each case considered the most frequent CAGE or PET tag (or the 5'-most TSS in the case of TSSs identified only from GENCODE data) as representative of that cluster for downstream analyses.

The above effort yielded 7,157 TSS clusters in the ENCODE regions. We classified these TSSs into three categories: known (present at the end of GENCODE-defined transcripts), novel (supported by other evidence) and unsupported. The novel TSSs were further subdivided on the basis of the nature of the supporting evidence (see Table 3 and Supplementary Information section 3.5), with all four of the resulting subtypes showing significant overlap with experimental evidence using the GSC statistic. Although there is a larger relative proportion of singleton tags in the novel category, when analysis is restricted to only singleton tags, the novel TSSs continue to have highly significant overlap with supporting evidence (see Supplementary Information section 3.5.1).

Correlating genomic features with chromatin structure and transcription factor binding. By measuring relative sensitivity to DNaseI digestion (see Supplementary Information section 3.3), we identified DNaseI hypersensitive sites throughout the ENCODE regions. DHSs and TSSs both reflect genomic regions thought to be enriched for regulatory information and many DHSs reside at or near TSSs. We partitioned DHSs into those within 2.5 kb of a TSS (958; 46.5%) and the remaining ones, which were classified as distal (1,102; 53.5%). We then cross-analysed the TSSs and DHSs with data sets relating to histone modifications, chromatin accessibility and sequence-specific transcription factor binding by summarizing these signals in aggregate relative to the distance from TSSs or DHSs. Figure 5 shows representative profiles of specific histone modifications, Pol II and selected transcription factor binding for the different categories of TSSs. Further profiles and statistical analysis of these studies can be found in Supplementary Information 3.6.

In the case of the three TSS categories (known, novel and unsupported), known and novel TSSs are both associated with similar signals for multiple factors (ranging from histone modifications through DNaseI accessibility), whereas unsupported TSSs are not.

Table 3 | Different categories of TSSs defined on the basis of support from different transcript-survey methods

Category	Transcript survey method	Number of TSS clusters (non-redundant)*	P value†	Singleton clusters‡ (%)
Known	GENCODE 5' ends	1,730	2×10^{-70}	25 (74 overall)
Novel	GENCODE sense exons	1,437	6×10^{-39}	64
	GENCODE antisense exons	521	3×10^{-8}	65
	Unbiased transcription survey	639	7×10^{-63}	71
	CpG island	164	4×10^{-90}	60
	None	2,666	-	83.4
Unsupported	None	2,666	-	83.4

* Number of TSS clusters with this support, excluding TSSs from higher categories.

† Probability of overlap between the transcript support and the PET/CAGE tags, as calculated by the Genome Structure Correction statistic (see Supplementary Information section 1.3).

‡ Per cent of clusters with only one tag. For the 'known' category this was calculated as the per cent of GENCODE 5' ends with tag support (25%) or overall (74%).

The enrichments seen with chromatin modifications and sequence-specific factors, along with the significant clustering of this evidence, indicate that the novel TSSs do not reflect false positives and probably use the same biological machinery as other promoters. Sequence-specific transcription factors show a marked increase in binding across the broad region that encompasses each TSS. This increase is notably symmetric, with binding equally likely upstream or downstream of a TSS (see Supplementary Information section 3.7 for an explanation of why this symmetrical signal is not an artefact of the analysis of the signals). Furthermore, there is enrichment of SMARCC1 binding (a member of the SWI/SNF chromatin-modifying complex), which persists across a broader extent than other factors. The broad signals with this factor indicate that the ChIP-chip results reflect both specific enrichment at the TSS and broader enrichments across ~5-kb regions (this is not due to technical issues, see Supplementary Information section 3.8).

We selected 577 GENCODE-defined TSSs at the 5' ends of a protein-coding transcript with over 3 exons, to assess expression status. Each transcript was classified as: (1) 'active' (gene on) or 'inactive' (gene off) on the basis of the unbiased transcript surveys, and (2) residing near a 'CpG island' or not ('non-CpG island') (see Supplementary Information section 3.17). As expected, the aggregate

signal of histone modifications is mainly attributable to active TSSs (Fig. 5), in particular those near CpG islands. Pronounced doublet peaks at the TSS can be seen with these large signals (similar to previous work in yeast⁴⁸) owing to the chromatin accessibility at the TSS. Many of the histone marks and Pol II signals are now clearly asymmetrical, with a persistent level of Pol II into the genic region, as expected. However, the sequence-specific factors remain largely symmetrically distributed. TSSs near CpG islands show a broader distribution of histone marks than those not near CpG islands (see Supplementary Information section 3.6). The binding of some transcription factors (E2F1, E2F4 and MYC) is extensive in the case of active genes, and is lower (or absent) in the case of inactive genes.

Chromatin signature of distal elements. Distal DHSs show characteristic patterns of histone modification that are the inverse of TSSs, with high H3K4me1 accompanied by lower levels of H3K4Me3 and H3Ac (Fig. 5). Many factors with high occupancy at TSSs (for example, E2F4) show little enrichment at distal DHSs, whereas other factors (for example, MYC) are enriched at both TSSs and distal DHSs⁴⁹. A particularly interesting observation is the relative enrichment of the insulator-associated factor CTCF⁵⁰ at both distal DHSs and TSSs; this contrasts with SWI/SNF components SMARCC2 and SMARCC1, which are TSS-centric. Such differential

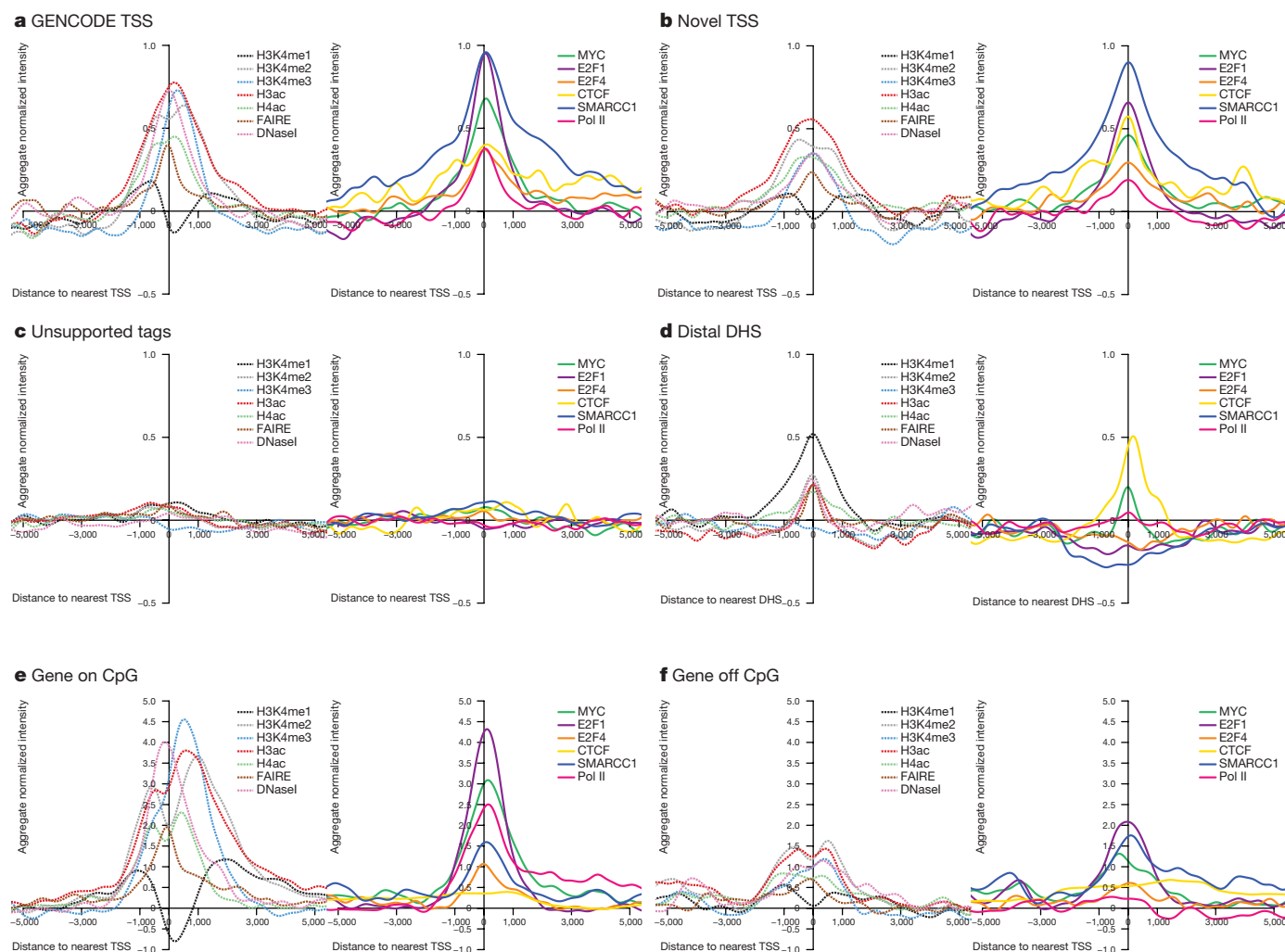


Figure 5 | Aggregate signals of tiling-array experiments from either ChIP-chip or chromatin structure assays, represented for different classes of TSSs and DHS. For each plot, the signal was first normalized with a mean of 0 and standard deviation of 1, and then the normalized scores were summed at each position for that class of TSS or DHS and smoothed using a kernel density method (see Supplementary Information section 3.6). For each class of sites there are two adjacent plots. The left plot depicts the data for general

factors: FAIRE and DNaseI sensitivity as assays of chromatin accessibility and H3K4me1, H3K4me2, H3K4me3, H3ac and H4ac histone modifications (as indicated); the right plot shows the data for additional factors, namely MYC, E2F1, E2F4, CTCF, SMARCC1 and Pol II. The columns provide data for the different classes of TSS or DHS (unsmoothed data and statistical analysis shown in Supplementary Information section 3.6).

behaviour of sequence-specific factors points to distinct biological differences, mediated by transcription factors, between distal regulatory sites and TSSs.

Unbiased maps of sequence-specific regulatory factor binding.

The previous section focused on specific positions defined by TSSs or DHSs. We then analysed sequence-specific transcription factor binding data in an unbiased fashion. We refer to regions with enriched binding of regulatory factors as RFBs. RFBs were identified on the basis of ChIP-chip data in two ways: first, each investigator developed and used their own analysis method(s) to define high-enrichment regions, and second (and independently), a stringent false discovery rate (FDR) method was applied to analyse all data using three cut-offs (1%, 5% and 10%). The laboratory-specific and FDR-based methods were highly correlated, particularly for regions with strong signals^{10,11}. For consistency, we used the results obtained with the FDR-based method (see Supplementary Information section 3.10). These RFBs can be used to find sequence motifs (see Supplementary Information section S3.11).

RFBs are associated with the 5' ends of transcripts. The distribution of RFBs is non-random (see ref. 10) and correlates with the positions of TSSs. We examined the distribution of specific RFBs relative to the known TSSs. Different transcription factors and histone modifications vary with respect to their association with TSSs (Fig. 6; see Supplementary Information section 3.12 for modelling of random expectation). Factors for which binding sites are most enriched at the 5' ends of genes include histone modifications, TAF1 and RNA Pol II with a hypo-phosphorylated carboxy-terminal domain⁵¹—confirming previous expectations. Surprisingly, we found that E2F1, a sequence-specific factor that regulates the expression of many genes at the G1 to S transition⁵², is also tightly associated with TSSs⁵²; this association is as strong as that of TAF1, the well-known TATA box-binding protein associated factor 1 (ref. 53). These results suggest that E2F1 has a more general role in transcription than previously suspected, similar to that for MYC^{54–56}. In contrast, the large-scale assays did not support the promoter binding that was found in smaller-scale studies (for example, on SIRT1 and SPI1 (PU1)).

Integration of data on sequence-specific factors. We expect that regulatory information is not dispersed independently across the genome, but rather is clustered into distinct regions⁵⁷. We refer to regions that contain multiple regulatory elements as 'regulatory clusters'. We sought to predict the location of regulatory clusters by

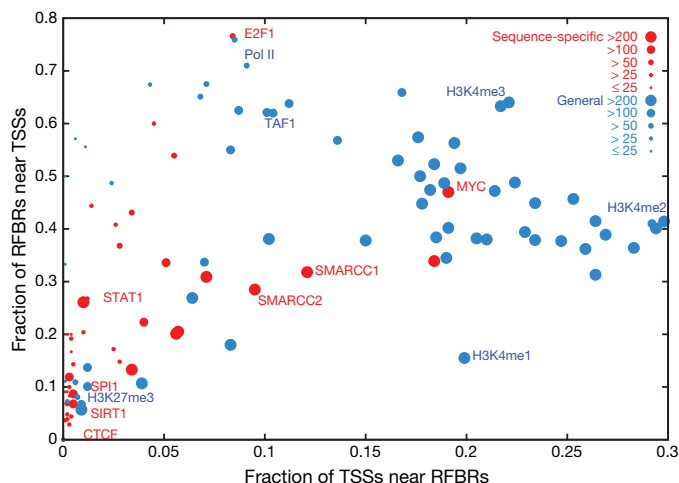


Figure 6 | Distribution of RFBs relative to GENCODE TSSs. Different RFBs from sequence-specific factors (red) or general factors (blue) are plotted showing their relative distribution near TSSs. The x axis indicates the proportion of TSSs close (within 2.5 kb) to the specified factor. The y axis indicates the proportion of RFBs close to TSSs. The size of the circle provides an indication of the number of RFBs for each factor. A handful of representative factors are labelled.

cross-integrating data generated using all transcription factor and histone modification assays, including results falling below an arbitrary threshold in individual experiments. Specifically, we used four complementary methods to integrate the data from 129 ChIP-chip data sets (see Supplementary Information section 3.13 and ref. 58). These four methods detect different classes of regulatory clusters and as a whole identified 1,393 clusters. Of these, 344 were identified by all four methods, with another 500 found by three methods (see Supplementary Information section 3.13.5). 67% of the 344 regulatory clusters identified by all four methods (or 65% of the full set of 1,393) reside within 2.5 kb of a known or novel TSS (as defined above; see Table 3 and Supplementary Information section 3.14 for a breakdown by category). Restricting this analysis to previously annotated TSSs (for example, RefSeq or Ensembl) reveals that roughly 25% of the regulatory clusters are close to a previously identified TSS. These results suggest that many of the regulatory clusters identified by integrating the ChIP-chip data sets are undiscovered promoters or are somehow associated with transcription in another fashion. To test these possibilities, sets of 126 and 28 non-GENCODE-based regulatory clusters were tested for promoter activity (see Supplementary Information section 3.15) and by RACE, respectively. These studies revealed that 24.6% of the 126 tested regulatory clusters had promoter activity and that 78.6% of the 28 regulatory clusters analysed by RACE yielded products consistent with a TSS⁵⁸. The ChIP-chip data sets were generated on a mixture of cell lines, predominantly HeLa and GM06990, and were different from the CAGE/PET data, meaning that tissue specificity contributes to the presence of unique TSSs and regulatory clusters. The large increase in promoter proximal regulatory clusters identified by including the additional novel TSSs coupled with the positive promoter and RACE assays suggests that most of the regulatory regions identifiable by these clustering methods represent bona fide promoters (see Supplementary Information section 3.16). Although the regulatory factor assays were more biased towards regions associated with promoters, many of the sites from these experiments would have previously been described as distal to promoters. This suggests that commonplace use of RefSeq- or Ensembl-based gene definition to define promoter proximity will dramatically overestimate the number of distal sites.

Predicting TSSs and transcriptional activity on the basis of chromatin structure. The strong association between TSSs and both histone modifications and DHSs prompted us to investigate whether the location and activity of TSSs could be predicted solely on the basis of chromatin structure information. We trained a support vector machine (SVM) by using histone modification data anchored around DHSs to discriminate between DHSs near TSSs and those distant from TSSs. We used a selected 2,573 DHSs, split roughly between TSS-proximal DHSs and TSS-distal DHSs, as a training set. The SVM performed well, with an accuracy of 83% (see Supplementary Information section 3.17). Using this SVM, we then predicted new TSSs using information about DHSs and histone modifications—of 110 high-scoring predicted TSSs, 81 resided within 2.5 kb of a novel TSS. As expected, these show a significant overlap to the novel TSS groups (defined above) but without a strong bias towards any particular category (see Supplementary Information section 3.17.1.5).

To investigate the relationship between chromatin structure and gene expression, we examined transcript levels in two cell lines using a transcript-tiling array. We compared this transcript data with the results of ChIP-chip experiments that measured histone modifications across the ENCODE regions. From this, we developed a variety of predictors of expression status using chromatin modifications as variables; these were derived using both decision trees and SVMs (see Supplementary Information section 3.17). The best of these correctly predicts expression status (transcribed versus non-transcribed) in 91% of cases. This success rate did not decrease dramatically when the predicting algorithm incorporated the results from one cell line to predict the expression status of another cell line. Interestingly, despite

the striking difference in histone modification enrichments in TSSs residing near versus those more distal to CpG islands (see Fig. 5 and Supplementary Information section 3.6), including information about the proximity to CpG islands did not improve the predictors. This suggests that despite the marked differences in histone modifications among these TSS classes, a single predictor can be made, using the interactions between the different histone modification levels.

In summary, we have integrated many data sets to provide a more complete view of regulatory information, both around specific sites (TSSs and DHSs) and in an unbiased manner. From analysing multiple data sets, we find 4,491 known and novel TSSs in the ENCODE regions, almost tenfold more than the number of established genes. This large number of TSSs might explain the extensive transcription described above; it also begins to change our perspective about regulatory information—without such a large TSS catalogue, many of the regulatory clusters would have been classified as residing distal to promoters. In addition to this revelation about the abundance of promoter-proximal regulatory elements, we also identified a considerable number of putative distal regulatory elements, particularly on the basis of the presence of DHSs. Our study of distal regulatory elements was probably most hindered by the paucity of data generated using distal-element-associated transcription factors; nevertheless, we clearly detected a set of distal-DHS-associated segments bound by CTCF or MYC. Finally, we showed that information about chromatin structure alone could be used to make effective predictions about both the location and activity of TSSs.

Replication

Overview. DNA replication must be carefully coordinated, both across the genome and with respect to development. On a larger scale, early replication in S phase is broadly correlated with gene density and transcriptional activity^{59–66}; however, this relationship is not universal, as some actively transcribed genes replicate late and vice versa^{61,64–68}. Importantly, the relationship between transcription and DNA replication emerges only when the signal of transcription is averaged over a large window (>100 kb)⁶³, suggesting that larger-scale chromosomal architecture may be more important than the activity of specific genes⁶⁹.

The ENCODE Project provided a unique opportunity to examine whether individual histone modifications on human chromatin can be correlated with the time of replication and whether such correlations support the general relationship of active, open chromatin with early replication. Our studies also tested whether segments showing interallelic variation in the time of replication have two different types of histone modifications consistent with an interallelic variation in chromatin state.

DNA replication data set. We mapped replication timing across the ENCODE regions by analysing Brd-U-labelled fractions from synchronized HeLa cells (collected at 2 h intervals throughout S phase) on tiling arrays (see Supplementary Information section 4.1). Although the HeLa cell line has a considerably altered karyotype, correlation of these data with other cell line data (see below) suggests the results are relevant to other cell types. The results are expressed as the time at which 50% of any given genomic position is replicated (TR50), with higher values signifying later replication times. In addition to the five ‘activating’ histone marks, we also correlated the TR50 with H3K27me₃, a modification associated with polycomb-mediated transcriptional repression^{70–74}. To provide a consistent comparison framework, the histone data were smoothed to 100-kb resolution, and then correlated with the TR50 data by a sliding window correlation analysis (see Supplementary Information section 4.2). The continuous profiles of the activating marks, histone H3K4 mono-, di-, and tri-methylation and histone H3 and H4 acetylation, are generally anti-correlated with the TR50 signal (Fig. 7a and Supplementary Information section 4.3). In contrast, H3K27me₃ marks show a predominantly positive correlation with late-replicating segments (Fig. 7a; see Supplementary Information section 4.3 for additional analysis).

Although most genomic regions replicate in a temporally specific window in S phase, other regions demonstrate an atypical pattern of replication (Pan-S) where replication signals are seen in multiple parts of S phase. We have suggested that such a pattern of replication stems from interallelic variation in the chromatin structure^{59,75}. If one allele is in active chromatin and the other in repressed chromatin, both types of modified histones are expected to be enriched in the Pan-S segments. An ENCODE region was classified as non-specific (or Pan-S) regions when >60% of the probes in a 10-kb window

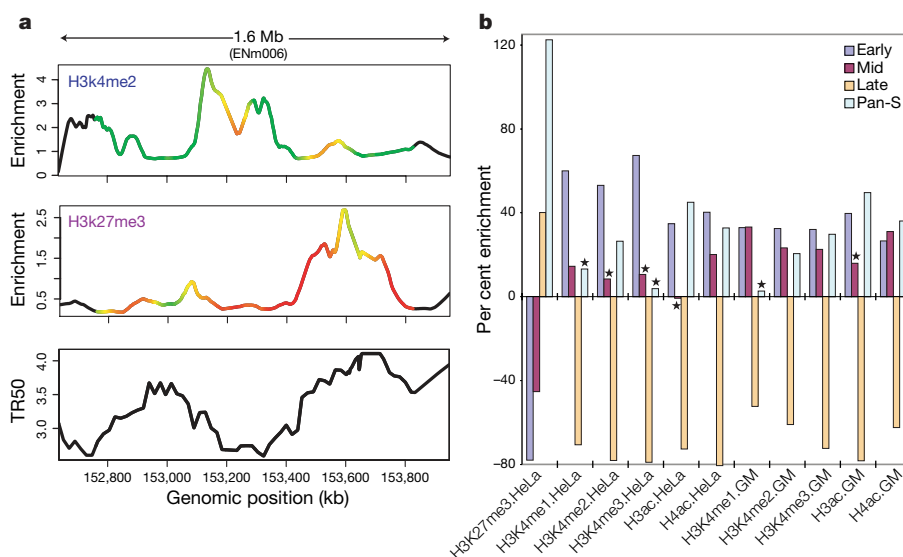


Figure 7 | Correlation between replication timing and histone modifications. **a**, Comparison of two histone modifications (H3K4me₂ and H3K27me₃), plotted as enrichment ratio from the Chip-chip experiments and the time for 50% of the DNA to replicate (TR50), indicated for ENCODE region ENm006. The colours on the curves reflect the correlation strength in a sliding 250-kb window. **b**, Differing levels of histone modification for

different TR50 partitions. The amounts of enrichment or depletion of different histone modifications in various cell lines are depicted (indicated along the bottom as ‘histone mark.cell line’; GM = GM06990). Asterisks indicate enrichments/depletions that are not significant on the basis of multiple tests. Each set has four partitions on the basis of replication timing: early, mid, late and Pan-S.

replicated in multiple intervals in S phase. The remaining regions were sub-classified into early-, mid- or late-replicating based on the average TR50 of the temporally specific probes within a 10-kb window⁷⁵. For regions of each class of replication timing, we determined the relative enrichment of various histone modification peaks in HeLa cells (Fig. 7b; Supplementary Information section 4.4). The correlations of activating and repressing histone modification peaks with TR50 are confirmed by this analysis (Fig. 7b). Intriguingly, the Pan-S segments are unique in being enriched for both activating (H3K4me2, H3ac and H4ac) and repressing (H3K27me3) histones, consistent with the suggestion that the Pan-S replication pattern arises from interallelic variation in chromatin structure and time of replication⁷⁵. This observation is also consistent with the Pan-S replication pattern seen for the H19/IGF2 locus, a known imprinted region with differential epigenetic modifications across the two alleles⁷⁶.

The extensive rearrangements in the genome of HeLa cells led us to ask whether the detected correlations between TR50 and chromatin state are seen with other cell lines. The histone modification data with GM06990 cells allowed us to test whether the time of replication of genomic segments in HeLa cells correlated with the chromatin state in GM06990 cells. Early- and late-replicating segments in HeLa cells are enriched and depleted, respectively, for activating marks in GM06990 cells (Fig. 7b). Thus, despite the presence of genomic rearrangements (see Supplementary Information section 2.12), the TR50 and chromatin state in HeLa cells are not far from a constitutive baseline also seen with a cell line from a different lineage. The enrichment of multiple activating histone modifications and the depletion of a repressive modification from segments that replicate early in S phase extends previous work in the field at a level of detail and scale not attempted before in mammalian cells. The duality of histone modification patterns in Pan-S areas of the HeLa genome, and the concordance of chromatin marks and replication time across two disparate cell lines (HeLa and GM06990) confirm the coordination of histone modifications with replication in the human genome.

Chromatin architecture and genomic domains

Overview. The packaging of genomic DNA into chromatin is intimately connected with the control of gene expression and other chromosomal processes. We next examined chromatin structure over a larger scale to ascertain its relation to transcription and other processes. Large domains (50 to >200 kb) of generalized DNaseI sensitivity have been detected around developmentally regulated gene clusters⁷⁷, prompting speculation that the genome is organized

into 'open' and 'closed' chromatin territories that represent higher-order functional domains. We explored how different chromatin features, particularly histone modifications, correlate with chromatin structure, both over short and long distances.

Chromatin accessibility and histone modifications. We used histone modification studies and DNaseI sensitivity data sets (introduced above) to examine general chromatin accessibility without focusing on the specific DHS sites (see Supplementary Information sections 3.1, 3.3 and 3.4). A fundamental difficulty in analysing continuous data across large genomic regions is determining the appropriate scale for analysis (for example, 2 kb, 5 kb, 20 kb, and so on). To address this problem, we developed an approach based on wavelet analysis, a mathematical tool pioneered in the field of signal processing that has recently been applied to continuous-value genomic analyses. Wavelet analysis provides a means for consistently transforming continuous signals into different scales, enabling the correlation of different phenomena independently at differing scales in a consistent manner.

Global correlations of chromatin accessibility and histone modifications. We computed the regional correlation between DNaseI sensitivity and each histone modification at multiple scales using a wavelet approach (Fig. 8 and Supplementary Information section 4.2). To make quantitative comparisons between different histone modifications, we computed histograms of correlation values between DNaseI sensitivity and each histone modification at several scales and then tested these for significance at specific scales. Figure 8c shows the distribution of correlation values at a 16-kb scale, which is considerably larger than individual *cis*-acting regulatory elements. At this scale, H3K4me2, H3K4me3 and H3ac show similarly high correlation. However, they are significantly distinguished from H3K4me1 and H4ac modifications ($P < 1.5 \times 10^{-35}$; see Supplementary Information section 4.5), which show lower correlation with DNaseI sensitivity. These results suggest that larger-scale relationships between chromatin accessibility and histone modifications are dominated by sub-regions in which higher average DNaseI sensitivity is accompanied by high levels of H3K4me2, H3K4me3 and H3ac modifications.

Local correlations of chromatin accessibility and histone modifications. Narrowing to a scale of ~2 kb revealed a more complex situation, in which H3K4me2 is the histone modification that is best correlated with DNaseI sensitivity. However, there is no clear combination of marks that correlate with DNaseI sensitivity in a way that is analogous to that seen at a larger scale (see Supplementary Information section 4.3). One explanation for the increased

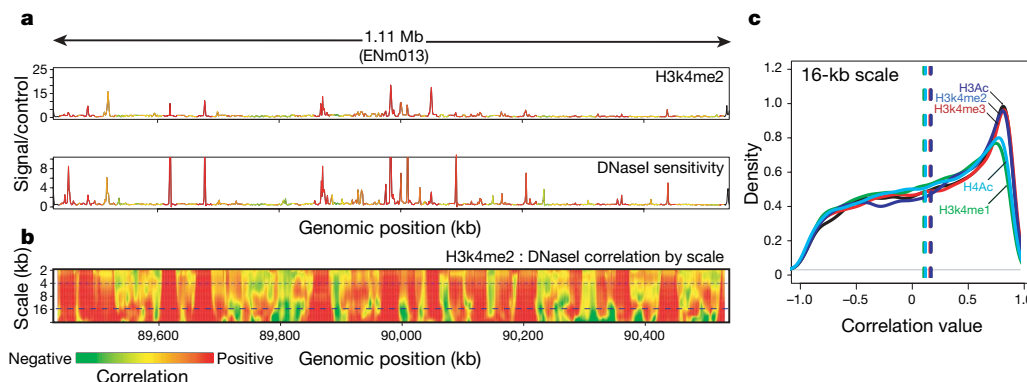


Figure 8 | Wavelet correlations of histone marks and DNaseI sensitivity. As an example, correlations between DNaseI sensitivity and H3K4me2 (both in the GM06990 cell line) over a 1.1-Mb region on chromosome 7 (ENCODE region ENm013) are shown. **a**, The relationship between histone modification H3K4me2 (upper plot) and DNaseI sensitivity (lower plot) is shown for ENCODE region ENm013. The curves are coloured with the strength of the local correlation at the 4-kb scale (top dashed line in panel **b**). **b**, The same data as in **a** are represented as a wavelet correlation. The

y axis shows the differing scales decomposed by the wavelet analysis from large to small scale (in kb); the colour at each point in the heatmap represents the level of correlation at the given scale, measured in a 20 kb window centred at the given position. **c**, Distribution of correlation values at the 16 kb scale between the indicated histone marks. The y axis is the density of these correlation values across ENCODE; all modifications show a peak at a positive-correlation value.

complexity at smaller scales is that there is a mixture of different classes of accessible chromatin regions, each having a different pattern of histone modifications. To examine this, we computed the degree to which local peaks in histone methylation or acetylation occur at DHSs (see Supplementary Information section 4.5.1). We found that 84%, 91% and 93% of significant peaks in H3K4 mono-, di- and tri-methylation, respectively, and 93% and 81% of significant peaks in H3ac and H4ac acetylation, respectively, coincided with DHSs (see Supplementary Information section 4.5). Conversely, a proportion of DHSs seemed not to be associated with significant peaks in H3K4 mono-, di- or tri-methylation (37%, 29% and 47%, respectively), nor with peaks in H3 or H4 acetylation (both 57%). Because only a limited number of histone modification marks were assayed, the possibility remains that some DHSs harbour other histone modifications. The absence of a more complete concordance between DHSs and peaks in histone acetylation is surprising given the widely accepted notion that histone acetylation has a central role in mediating chromatin accessibility by disrupting higher-order chromatin folding.

DNA structure at DHSs. The observation that distinctive hydroxyl radical cleavage patterns are associated with specific DNA structures⁷⁸ prompted us to investigate whether DHS subclasses differed with respect to their local DNA structure. Conversely, because different DNA sequences can give rise to similar hydroxyl radical cleavage patterns⁷⁹, genomic regions that adopt a particular local structure do not necessarily have the same nucleotide sequence. Using a Gibbs sampling algorithm on hydroxyl radical cleavage patterns of 3,150 DHSs⁸⁰, we discovered an 8-base segment with a conserved cleavage signature (CORCS; see Supplementary Information section 4.6). The underlying DNA sequences that give rise to this pattern have little primary sequence similarity despite this similar structural pattern. Furthermore, this structural element is strongly enriched in promoter-proximal DHSs (11.3-fold enrichment compared to the rest of the ENCODE regions) relative to promoter-distal DHSs (1.5-fold enrichment); this element is enriched 10.9-fold in CpG islands, but is higher still (26.4-fold) in CpG islands that overlap a DHS.

Large-scale domains in the ENCODE regions. The presence of extensive correlations seen between histone modifications, DNaseI

sensitivity, replication, transcript density and protein factor binding led us to investigate whether all these features are organized systematically across the genome. To test this, we performed an unsupervised training of a two-state HMM with inputs from these different features (see Supplementary Information section 4.7 and ref. 81). No other information except for the experimental variables was used for the HMM training routines. We consistently found that one state ('active') generally corresponded to domains with high levels of H3ac and RNA transcription, low levels of H3K27me3 marks, and early replication timing, whereas the other state ('repressed') reflected domains with low H3ac and RNA, high H3K27me3, and late replication (see Fig. 9). In total, we identified 70 active regions spanning 11.4 Mb and 82 inactive regions spanning 17.8 Mb (median size 136 kb versus 104 kb respectively). The active domains are markedly enriched for GENCODE TSSs, CpG islands and Alu repetitive elements ($P < 0.0001$ for each), whereas repressed regions are significantly enriched for LINE1 and LTR transposons ($P < 0.001$). Taken together, these results demonstrate remarkable concordance between ENCODE functional data types and provide a view of higher-order functional domains defined by a broader range of factors at a markedly higher resolution than was previously available⁸².

Evolutionary constraint and population variability

Overview. Functional genomic sequences can also be identified by examining evolutionary changes across multiple extant species and within the human population. Indeed, such studies complement experimental assays that identify specific functional elements^{83–85}. Evolutionary constraint (that is, the rejection of mutations at a particular location) can be measured by either (i) comparing observed substitutions to neutral rates calculated from multi-sequence alignments^{86–88}, or (ii) determining the presence and frequency of intra-species polymorphisms. Importantly, both approaches are indifferent to any specific function that the constrained sequence might confer.

Previous studies comparing the human, mouse, rat and dog genomes examined bulk evolutionary properties of all nucleotides in the genome, and provided little insight about the precise positions of constrained bases. Interestingly, these studies indicated that the

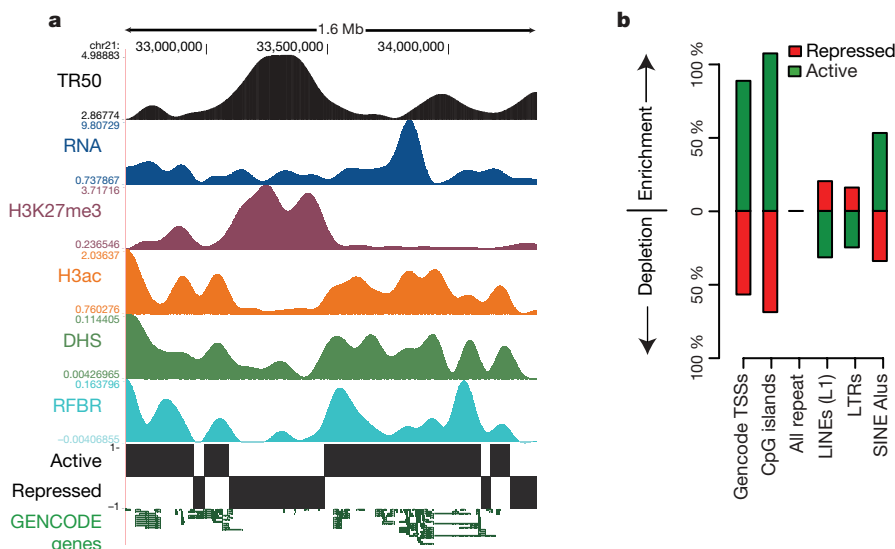


Figure 9 | Higher-order functional domains in the genome. The general concordance of multiple data types is shown for an illustrative ENCODE region (ENm005). **a**, Domains were determined by simultaneous HMM segmentation of replication time (TR50; black), bulk RNA transcription (blue), H3K27me3 (purple), H3ac (orange), DHS density (green), and RFBR density (light blue) measured continuously across the 1.6-Mb ENm005. All data were generated using HeLa cells. The histone, RNA, DHS and RFBR signals are wavelet-smoothed to an approximately 60-kb scale (see

Supplementary Information section 4.7). The HMM segmentation is shown as the blocks labelled 'active' and 'repressed' and the structure of GENCODE genes (not used in the training) is shown at the end. **b**, Enrichment or depletion of annotated sequence features (GENCODE TSSs, CpG islands, LINE1 repeats, Alu repeats, and non-exonic constrained sequences (CSs)) in active versus repressed domains. Note the marked enrichment of TSSs, CpG islands and Alus in active domains, and the enrichment of LINE and LTRs in repressed domains.

majority of constrained bases reside within the non-coding portion of the human genome. Meanwhile, increasingly rich data sets of polymorphisms across the human genome have been used extensively to establish connections between genetic variants and disease, but far fewer analyses have sought to use such data for assessing functional constraint⁸⁵.

The ENCODE Project provides an excellent opportunity for more fully exploiting inter- and intra-species sequence comparisons to examine genome function in the context of extensive experimental studies on the same regions of the genome. We consolidated the experimentally derived information about the ENCODE regions and focused our analyses on 11 major classes of genomic elements. These classes are listed in Table 4 and include two non-experimentally derived data sets: ancient repeats (ARs; mobile elements that inserted early in the mammalian lineage, have subsequently become dormant, and are assumed to be neutrally evolving) and constrained sequences (CSs; regions that evolve detectably more slowly than neutral sequences).

Comparative sequence data sets and analysis. We generated 206 Mb of genomic sequence orthologous to the ENCODE regions from 14 mammalian species using a targeted strategy that involved isolating⁸⁹ and sequencing⁹⁰ individual bacterial artificial chromosome clones. For an additional 14 vertebrate species, we used 340 Mb of orthologous genomic sequence derived from genome-wide sequencing efforts^{3–8,91–93}. The orthologous sequences were aligned using three alignment programs: TBA⁹⁴, MAVID⁹⁵ and MLAGAN⁹⁶. Four independent methods that generated highly concordant results⁹⁷ were then used to identify sequences under constraint (PhastCons⁸⁸, GERP⁸⁷, SCONe⁹⁸ and BinCons⁸⁶). From these analyses, we developed a high-confidence set of 'constrained sequences' that correspond to 4.9% of the nucleotides in the ENCODE regions. The threshold for determining constraint was set using a FDR rate of 5% (see ref. 97); this level is similar to previous estimates of the fraction of the human genome under mammalian constraint^{4,86–88} but the FDR rate was not chosen to fit this result. The median length of these constrained sequences is 19 bases, with the minimum being 8 bases—roughly the size of a typical transcription factor binding site. These analyses, therefore, provide a resolution of constrained sequences that is substantially better than that currently available using only whole-genome vertebrate sequences^{99–102}.

Intra-species variation studies mainly used SNP data from Phases I and II, and the 10 re-sequenced regions in ENCODE regions with 48 individuals of the HapMap Project¹⁰³; nucleotide insertion or deletion (indel) data were from the SNP Consortium and HapMap. We also examined the ENCODE regions for the presence of overlaps with known segmental duplications¹⁰⁴ and CNVs.

Experimentally identified functional elements and constrained sequences. We first compared the detected constrained sequences

with the positions of experimentally identified functional elements. A total of 40% of the constrained bases reside within protein-coding exons and their associated untranslated regions (Fig. 10) and, in agreement with previous genome-wide estimates, the remaining constrained bases do not overlap the mature transcripts of protein-coding genes^{4,5,88,105,106}. When we included the other experimental annotations, we found that an additional 20% of the constrained bases overlap experimentally identified non-coding functional regions, although far fewer of these regions overlap constrained sequences compared to coding exons (see below). Most experimental annotations are significantly different from a random expectation for both base-pair or element-level overlaps (using the GSC statistic, see Supplementary Information section 1.3), with a more striking deviation when considering elements (Fig. 11). The exceptions to this are pseudogenes, Un.TxFrags and RxFrags. The increase in significance moving from base-pair measures to the element level suggests that discrete islands of constrained sequence exist within experimentally identified functional elements, with the surrounding bases apparently not showing evolutionary constraint. This notion is discussed in greater detail in ref. 97.

We also examined measures of human variation (heterozygosity, derived allele-frequency spectra and indel rates) within the sequences of the experimentally identified functional elements (Fig. 12). For these studies, ARs were used as a marker for neutrally evolving sequence. Most experimentally identified functional elements are associated with lower heterozygosity compared to ARs, and a few have lower indel rates compared with ARs. Striking outliers are 3' UTRs, which have dramatically increased indel rates without an obvious cause. This is discussed in more depth in ref. 107.

These findings indicate that the majority of the evolutionarily constrained, experimentally identified functional elements show evidence of negative selection both across mammalian species and within the human population. Furthermore, we have assigned at least one molecular function to the majority (60%) of all constrained bases in the ENCODE regions.

Conservation of regulatory elements. The relationship between individual classes of regulatory elements and constrained sequences varies considerably, ranging from cases where there is strong evolutionary constraint (for example, pan-vertebrate ultraconserved regions^{108,109}) to examples of regulatory elements that are not conserved between orthologous human and mouse genes¹¹⁰. Within the ENCODE regions, 55% of RFBs overlap the high-confidence

Table 4 | Eleven classes of genomic elements subjected to evolutionary and population-genetics analyses

Abbreviation	Description
CDS	Coding exons, as annotated by GENCODE
5' UTR	5' untranslated region, as annotated by GENCODE
3' UTR	3' untranslated region, as annotated by GENCODE
Un.TxFrag	Unannotated region detected by RNA hybridization to tiling array (that is, unannotated TxFrag)
RxFrag	Region detected by RACE and analysis on tiling array
Pseudogene	Pseudogene identified by consensus pseudogene analysis
RFB	Regulatory factor binding region identified by ChIP-chip assay
RFB-SeqSp	Regulatory factor binding region identified only by ChIP-chip assays for factors with known sequence-specificity
DHS	DNaseI hypersensitive sites found in multiple tissues
FAIRE	Region of open chromatin identified by the FAIRE assay
TSS	Transcription start site
AR	Ancient repeat inserted early in the mammalian lineage and presumed to be neutrally evolving
CS	Constrained sequence identified by analysing multi-sequence alignments

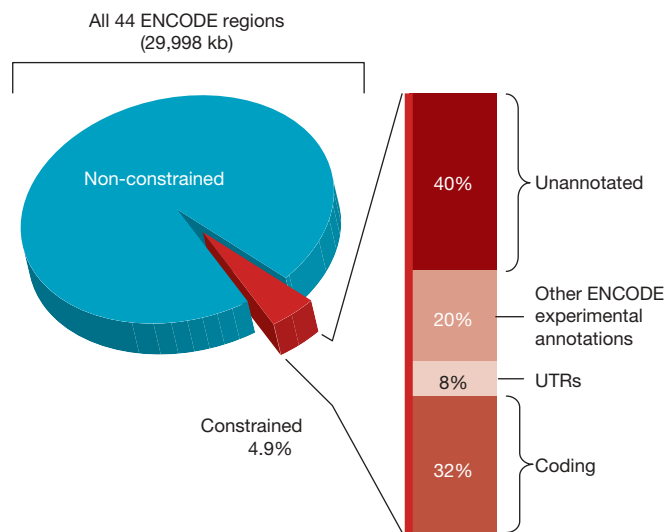


Figure 10 | Relative proportion of different annotations among constrained sequences. The 4.9% of bases in the ENCODE regions identified as constrained is subdivided into the portions that reflect known coding regions, UTRs, other experimentally annotated regions, and unannotated sequence.

constrained sequences. As expected, RFBs have many unconstrained bases, presumably owing to the small size of the specific binding site. We investigated whether the binding sites in RFBs could be further delimited using information about evolutionary constraint. For 7 out of 17 factors with either known TRANSFAC or Jaspar motifs, our ChIP-chip data revealed a marked enrichment of the appropriate motif within the constrained versus the unconstrained portions of the RFBs (see Supplementary Information section 5.1). This enrichment was seen for levels of stringency used for defining ChIP-chip-positive sites (1% and 5% FDR level), indicating that combining sequence constraint and ChIP-chip data may provide a highly sensitive means for detecting factor binding sites in the human genome.

Experimentally identified functional elements and genetic variation. The above studies focus on purifying (negative) selection. We used nucleotide variation to detect potential signals of adaptive (positive) selection. We modified the standard McDonald–Kreitman test (MK-test^{111,112}) and the Hudson–Kreitman–Aguade (HKA)¹¹³ test (see Supplementary Information section 5.2.1), to examine whether an entire set of sequence elements shows an excess of polymorphisms or an excess of inter-species divergence. We found that constrained sequences and coding exons have an excess of polymorphisms (consistent with purifying selection), whereas 5' UTRs

show evidence of an excess of divergence (with a portion probably reflecting positive selection). In general, non-coding genomic regions show more variation, with both a large number of segments that undergo purifying selection and regions that are fast evolving.

We also examined structural variation (that is, CNVs, inversions and translocations¹¹⁴; see Supplementary Information section 5.2.2). Within these polymorphic regions, we encountered significant overrepresentation of CDSs, TxFrags, and intra-species constrained sequences ($P < 10^{-3}$, Fig. 13), and also detected a statistically significant under-representation of ARs ($P = 10^{-3}$). A similar overrepresentation of CDSs and intra-species constrained sequences was found within non-polymorphic segmental duplications.

Unexplained constrained sequences. Despite the wealth of complementary data, 40% of the ENCODE-region sequences identified as constrained are not associated with any experimental evidence of function. There is no evidence indicating that mutational cold spots account for this constraint; they have similar measures of constraint to experimentally identified elements and harbour equal proportions of SNPs. To characterize further the unexplained constrained sequences, we examined their clustering and phylogenetic distribution. These sequences are not uniformly distributed across most ENCODE regions, and even in most ENCODE regions the distribution is different from constrained sequences within experimentally identified functional elements (see Supplementary Information section 5.3). The large fraction of constrained sequence that does not match any experimentally identified elements is not surprising considering that only a limited set of transcription factors, cell lines and biological conditions have thus far been examined.

Unconstrained experimentally identified functional elements. In contrast, an unexpectedly large fraction of experimentally identified functional elements show no evidence of evolutionary constraint ranging from 93% for Un.TxFrags to 12% for CDS. For most types of non-coding functional elements, roughly 50% of the individual elements seemed to be unconstrained across all mammals.

There are two methodological reasons that might explain the apparent excess of unconstrained experimentally identified functional elements: the underestimation of sequence constraint or overestimation of experimentally identified functional elements. We do not believe that either of these explanations fully accounts for the large and varied levels of unconstrained experimentally functional sequences. The set of constrained bases analysed here is highly accurate and complete due to the depth of the multiple alignment. Both by bulk fitting procedures and by comparison of SNP frequencies to constraint there is clearly a proportion of constrained bases not captured in the defined 4.9% of constrained sequences, but it is small (see

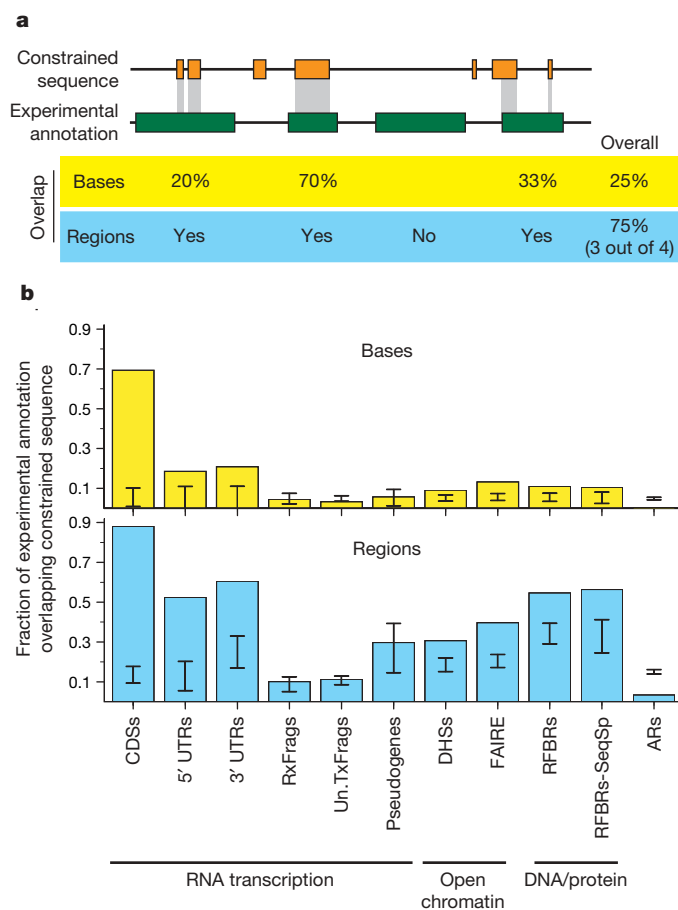
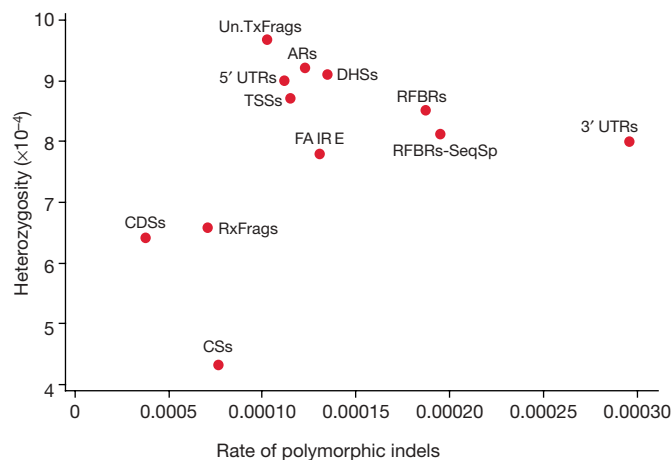


Figure 11 | Overlap of constrained sequences and various experimental annotations. **a**, A schematic depiction shows the different tests used for assessing overlap between experimental annotations and constrained sequences, both for individual bases and for entire regions. **b**, Observed fraction of overlap, depicted separately for bases and regions. The results are shown for selected experimental annotations. The internal bars indicate 95% confidence intervals of randomized placement of experimental elements using the GSC methodology to account for heterogeneity in the data sets. When the bar overlaps the observed value one cannot reject the hypothesis that these overlaps are consistent with random placements.



Supplementary Information section 5.4 and S5.5). More aggressive schemes to detect constraint only marginally increase the overlap with experimentally identified functional elements, and do so with considerably less specificity. Similarly, all experimental findings have been independently validated and, for the least constrained experimentally identified functional elements (Un.TxFrags and binding sites of sequence-specific factors), there is both internal validation and cross-validation from different experimental techniques. This suggests that there is probably not a significant overestimation of experimentally identified functional elements. Thus, these two explanations may contribute to the general observation about unconstrained functional elements, but cannot fully explain it.

Instead, we hypothesize five biological reasons to account for the presence of large amounts of unconstrained functional elements. The first two are particular to certain biological assays in which the elements being measured are connected to but do not coincide with the analysed region. An example of this is the parent transcript of an miRNA, where the current assays detect the exons (some of which are not under evolutionary selection), whereas the intronic miRNA actually harbours the constrained bases. Nevertheless, the transcript sequence provides the critical coupling between the regulated promoter and the miRNA. The sliding of transcription factors (which might bind a specific sequence but then migrate along the DNA) or the processivity of histone modifications across chromatin are more exotic examples of this. A related, second hypothesis is that delocalized behaviours of the genome, such as general chromatin accessibility, may be maintained by some biochemical processes (such as transcription of intergenic regions or specific factor binding) without the requirement for specific sequence elements. These two explanations of both connected components and diffuse components related to, but not coincident with, constrained sequences are particularly relevant for the considerable amount of unannotated and unconstrained transcripts.

The other three hypotheses may be more general—the presence of neutral (or near neutral) biochemical elements, of lineage-specific functional elements, and of functionally conserved but non-orthologous elements. We believe there is a considerable proportion of neutral biochemically active elements that do not confer a selective advantage or disadvantage to the organism. This neutral pool of sequence elements may turn over during evolutionary time,

emerging via certain mutations and disappearing by others. The size of the neutral pool would largely be determined by the rate of emergence and extinction through chance events; low information-content elements, such as transcription factor-binding sites¹¹⁰ will have larger neutral pools. Second, from this neutral pool, some elements might occasionally acquire a biological role and so come under evolutionary selection. The acquisition of a new biological role would then create a lineage-specific element. Finally, a neutral element from the general pool could also become a peer of an existing selected functional element and either of the two elements could then be removed by chance. If the older element is removed, the newer element has, in essence, been conserved without using orthologous bases, providing a conserved function in the absence of constrained sequences. For example, a common HNF4A binding site in the human and mouse genomes may not reflect orthologous human and mouse bases, though the presence of an HNF4A site in that region was evolutionarily selected for in both lineages. Note that both the neutral turnover of elements and the ‘functional peering’ of elements has been suggested for *cis*-acting regulatory elements in *Drosophila*^{115,116} and mammals¹¹⁰. Our data support these hypotheses, and we have generalized this idea over many different functional elements. The presence of conserved function encoded by conserved orthologous bases is a commonplace assumption in comparative genomics; our findings indicate that there could be a sizable set of functionally conserved but non-orthologous elements in the human genome, and that these seem unconstrained across mammals. Functional data akin to the ENCODE Project on other related species, such as mouse, would be critical to understanding the rate of such functionally conserved but non-orthologous elements.

Conclusion

The generation and analyses of over 200 experimental data sets from studies examining the 44 ENCODE regions provide a rich source of functional information for 30 Mb of the human genome. The first conclusion of these efforts is that these data are remarkably informative. Although there will be ongoing work to enhance existing assays, invent new techniques and develop new data-analysis methods, the generation of genome-wide experimental data sets akin to the ENCODE pilot phase would provide an impressive platform for future genome exploration efforts. This now seems feasible in light of throughput improvements of many of the assays and the ever-declining costs of whole-genome tiling arrays and DNA sequencing. Such genome-wide functional data should be acquired and released openly, as has been done with other large-scale genome projects, to ensure its availability as a new foundation for all biologists studying the human genome. It is these biologists who will often provide the critical link from biochemical function to biological role for the identified elements.

The scale of the pilot phase of the ENCODE Project was also sufficiently large and unbiased to reveal important principles about the organization of functional elements in the human genome. In many cases, these principles agree with current mechanistic models. For example, we observe trimethylation of H3K4 enriched near active genes, and have improved the ability to accurately predict gene activity based on this and other histone modifications. However, we also uncovered some surprises that challenge the current dogma on biological mechanisms. The generation of numerous intercalated transcripts spanning the majority of the genome has been repeatedly suggested^{13,14}, but this phenomenon has been met with mixed opinions about the biological importance of these transcripts. Our analyses of numerous orthogonal data sets firmly establish the presence of these transcripts, and thus the simple view of the genome as having a defined set of isolated loci transcribed independently does not seem to be accurate. Perhaps the genome encodes a network of transcripts, many of which are linked to protein-coding transcripts and to the majority of which we cannot (yet) assign a biological role. Our perspective of transcription and genes may have to evolve and also poses

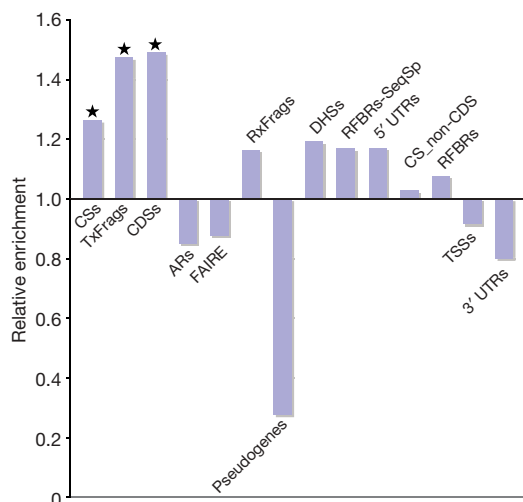


Figure 13 | CNV enrichment. The relative enrichment of different experimental annotations in the ENCODE regions associated with CNVs. CS, non-CDS are constrained sequences outside of coding regions. A value of 1 or less indicates no enrichment, and values greater than 1 show enrichment. Starred columns are cases that are significant on the basis of this enrichment being found in less than 5% of randomizations that matched each element class for length and density of features.

some interesting mechanistic questions. For example, how are splicing signals coordinated and used when there are so many overlapping primary transcripts? Similarly, to what extent does this reflect neutral turnover of reproducible transcripts with no biological role?

We gained subtler but equally important mechanistic findings relating to transcription, replication and chromatin modification. Transcription factors previously thought to primarily bind promoters bind more generally, and those which do bind to promoters are equally likely to bind downstream of a TSS as upstream. Interestingly, many elements that previously were classified as distal enhancers are, in fact, close to one of the newly identified TSSs; only about 35% of sites showing evidence of binding by multiple transcription factors are actually distal to a TSS. This need not imply that most regulatory information is confined to classic promoters, but rather it does suggest that transcription and regulation are coordinated actions beyond just the traditional promoter sequences. Meanwhile, although distal regulatory elements could be identified in the ENCODE regions, they are currently difficult to classify, in part owing to the lack of a broad set of transcription factors to use in analysing such elements. Finally, we now have a much better appreciation of how DNA replication is coordinated with histone modifications.

At the outset of the ENCODE Project, many believed that the broad collection of experimental data would nicely dovetail with the detailed evolutionary information derived from comparing multiple mammalian sequences to provide a neat 'dictionary' of conserved genomic elements, each with a growing annotation about their biochemical function(s). In one sense, this was achieved; the majority of constrained bases in the ENCODE regions are now associated with at least some experimentally derived information about function. However, we have also encountered a remarkable excess of experimentally identified functional elements lacking evolutionary constraint, and these cannot be dismissed for technical reasons. This is perhaps the biggest surprise of the pilot phase of the ENCODE Project, and suggests that we take a more 'neutral' view of many of the functions conferred by the genome.

METHODS

The methods are described in the Supplementary Information, with more technical details for each experiment often found in the references provided in Table 1. The Supplementary Information sections are arranged in the same order as the manuscript (with similar headings to facilitate cross-referencing). The first page of Supplementary Information also has an index to aid navigation. Raw data are available in ArrayExpress, GEO or EMBL/GenBank archives as appropriate, as detailed in Supplementary Information section 1.1. Processed data are also presented in a user-friendly manner at the UCSC Genome Browser's ENCODE portal (<http://genome.ucsc.edu/ENCODE/>).

Received 2 March; accepted 23 April 2007.

- International Human Genome Sequencing Consortium. Initial sequencing and analysis of the human genome. *Nature* **409**, 860–921 (2001).
- Venter, J. C. *et al.* The sequence of the human genome. *Science* **291**, 1304–1351 (2001).
- International Human Genome Sequencing Consortium. Finishing the euchromatic sequence of the human genome. *Nature* **431**, 931–945 (2004).
- Mouse Genome Sequencing Consortium. Initial sequencing and comparative analysis of the mouse genome. *Nature* **420**, 520–562 (2002).
- Rat Genome Sequencing Project Consortium. Genome sequence of the Brown Norway rat yields insights into mammalian evolution. *Nature* **428**, 493–521 (2004).
- Lindblad-Toh, K. *et al.* Genome sequence, comparative analysis and haplotype structure of the domestic dog. *Nature* **438**, 803–819 (2005).
- International Chicken Genome Sequencing Consortium. Sequence and comparative analysis of the chicken genome provide unique perspectives on vertebrate evolution. *Nature* **432**, 695–716 (2004).
- Chimpanzee Sequencing and Analysis Consortium. Initial sequence of the chimpanzee genome and comparison with the human genome. *Nature* **437**, 69–87 (2005).
- ENCODE Project Consortium. The ENCODE (ENCyclopedia Of DNA Elements) Project. *Science* **306**, 636–640 (2004).
- Zhang, Z. D. *et al.* Statistical analysis of the genomic distribution and correlation of regulatory elements in the ENCODE regions. *Genome Res.* **17**, 787–797 (2007).
- Euskirchen, G. M. *et al.* Mapping of transcription factor binding regions in mammalian cells by ChIP: comparison of array and sequencing based technologies. *Genome Res.* **17**, 898–909 (2007).
- Willingham, A. T. & Gingeras, T. R. TUF love for "junk" DNA. *Cell* **125**, 1215–1220 (2006).
- Carninci, P. *et al.* Genome-wide analysis of mammalian promoter architecture and evolution. *Nature Genet.* **38**, 626–635 (2006).
- Cheng, J. *et al.* Transcriptional maps of 10 human chromosomes at 5-nucleotide resolution. *Science* **308**, 1149–1154 (2005).
- Bertone, P. *et al.* Global identification of human transcribed sequences with genome tiling arrays. *Science* **306**, 2242–2246 (2004).
- Guigó, R. *et al.* EGASP: the human ENCODE Genome Annotation Assessment Project. *Genome Biol.* **7**, (Suppl. 1; S2) 1–31 (2006).
- Denoeud, F. *et al.* Prominent use of distal 5' transcription start sites and discovery of a large number of additional exons in ENCODE regions. *Genome Res.* **17**, 746–759 (2007).
- Tress, M. L. *et al.* The implications of alternative splicing in the ENCODE protein complement. *Proc. Natl Acad. Sci. USA* **104**, 5495–5500 (2007).
- Rozowsky, J. *et al.* The DART classification of unannotated transcription within ENCODE regions: Associating transcription with known and novel loci. *Genome Res.* **17**, 732–745 (2007).
- Kapranov, P. *et al.* Examples of the complex architecture of the human transcriptome revealed by RACE and high-density tiling arrays. *Genome Res.* **15**, 987–997 (2005).
- Balakirev, E. S. & Ayala, F. J. Pseudogenes: are they "junk" or functional DNA? *Annu. Rev. Genet.* **37**, 123–151 (2003).
- Mighell, A. J., Smith, N. R., Robinson, P. A. & Markham, A. F. Vertebrate pseudogenes. *FEBS Lett.* **468**, 109–114 (2000).
- Zheng, D. *et al.* Pseudogenes in the ENCODE regions: Consensus annotation, analysis of transcription and evolution. *Genome Res.* **17**, 839–851 (2007).
- Zheng, D. *et al.* Integrated pseudogene annotation for human chromosome 22: evidence for transcription. *J. Mol. Biol.* **349**, 27–45 (2005).
- Harrison, P. M., Zheng, D., Zhang, Z., Carriero, N. & Gerstein, M. Transcribed processed pseudogenes in the human genome: an intermediate form of expressed retrosequence lacking protein-coding ability. *Nucleic Acids Res.* **33**, 2374–2383 (2005).
- Washietl, S. *et al.* Structured RNAs in the ENCODE selected regions of the human genome. *Genome Res.* **17**, 852–864 (2007).
- Carninci, P. *et al.* The transcriptional landscape of the mammalian genome. *Science* **309**, 1559–1563 (2005).
- Runte, M. *et al.* The IC-SNURF-SNRPN transcript serves as a host for multiple small nucleolar RNA species and as an antisense RNA for UBE3A. *Hum. Mol. Genet.* **10**, 2687–2700 (2001).
- Seidl, C. I., Stricker, S. H. & Barlow, D. P. The imprinted *Air* ncRNA is an atypical RNAPII transcript that evades splicing and escapes nuclear export. *EMBO J.* **25**, 3565–3575 (2006).
- Parra, G. *et al.* Tandem chimerism as a means to increase protein complexity in the human genome. *Genome Res.* **16**, 37–44 (2006).
- Maston, G. A., Evans, S. K. & Green, M. R. Transcriptional regulatory elements in the human genome. *Annu. Rev. Genomics Hum. Genet.* **7**, 29–59 (2006).
- Trinklein, N. D., Aldred, S. J., Saldanha, A. J. & Myers, R. M. Identification and functional analysis of human transcriptional promoters. *Genome Res.* **13**, 308–312 (2003).
- Cooper, S. J., Trinklein, N. D., Anton, E. D., Nguyen, L. & Myers, R. M. Comprehensive analysis of transcriptional promoter structure and function in 1% of the human genome. *Genome Res.* **16**, 1–10 (2006).
- Cawley, S. *et al.* Unbiased mapping of transcription factor binding sites along human chromosomes 21 and 22 points to widespread regulation of noncoding RNAs. *Cell* **116**, 499–509 (2004).
- Yelin, R. *et al.* Widespread occurrence of antisense transcription in the human genome. *Nature Biotechnol.* **21**, 379–386 (2003).
- Katayama, S. *et al.* Antisense transcription in the mammalian transcriptome. *Science* **309**, 1564–1566 (2005).
- Ren, B. *et al.* Genome-wide location and function of DNA binding proteins. *Science* **290**, 2306–2309 (2000).
- Iyer, V. R. *et al.* Genomic binding sites of the yeast cell-cycle transcription factors SBF and MBF. *Nature* **409**, 533–538 (2001).
- Horak, C. E. *et al.* GATA-1 binding sites mapped in the β -globin locus by using mammalian ChIP-chip analysis. *Proc. Natl Acad. Sci. USA* **99**, 2924–2929 (2002).
- Wei, C. L. *et al.* A global map of p53 transcription-factor binding sites in the human genome. *Cell* **124**, 207–219 (2006).
- Kim, J., Bhinge, A. A., Morgan, X. C. & Iyer, V. R. Mapping DNA–protein interactions in large genomes by sequence tag analysis of genomic enrichment. *Nature Methods* **2**, 47–53 (2005).
- Dorschner, M. O. *et al.* High-throughput localization of functional elements by quantitative chromatin profiling. *Nature Methods* **1**, 219–225 (2004).
- Sabo, P. J. *et al.* Genome-scale mapping of DNase I sensitivity *in vivo* using tiling DNA microarrays. *Nature Methods* **3**, 511–518 (2006).
- Crawford, G. E. *et al.* DNase-chip: a high-resolution method to identify DNase I hypersensitive sites using tiled microarrays. *Nature Methods* **3**, 503–509 (2006).
- Hogan, G. J., Lee, C. K. & Lieb, J. D. Cell cycle-specified fluctuation of nucleosome occupancy at gene promoters. *PLoS Genet.* **2**, e158 (2006).

46. Koch, C. M. *et al.* The landscape of histone modifications across 1% of the human genome in five human cell lines. *Genome Res.* **17**, 691–707 (2007).
47. Smale, S. T. & Kadonaga, J. T. The RNA polymerase II core promoter. *Annu. Rev. Biochem.* **72**, 449–479 (2003).
48. Mito, Y., Henikoff, J. G. & Henikoff, S. Genome-scale profiling of histone H3.3 replacement patterns. *Nature Genet.* **37**, 1090–1097 (2005).
49. Heintzman, N. D. *et al.* Distinct and predictive chromatin signatures of transcriptional promoters and enhancers in the human genome. *Nature Genet.* **39**, 311–318 (2007).
50. Yusufzai, T. M., Tagami, H., Nakatani, Y. & Felsenfeld, G. CTCF tethers an insulator to subnuclear sites, suggesting shared insulator mechanisms across species. *Mol. Cell* **13**, 291–298 (2004).
51. Kim, T. H. *et al.* Direct isolation and identification of promoters in the human genome. *Genome Res.* **15**, 830–839 (2005).
52. Bieda, M., Xu, X., Singer, M. A., Green, R. & Farnham, P. J. Unbiased location analysis of E2F1-binding sites suggests a widespread role for E2F1 in the human genome. *Genome Res.* **16**, 595–605 (2006).
53. Ruppert, S., Wang, E. H. & Tjian, R. Cloning and expression of human TAF_{II}250: a TBP-associated factor implicated in cell-cycle regulation. *Nature* **362**, 175–179 (1993).
54. Fernandez, P. C. *et al.* Genomic targets of the human c-Myc protein. *Genes Dev.* **17**, 1115–1129 (2003).
55. Li, Z. *et al.* A global transcriptional regulatory role for c-Myc in Burkitt's lymphoma cells. *Proc. Natl Acad. Sci. USA* **100**, 8164–8169 (2003).
56. Orian, A. *et al.* Genomic binding by the *Drosophila* Myc, Max, Mad/Mnt transcription factor network. *Genes Dev.* **17**, 1101–1114 (2003).
57. de Laat, W. & Grosveld, F. Spatial organization of gene expression: the active chromatin hub. *Chromosome Res.* **11**, 447–459 (2003).
58. Trinklein, N. D. *et al.* Integrated analysis of experimental datasets reveals many novel promoters in 1% of the human genome. *Genome Res.* **17**, 720–731 (2007).
59. Jeon, Y. *et al.* Temporal profile of replication of human chromosomes. *Proc. Natl Acad. Sci. USA* **102**, 6419–6424 (2005).
60. Woodfine, K. *et al.* Replication timing of the human genome. *Hum. Mol. Genet.* **13**, 191–202 (2004).
61. White, E. J. *et al.* DNA replication-timing analysis of human chromosome 22 at high resolution and different developmental states. *Proc. Natl Acad. Sci. USA* **101**, 17771–17776 (2004).
62. Schubeler, D. *et al.* Genome-wide DNA replication profile for *Drosophila melanogaster*: a link between transcription and replication timing. *Nature Genet.* **32**, 438–442 (2002).
63. MacAlpine, D. M., Rodriguez, H. K. & Bell, S. P. Coordination of replication and transcription along a *Drosophila* chromosome. *Genes Dev.* **18**, 3094–3105 (2004).
64. Gilbert, D. M. Replication timing and transcriptional control: beyond cause and effect. *Curr. Opin. Cell Biol.* **14**, 377–383 (2002).
65. Schwaiger, M. & Schubeler, D. A question of timing: emerging links between transcription and replication. *Curr. Opin. Genet. Dev.* **16**, 177–183 (2006).
66. Hatton, K. S. *et al.* Replication program of active and inactive multigene families in mammalian cells. *Mol. Cell Biol.* **8**, 2149–2158 (1988).
67. Gartler, S. M., Goldstein, L., Tyler-Freer, S. E. & Hansen, R. S. The timing of *XIST* replication: dominance of the domain. *Hum. Mol. Genet.* **8**, 1085–1089 (1999).
68. Azuara, V. *et al.* Heritable gene silencing in lymphocytes delays chromatid resolution without affecting the timing of DNA replication. *Nature Cell Biol.* **5**, 668–674 (2003).
69. Cohen, S. M., Furey, T. S., Doggett, N. A. & Kaufman, D. G. Genome-wide sequence and functional analysis of early replicating DNA in normal human fibroblasts. *BMC Genomics* **7**, 301 (2006).
70. Cao, R. *et al.* Role of histone H3 lysine 27 methylation in Polycomb-group silencing. *Science* **298**, 1039–1043 (2002).
71. Muller, J. *et al.* Histone methyltransferase activity of a *Drosophila* Polycomb group repressor complex. *Cell* **111**, 197–208 (2002).
72. Bracken, A. P., Dietrich, N., Pasini, D., Hansen, K. H. & Helin, K. Genome-wide mapping of Polycomb target genes unravels their roles in cell fate transitions. *Genes Dev.* **20**, 1123–1136 (2006).
73. Kirmizis, A. *et al.* Silencing of human polycomb target genes is associated with methylation of histone H3 Lys 27. *Genes Dev.* **18**, 1592–1605 (2004).
74. Lee, T. I. *et al.* Control of developmental regulators by Polycomb in human embryonic stem cells. *Cell* **125**, 301–313 (2006).
75. Karnani, N., Taylor, C., Malhotra, A. & Dutta, A. Pan-S replication patterns and chromosomal domains defined by genome tiling arrays of human chromosomes. *Genome Res.* **17**, 865–876 (2007).
76. Delaval, K., Wagschal, A. & Feil, R. Epigenetic deregulation of imprinting in congenital diseases of aberrant growth. *Bioessays* **28**, 453–459 (2006).
77. Dillon, N. Gene regulation and large-scale chromatin organization in the nucleus. *Chromosome Res.* **14**, 117–126 (2006).
78. Burkhoff, A. M. & Tullius, T. D. Structural details of an adenine tract that does not cause DNA to bend. *Nature* **331**, 455–457 (1988).
79. Price, M. A. & Tullius, T. D. How the structure of an adenine tract depends on sequence context: a new model for the structure of T_nA_n DNA sequences. *Biochemistry* **32**, 127–136 (1993).
80. Greenbaum, J. A., Parker, S. C. J. & Tullius, T. D. Detection of DNA structural motifs in functional genomic elements. *Genome Res.* **17**, 940–946 (2007).
81. Thurman, R. E., Day, N., Noble, W. S. & Stamatoyannopoulos, J. A. Identification of higher-order functional domains in the human ENCODE regions. *Genome Res.* **17**, 917–927 (2007).
82. Gilbert, N. *et al.* Chromatin architecture of the human genome: gene-rich domains are enriched in open chromatin fibers. *Cell* **118**, 555–566 (2004).
83. Nobrega, M. A., Ovcharenko, I., Afzal, V. & Rubin, E. M. Scanning human gene deserts for long-range enhancers. *Science* **302**, 413 (2003).
84. Woolfe, A. *et al.* Highly conserved non-coding sequences are associated with vertebrate development. *PLoS Biol.* **3**, e7 (2005).
85. Drake, J. A. *et al.* Conserved noncoding sequences are selectively constrained and not mutation cold spots. *Nature Genet.* **38**, 223–227 (2006).
86. Margulies, E. H., Blanchette, M., NISC Comparative Sequencing Program, Haussler, D. & Green, E. D. Identification and characterization of multi-species conserved sequences. *Genome Res.* **13**, 2507–2518 (2003).
87. Cooper, G. M. *et al.* Distribution and intensity of constraint in mammalian genomic sequence. *Genome Res.* **15**, 901–913 (2005).
88. Siepel, A. *et al.* Evolutionarily conserved elements in vertebrate, insect, worm, and yeast genomes. *Genome Res.* **15**, 1034–1050 (2005).
89. Thomas, J. W. *et al.* Parallel construction of orthologous sequence-ready clone contig maps in multiple species. *Genome Res.* **12**, 1277–1285 (2002).
90. Blakesley, R. W. *et al.* An intermediate grade of finished genomic sequence suitable for comparative analyses. *Genome Res.* **14**, 2235–2244 (2004).
91. Aparicio, S. *et al.* Whole-genome shotgun assembly and analysis of the genome of *Fugu rubripes*. *Science* **297**, 1301–1310 (2002).
92. Jaillon, O. *et al.* Genome duplication in the teleost fish *Tetraodon nigroviridis* reveals the early vertebrate proto-karyotype. *Nature* **431**, 946–957 (2004).
93. Margulies, E. H. *et al.* An initial strategy for the systematic identification of functional elements in the human genome by low-redundancy comparative sequencing. *Proc. Natl Acad. Sci. USA* **102**, 4795–4800 (2005).
94. Blanchette, M. *et al.* Aligning multiple genomic sequences with the threaded blockset aligner. *Genome Res.* **14**, 708–715 (2004).
95. Bray, N. & Pachter, L. MAVID: constrained ancestral alignment of multiple sequences. *Genome Res.* **14**, 693–699 (2004).
96. Brudno, M. *et al.* LAGAN and Multi-LAGAN: efficient tools for large-scale multiple alignment of genomic DNA. *Genome Res.* **13**, 721–731 (2003).
97. Margulies, E. H. *et al.* Relationship between evolutionary constraint and genome function for 1% of the human genome. *Genome Res.* **17**, 760–774 (2007).
98. Asthana, S., Roytberg, M., Stamatoyannopoulos, J. A. & Sunyaev, S. Analysis of sequence conservation at nucleotide resolution. *PLoS Comp. Biol.* (submitted).
99. Cooper, G. M., Brudno, M., Green, E. D., Batzoglou, S. & Sidow, A. Quantitative estimates of sequence divergence for comparative analyses of mammalian genomes. *Genome Res.* **13**, 813–820 (2003).
100. Eddy, S. R. A model of the statistical power of comparative genome sequence analysis. *PLoS Biol.* **3**, e10 (2005).
101. Stone, E. A., Cooper, G. M. & Sidow, A. Trade-offs in detecting evolutionarily constrained sequence by comparative genomics. *Annu. Rev. Genomics Hum. Genet.* **6**, 143–164 (2005).
102. McAuliffe, J. D., Jordan, M. I. & Pachter, L. Subtree power analysis and species selection for comparative genomics. *Proc. Natl Acad. Sci. USA* **102**, 7900–7905 (2005).
103. International HapMap Consortium. A haplotype map of the human genome. *Nature* **437**, 1299–1320 (2005).
104. Cheng, Z. *et al.* A genome-wide comparison of recent chimpanzee and human segmental duplications. *Nature* **437**, 88–93 (2005).
105. Cooper, G. M. *et al.* Characterization of evolutionary rates and constraints in three mammalian genomes. *Genome Res.* **14**, 539–548 (2004).
106. Dermitzakis, E. T., Reymond, A. & Antonarakis, S. E. Conserved non-genic sequences - an unexpected feature of mammalian genomes. *Nature Rev. Genet.* **6**, 151–157 (2005).
107. Clark, T. G. *et al.* Small insertions/deletions and functional constraint in the ENCODE regions. *Genome Biol.* (submitted) (2007).
108. Bejerano, G. *et al.* Ultraconserved elements in the human genome. *Science* **304**, 1321–1325 (2004).
109. Woolfe, A. *et al.* Highly conserved non-coding sequences are associated with vertebrate development. *PLoS Biol.* **3**, e7 (2005).
110. Dermitzakis, E. T. & Clark, A. G. Evolution of transcription factor binding sites in Mammalian gene regulatory regions: conservation and turnover. *Mol. Biol. Evol.* **19**, 1114–1121 (2002).
111. McDonald, J. H. & Kreitman, M. Adaptive protein evolution at the *Adh* locus in *Drosophila*. *Nature* **351**, 652–654 (1991).
112. Andolfatto, P. Adaptive evolution of non-coding DNA in *Drosophila*. *Nature* **437**, 1149–1152 (2005).
113. Hudson, R. R., Kreitman, M. & Aguade, M. A test of neutral molecular evolution based on nucleotide data. *Genetics* **116**, 153–159 (1987).
114. Feuk, L., Carson, A. R. & Scherer, S. W. Structural variation in the human genome. *Nature Rev. Genet.* **7**, 85–97 (2006).
115. Ludwig, M. Z. *et al.* Functional evolution of a cis-regulatory module. *PLoS Biol.* **3**, e93 (2005).
116. Ludwig, M. Z. & Kreitman, M. Evolutionary dynamics of the enhancer region of even-skipped in *Drosophila*. *Mol. Biol. Evol.* **12**, 1002–1011 (1995).
117. Harrow, J. *et al.* GENCODE: producing a reference annotation for ENCODE. *Genome Biol.* **7**, (Suppl. 1; S4) 1–9 (2006).

118. Emanuelsson, O. *et al.* Assessing the performance of different high-density tiling microarray strategies for mapping transcribed regions of the human genome. *Genome Res.* advance online publication, doi: 10.1101/gr.5014606 (21 November 2006).
119. Kapranov, P. *et al.* Large-scale transcriptional activity in chromosomes 21 and 22. *Science* **296**, 916–919 (2002).
120. Bhinghe, A. A., Kim, J., Euskirchen, G., Snyder, M. & Iyer, V. R. Mapping the chromosomal targets of STAT1 by Sequence Tag Analysis of Genomic Enrichment (STAGE). *Genome Res.* **17**, 910–916 (2007).
121. Ng, P. *et al.* Gene identification signature (GIS) analysis for transcriptome characterization and genome annotation. *Nature Methods* **2**, 105–111 (2005).
122. Giresi, P. G., Kim, J., McDaniel, R. M., Iyer, V. R. & Lieb, J. D. FAIRE (Formaldehyde-Assisted Isolation of Regulatory Elements) isolates active regulatory elements from human chromatin. *Genome Res.* **17**, 877–885 (2006).
123. Rada-Iglesias, A. *et al.* Binding sites for metabolic disease related transcription factors inferred at base pair resolution by chromatin immunoprecipitation and genomic microarrays. *Hum. Mol. Genet.* **14**, 3435–3447 (2005).
124. Kim, T. H. *et al.* A high-resolution map of active promoters in the human genome. *Nature* **436**, 876–880 (2005).
125. Halees, A. S. & Weng, Z. PromoSer: improvements to the algorithm, visualization and accessibility. *Nucleic Acids Res.* **32**, W191–W194 (2004).
126. Bajic, V. B. *et al.* Performance assessment of promoter predictions on ENCODE regions in the EGASP experiment. *Genome Biol.* **7**, (Suppl 1; S3) 1–13 (2006).
127. Zheng, D. & Gerstein, M. B. A computational approach for identifying pseudogenes in the ENCODE regions. *Genome Biol.* **7**, S13.1–S13.10 (2006).
128. Stranger, B. E. *et al.* Genome-wide associations of gene expression variation in humans. *PLoS Genet* **1**, e78 (2005).
129. Turner, B. M. Reading signals on the nucleosome with a new nomenclature for modified histones. *Nature Struct. Mol. Biol.* **12**, 110–112 (2005).

Supplementary Information is linked to the online version of the paper at www.nature.com/nature.

Acknowledgements We thank D. Leja for providing graphical expertise and support. Funding support is acknowledged from the following sources: National Institutes of Health, The European Union BioSapiens NoE, Affymetrix, Swiss National Science Foundation, the Spanish Ministerio de Educación y Ciencia, Spanish Ministry of Education and Science, CIBERESP, Genome Spain and Generalitat de Catalunya, Ministry of Education, Culture, Sports, Science and Technology of Japan, the NCCR Frontiers in Genetics, the Jérôme Lejeune Foundation, the Childcare Foundation, the Novartis Foundations, the Danish Research Council, the Swedish Research Council, the Knut and Alice Wallenberg Foundation, the Wellcome Trust, the Howard Hughes Medical Institute, the Bio-X Institute, the RIKEN Institute, the US Army, National Science Foundation, the Deutsche Forschungsgemeinschaft, the Austrian Gen-AU program, the BBSRC and The European Molecular Biology Laboratory. We thank the Barcelona SuperComputing Center and the NIH Biowulf cluster for computer facilities. The Consortium thanks the ENCODE Scientific Advisory Panel for their advice on the project: G. Weinstock, M. Cherry, G. Churchill, M. Eisen, S. Elgin, J. Lis, J. Rine, M. Vidal and P. Zamore.

Author Information Reprints and permissions information is available at www.nature.com/reprints. The authors declare no competing financial interests. The list of individual authors is divided among the six main analysis groups and five organizational groups. Correspondence and requests for materials should be addressed to the co-chairs of the ENCODE analysis groups (listed in the Analysis Coordination group) E. Birney (birney@ebi.ac.uk); J. A. Stamatoyannopoulos (jstam@u.washington.edu); A. Dutta (ad8q@virginia.edu); R. Guigó (ruguigo@imim.es); T. R. Gingeras (Tom_Gingeras@affymetrix.com); E. H. Margulies (elliott@nhgri.nih.gov); Z. Weng (zhiping@bu.edu); M. Snyder (michael.snyder@yale.edu); E. T. Dermitzakis (md4@sanger.ac.uk) or collectively (encode_chairs@ebi.ac.uk).

The ENCODE Project Consortium

Analysis Coordination Ewan Birney¹, John A. Stamatoyannopoulos², Anindya Dutta³, Roderic Guigó^{4,5}, Thomas R. Gingeras⁶, Elliott H. Margulies⁷, Zhiping Weng^{8,9}, Michael Snyder^{10,11} & Emmanouil T. Dermitzakis¹²

Chromatin and Replication John A. Stamatoyannopoulos², Robert E. Thurman^{2,13}, Michael S. Kuehn^{2,13}, Christopher M. Taylor³, Shane Neph², Christoph M. Koch¹², Saurabh Asthana¹⁴, Ankit Malhotra³, Ivan Adzhubei¹⁴, Jason A. Greenbaum¹⁵, Robert M. Andrews¹², Paul Flicek¹, Patrick J. Boyle³, Hua Cao¹³, Nigel P. Carter¹², Gayle K. Clelland¹², Sean Davis¹⁶, Nathan Day², Pawandeep Dhami¹², Shane C. Dillon¹², Michael O. Dorschner², Heike Fiegler¹², Paul G. Giresi¹⁷, Jeff Goldy², Michael Hawrylycz¹⁸, Andrew Haydock², Richard Humbert², Keith D. James¹², Brett E. Johnson¹³, Ericka M. Johnson¹³, Tristan T. Frum¹³, Elizabeth R. Rosenzweig¹³, Neerja Karnani¹³, Kirsten Lee², Gregory C. Lefebvre¹², Patrick A. Navas¹³, Fidencio Neri², Stephen C. J. Parker¹⁵, Peter J. Sabo², Richard Sandstrom², Anthony Shafer², David Vetrie¹², Molly Weaver², Sarah

Wilcox¹², Man Yu¹³, Francis S. Collins⁷, Job Dekker¹⁹, Jason D. Lieb¹⁷, Thomas D. Tullius¹⁵, Gregory E. Crawford²⁰, Shamil Sunyaev¹⁴, William S. Noble², Ian Dunham¹² & Anindya Dutta³

Genes and Transcripts Roderic Guigó^{4,5}, France Denoeud⁵, Alexandre Reymond^{21,22}, Philipp Kapranov⁶, Joel Rozowsky¹¹, Deyou Zheng¹¹, Robert Castelo⁵, Adam Frankish¹², Jennifer Harrow¹², Srinka Ghosh⁶, Albin Sandelin²³, Ivo L. Hofacker²⁴, Robert Baertsch^{25,26}, Damian Keefe¹, Paul Flicek¹, Sujit Dike⁶, Jill Cheng⁶, Heather A. Hirsch²⁷, Edward A. Sekinger²⁷, Julien Lagarde⁵, Josep F. Abril^{5,28}, Atif Shahab²⁹, Christoph Flamm^{24,30}, Claudia Fried³⁰, Jörg Hackermüller³², Jana Hertel³⁰, Manja Lindemeyer³⁰, Kristin Missa^{30,31}, Andrea Tanzer^{24,30}, Stefan Washietl²⁴, Jan Korbel¹¹, Olof Emanuelsson¹¹, Jakob S. Pedersen²⁶, Nancy Holroyd¹², Ruth Taylor¹², David Swarbreck¹², Nicholas Matthews¹², Mark C. Dickson³³, Daryl J. Thomas^{25,26}, Matthew T. Weirauch²⁵, James Gilbert¹², Jorg Drenkow⁶, Ian Bell⁶, XiaoDong Zhao³⁴, K.G. Srinivasan³⁴, Wing-Kin Sung³⁴, Hong Sain Ooi³⁴, Kuo Ping Chiu³⁴, Sylvain Foissac⁴, Tyler Alioto⁴, Michael Brent³⁵, Lior Pachter³⁶, Michael L. Tress³⁷, Alfonso Valencia³⁷, Siew Woh Choo³⁴, Chiou Yu Choo³⁴, Catherine Ucla²², Caroline Manzano²², Carine Wyss²², Evelyn Cheung⁶, Taane G. Clark³⁸, James B. Brown³⁹, Madhavan Ganesh⁶, Sandeep Patel⁶, Hari Tammana⁶, Jacqueline Chrast²¹, Charlotte N. Henriksen²¹, Chikatoshi Kai²³, Jun Kawai^{23,40}, Ugrappa Nagalakshmi¹⁰, Jiaqian Wu¹⁰, Zheng Lian⁴¹, Jin Lian⁴¹, Peter Newburger⁴², Xueqing Zhang⁴², Peter Bickel⁴³, John S. Mattick⁴⁴, Piero Carninci⁴⁰, Yoshihide Hayashizaki^{23,40}, Sherman Weissman⁴¹, Emmanouil T. Dermitzakis¹², Elliott H. Margulies⁷, Tim Hubbard¹², Richard M. Myers³³, Jane Rogers¹², Peter F. Stadler^{24,30,45}, Todd M. Lowe²⁵, Chia-Lin Wei³⁴, Yijun Ruan³⁴, Michael Snyder^{10,11}, Ewan Birney¹, Kevin Struhl²⁷, Mark Gerstein^{11,46,47}, Stylianos E. Antonarakis²² & Thomas R. Gingeras⁶

Integrated Analysis and Manuscript Preparation James B. Brown³⁹, Paul Flicek¹, Yutao Fu⁸, Damian Keefe¹, Ewan Birney¹, France Denoeud⁵, Mark Gerstein^{11,46,47}, Eric D. Green^{7,48}, Philipp Kapranov⁶, Ulas Karaöz⁸, Richard M. Myers³³, William S. Noble², Alexandre Reymond^{21,22}, Joel Rozowsky¹¹, Kevin Struhl²⁷, Adam Siepel^{25,26}, John A. Stamatoyannopoulos², Christopher M. Taylor³, James Taylor^{49,50}, Robert E. Thurman²¹³, Thomas D. Tullius¹⁵, Stefan Washietl²⁴ & Deyou Zheng¹¹

Management Group Laura A. Liefer⁵¹, Kris A. Wetterstrand⁵¹, Peter J. Good⁵¹, Elise A. Feingold⁵¹, Mark S. Guyer⁵¹ & Francis S. Collins⁵²

Multi-species Sequence Analysis Elliott H. Margulies⁷, Gregory M. Cooper³³, George Asimenos⁵³, Daryl J. Thomas^{25,26}, Colin N. Dewey⁵⁴, Adam Siepel^{25,26}, Ewan Birney¹, Damian Keefe¹, Minmei Hou^{49,50}, James Taylor^{49,50}, Sergey Nikolaev²², Juan I. Montoya-Burgos⁵⁵, Ari Löytynoja¹, Simon Whelan¹, Fabio Pardi¹, Tim Massingham¹, James B. Brown³⁹, Haiyan Huang⁴³, Nancy R. Zhang^{43,56}, Peter Bickel⁴³, Ian Holmes⁵⁷, James C. Mullikin⁴⁸, Abel Ureta-Vidal¹, Benedict Paten¹, Michael Seringhaus¹, Deanna Church⁵⁸, Kate Rosenbloom²⁶, W. James Kent^{25,26}, Eric A. Stone³³, NISC Comparative Sequencing Program^{*}, Baylor College of Medicine Human Genome Sequencing Center^{*}, Washington University Genome Sequencing Center^{*}, Broad Institute^{*}, Children's Hospital Oakland Research Institute^{*}, Mark Gerstein^{11,46,47}, Stylianos E. Antonarakis²², Serafim Batzoglou⁵³, Nick Goldman¹, Ross C. Hardison^{50,59}, David Haussler^{25,26,60}, Webb Miller^{49,50,61}, Lior Pachter³⁶, Eric D. Green^{7,48} & Arend Sidow^{33,62}

Transcriptional Regulatory Elements Zhiping Weng^{8,9}, Nathan D. Trinklein³³, Yutao Fu⁸, Zhengdong D. Zhang¹¹, Ulas Karaöz⁸, Leah Barrera⁶⁸, Rhona Stuart⁶⁸, Deyou Zheng¹¹, Srinka Ghosh⁶, Paul Flicek¹, David C. King^{50,59}, James Taylor^{49,50}, Adam Ameur⁶⁹, Stefan Enroth⁶⁹, Mark C. Bieda⁷⁰, Christoph M. Koch¹², Heather A. Hirsch²⁷, Chia-Lin Wei³⁴, Jill Cheng⁶, Jonghwan Kim⁷¹, Akshay A. Bhinge⁷¹, Paul G. Giresi¹⁷, Nan Jiang⁷², Jun Liu³⁴, Fei Yao³⁴, Wing-Kin Sung³⁴, Kuo Ping Chiu³⁴, Vinsensius B. Vega³⁴, Charlie W. Lee³⁴, Patrick Ng³⁴, Atif Shahab²⁹, Edward A. Sekinger²⁷, Annie Yang²⁷, Zarmik Moqtaderi²⁷, Zhou Zhu²⁷, Xiaoqin Xu⁷⁰, Sharon Squazzo⁷⁰, Matthew J. Oberley⁷³, David Inman⁷³, Michael A. Singer⁷², Todd A. Richmond⁷², Kyle J. Munn^{72,74}, Alvaro Rada-Iglesias⁷⁴, Ola Wallerman⁷⁴, Jan Komorowski⁶⁹, Gayle K. Clelland¹², Sarah Wilcox¹², Shane C. Dillon¹², Robert M. Andrews¹², Joanna C. Fowler¹², Philippe Couttet¹², Keith D. James¹², Gregory C. Lefebvre¹², Alexander W. Bruce¹², Oliver M. Dovey¹², Peter D. Ellis¹², Pawandeep Dhami¹², Cordelia F. Langford¹², Nigel P. Carter¹², David Vetrie¹², Philipp Kapranov⁶, David A. Nix⁶, Ian Bell⁶, Sandeep Patel⁶, Joel Rozowsky¹¹, Ghia Euskirchen¹⁰, Stephen Hartman¹⁰, Jin Lian⁴¹, Jiaqian Wu¹⁰, Alexander E. Urban¹⁰, Peter Kraus¹⁰, Sara Van Calcar⁶⁸, Nate Heintzman⁶⁸, Tae Hoon Kim⁶⁸, Kun Wang⁶⁸, Chunxu Qu⁶⁸, Gary Hon⁶⁸, Rosa Luna⁷⁵, Christopher K. Glass⁷⁵, M. Geoff Rosenfeld⁷⁵, Shelley Force Aldred³³, Sara J. Cooper³³, Anason Halees⁸, Jane M. Lin⁹, Hennady P. Shulha⁹, Xiaoling Zhang⁸, Mousheng Xu⁸, Jaafar N. S. Haidar⁹, Yong Yu⁹, Ewan Birney¹, Sherman Weissman⁴¹, Yijun Ruan³⁴, Jason D. Lieb¹⁷, Vishwanath R. Iyer⁷¹, Roland D. Green⁷², Thomas R. Gingeras⁶, Claes Wadelius⁷⁴, Ian Dunham¹², Kevin Struhl²⁷, Ross C. Hardison^{50,59}, Mark Gerstein^{11,46,47}, Peggy J. Farnham⁷⁰, Richard M. Myers³³, Bing Ren⁶⁸ & Michael Snyder^{10,11}

UCSC Genome Browser Daryl J. Thomas^{25,26}, Kate Rosenbloom²⁶, Rachel A. Harte²⁶, Angie S. Hinrichs²⁶, Heather Trumbower²⁶, Hiram Clawson²⁶, Jennifer Hillman-Jackson²⁶, Ann S. Zweig²⁶, Kayla Smith²⁶, Archana Thakapallayil²⁶, Galt Barber²⁶, Robert M. Kuhn²⁶, Donna Karolchik²⁶, David Haussler^{25,26,60} & W. James Kent^{25,26}

Variation Emmanouil T. Dermitzakis¹², Lluís Armengol⁷⁶, Christine P. Bird¹², Taane G. Clark³⁸, Gregory M. Cooper³³, Paul I. W. de Bakker⁷⁷, Andrew D. Kern²⁶, Nuria Lopez-Bigas⁵, Joel D. Martin^{50,59}, Barbara E. Stranger¹², Daryl J. Thomas^{25,26}, Abigail Woodroffe⁷⁸, Serafim Batzoglou⁵³, Eugene Davydov⁵³, Antigone Dimas¹², Eduardo Eyra⁵, Ingileif B. Hallgrímsson⁷⁹, Ross C. Hardison^{50,59}, Julian Huppert¹², Arend Sidow^{33,62}, James Taylor^{49,50}, Heather Trumbower²⁶, Michael C. Zody⁷⁷, Roderic Guigó^{4,5}, James C. Mullikin⁷, Gonçalo R. Abecasis⁷⁸, Xavier Estivill^{76,80} & Ewan Birney¹.

***NISC Comparative Sequencing Program** Gerard G. Bouffard^{7,48}, Xiaobin Guan⁴⁸, Nancy F. Hansen⁴⁸, Jacquelyn R. Idol⁷, Valerie V.B. Maduro⁷, Baishali Maskeri⁴⁸, Jennifer C. McDowell⁴⁸, Morgan Park⁴⁸, Pamela J. Thomas⁴⁸, Alice C. Young⁴⁸ & Robert W. Blakesley^{7,48} **Baylor College of Medicine, Human Genome Sequencing Center** Donna M. Muzny⁶³, Erica Sodergren⁶³, David A. Wheeler⁶³, Kim C. Worley⁶³, Huaiyang Jiang⁶³, George M. Weinstock⁶³ & Richard A. Gibbs⁶³; **Washington University Genome Sequencing Center** Tina Graves⁶⁴, Robert Fulton⁶⁴, Elaine R. Mardis⁶⁴ & Richard K. Wilson⁶⁴ **Broad Institute** Michele Clamp⁶⁵, James Cuff⁶⁵, Sante Gnerre⁶⁵, David B. Jaffe⁶⁵, Jean L. Chang⁶⁵, Kerstin Lindblad-Toh⁶⁵ & Eric S. Lander^{65,66} **Children's Hospital Oakland Research Institute** Maxim Koriabine⁶⁷, Mikhail Nedefov⁶⁷, Kazutoyo Osoegawa⁶⁷, Yuko Yoshinaga⁶⁷, Baoli Zhu⁶⁷ & Pieter J. de Jong⁶⁷

Affiliations for participants: ¹EMBL-European Bioinformatics Institute, Wellcome Trust Genome Campus, Hinxton, Cambridge, CB10 1SD, UK. ²Department of Genome Sciences, 1705 NE Pacific Street, Box 357730, University of Washington, Seattle, Washington 98195, USA. ³Department of Biochemistry and Molecular Genetics, Jordan 1240, Box 800733, 1300 Jefferson Park Ave, University of Virginia School of Medicine, Charlottesville, Virginia 22908, USA. ⁴Genomic Bioinformatics Program, Center for Genomic Regulation, ⁵Research Group in Biomedical Informatics, Institut Municipal d'Investigació Mèdica/Universitat Pompeu Fabra, c/o Dr. Aiguader 88, Barcelona Biomedical Research Park Building, 08003 Barcelona, Catalonia, Spain. ⁶Affymetrix, Inc., Santa Clara, California 95051, USA. ⁷Genome Technology Branch, National Human Genome Research Institute, National Institutes of Health, Bethesda, Maryland 20892, USA. ⁸Bioinformatics Program, Boston University, 24 Cummington St., Boston, Massachusetts 02215, USA. ⁹Biomedical Engineering Department, Boston University, 44 Cummington St., Boston, Massachusetts 02215, USA. ¹⁰Department of Molecular, Cellular, and Developmental Biology, Yale University, New Haven, Connecticut 06520, USA. ¹¹Department of Molecular Biophysics and Biochemistry, Yale University, PO Box 208114, New Haven, Connecticut 06520, USA. ¹²The Wellcome Trust Sanger Institute, Wellcome Trust Genome Campus, Hinxton, Cambridge, CB10 1SA, UK. ¹³Division of Medical Genetics, 1705 NE Pacific Street, Box 357720, University of Washington, Seattle, Washington 98195, USA. ¹⁴Division of Genetics, Brigham and Women's Hospital and Harvard Medical School, 77 Avenue Louis Pasteur, Boston, Massachusetts 02115, USA. ¹⁵Department of Chemistry and Program in Bioinformatics, Boston University, 590 Commonwealth Avenue, Boston, Massachusetts 02215, USA. ¹⁶Genetics Branch, Center for Cancer Research, National Cancer Institute, National Institutes of Health, Bethesda, Maryland 20892, USA. ¹⁷Department of Biology and Carolina Center for Genome Sciences, CB# 3280, 202 Fordham Hall, The University of North Carolina at Chapel Hill, Chapel Hill, North Carolina 27599, USA. ¹⁸Allen Institute for Brain Sciences, 551 North 34th Street, Seattle, Washington 98103, USA. ¹⁹Program in Gene Function and Expression and Department of Biochemistry and Molecular Pharmacology, University of Massachusetts Medical School, 364 Plantation Street, Worcester, Massachusetts 01605, USA. ²⁰Institute for Genome Sciences & Policy and Department of Pediatrics, 101 Science Drive, Duke University, Durham, North Carolina 27708, USA. ²¹Center for Integrative Genomics, University of Lausanne, Genopode building, 1015 Lausanne, Switzerland. ²²Department of Genetic Medicine and Development, University of Geneva Medical School, 1211 Geneva, Switzerland. ²³Genome Exploration Research Group, RIKEN Genomic Sciences Center (GSC), RIKEN Yokohama Institute, 1-7-22, Suehiro-cho, Tsurumi-ku, Yokohama, Kanagawa, 230-0045, Japan. ²⁴Institute for Theoretical Chemistry, University of Vienna, Währingerstraße 17, A-1090 Wien, Austria. ²⁵Department of Biomolecular Engineering, University of California, Santa Cruz, 1156 High Street, Santa Cruz, California 95064, USA. ²⁶Center for Biomolecular Science and Engineering, Engineering 2, Suite 501, Mail Stop CBSE/ITI, University of California, Santa Cruz, California 95064, USA. ²⁷Department of Biological Chemistry & Molecular Pharmacology, Harvard Medical School, 240 Longwood Avenue, Boston, Massachusetts 02115, USA. ²⁸Department of Genetics, Facultat de Biologia, Universitat de Barcelona, Av Diagonal, 645, 08028, Barcelona, Catalonia, Spain. ²⁹Bioinformatics Institute, 30 Biopolis Street, #07-01 Matrix, Singapore, 138671, Singapore. ³⁰Bioinformatics Group, Department of Computer Science, ³¹Interdisciplinary Center of Bioinformatics, University of Leipzig, Härtelstraße 16-18, D-04107 Leipzig, Germany. ³²Fraunhofer Institut für Zelltherapie und Immunologie - IZI, Deutscher Platz 5e, D-04103 Leipzig, Germany. ³³Department of Genetics, Stanford University School of Medicine, Stanford,

California 94305, USA. ³⁴Genome Institute of Singapore, 60 Biopolis Street, Singapore 138672, Singapore. ³⁵Laboratory for Computational Genomics, Washington University, Campus Box 1045, Saint Louis, Missouri 63130, USA. ³⁶Department of Mathematics and Computer Science, University of California, Berkeley, California 94720, USA. ³⁷Spanish National Cancer Research Centre, CNIO, Madrid, E-28029, Spain. ³⁸Department of Epidemiology and Public Health, Imperial College, St Mary's Campus, Norfolk Place, London W2 1PG, UK. ³⁹Department of Applied Science & Technology, University of California, Berkeley, California 94720, USA. ⁴⁰Genome Science Laboratory, Discovery and Research Institute, RIKEN Wako Institute, 2-1 Hirasawa, Wako, Saitama, 351-0198, Japan. ⁴¹Department of Genetics, Yale University School of Medicine, 333 Cedar Street, New Haven, Connecticut 06510, USA. ⁴²Department of Pediatrics, University of Massachusetts Medical School, 55 Lake Avenue, North Worcester, Massachusetts 01605, USA. ⁴³Department of Statistics, University of California, Berkeley, California 94720, USA. ⁴⁴Institute for Molecular Bioscience, University of Queensland, St. Lucia, QLD 4072, Australia. ⁴⁵The Santa Fe Institute, 1399 Hyde Park Road, Santa Fe, New Mexico 87501, USA. ⁴⁶Department of Computer Science, Yale University, PO Box 208114, New Haven, Connecticut 06520-8114, USA. ⁴⁷Program in Computational Biology & Bioinformatics, Yale University, PO Box 208114, New Haven, Connecticut 06520-8114, USA. ⁴⁸NIH Intramural Sequencing Center, National Human Genome Research Institute, National Institutes of Health, Bethesda, Maryland 20892, USA. ⁴⁹Department of Computer Science and Engineering, The Pennsylvania State University, University Park, Pennsylvania 16802, USA. ⁵⁰Center for Comparative Genomics and Bioinformatics, Huck Institutes for Life Sciences, The Pennsylvania State University, University Park, Pennsylvania 16802, USA. ⁵¹Division of Extramural Research, National Human Genome Research Institute, National Institute of Health, 5635 Fishers Lane, Suite 4076, Bethesda, Maryland 20892-9305, USA. ⁵²Office of the Director, National Human Genome Research Institute, National Institute of Health, 31 Center Drive, Suite 4B09, Bethesda, Maryland 20892-2152, USA. ⁵³Department of Computer Science, Stanford University, Stanford, California 94305, USA. ⁵⁴Department of Biostatistics and Medical Informatics, University of Wisconsin-Madison, 6720 MSC, 1300 University Ave, Madison, Wisconsin 53706, USA. ⁵⁵Department of Zoology and Animal Biology, Faculty of Sciences, University of Geneva, 1205 Geneva, Switzerland. ⁵⁶Department of Statistics, Stanford University, Stanford, California 94305, USA. ⁵⁷Department of Bioengineering, University of California, Berkeley, California 94720-1762, USA. ⁵⁸National Center for Biotechnology Information, National Institutes of Health, Bethesda, Maryland 20894, USA. ⁵⁹Department of Biochemistry and Molecular Biology, Huck Institutes of Life Sciences, The Pennsylvania State University, University Park, Pennsylvania 16802, USA. ⁶⁰Howard Hughes Medical Institute, University of California, Santa Cruz, California 95064, USA. ⁶¹Department of Biology, The Pennsylvania State University, University Park, Pennsylvania 16802, USA. ⁶²Department of Pathology, Stanford University School of Medicine, Stanford, California 94305, USA. ⁶³Human Genome Sequencing Center and Department of Molecular and Human Genetics, Baylor College of Medicine, Houston, Texas 77030, USA. ⁶⁴Genome Sequencing Center, Washington University School of Medicine, Campus Box 8501, 4444 Forest Park Avenue, Saint Louis, Missouri 63108, USA. ⁶⁵Broad Institute of Harvard University and Massachusetts Institute of Technology, 320 Charles Street, Cambridge, Massachusetts 02141, USA. ⁶⁶Whitehead Institute for Biomedical Research, 9 Cambridge Center, Cambridge, Massachusetts 02142, USA. ⁶⁷Children's Hospital Oakland Research Institute, BACPAC Resources, 747 52nd Street, Oakland, California 94609, USA. ⁶⁸Ludwig Institute for Cancer Research, 9500 Gilman Drive, La Jolla, California 92093-0653, USA. ⁶⁹The Linnaeus Centre for Bioinformatics, Uppsala University, BMC, Box 598, SE-75121 Uppsala, Sweden. ⁷⁰Department of Pharmacology and the Genome Center, University of California, Davis, California 95616, USA. ⁷¹Institute for Cellular & Molecular Biology, The University of Texas at Austin, 1 University Station A4800, Austin, Texas 78712, USA. ⁷²NimbleGen Systems, Inc., 1 Science Court, Madison, Wisconsin 53711, USA. ⁷³University of Wisconsin Medical School, Madison, Wisconsin 53706, USA. ⁷⁴Department of Genetics and Pathology, Rudbeck Laboratory, Uppsala University, SE-75185 Uppsala, Sweden. ⁷⁵University of California, San Diego School of Medicine, 9500 Gilman Drive, La Jolla, California 92093, USA. ⁷⁶Genes and Disease Program, Center for Genomic Regulation, c/o Dr. Aiguader 88, Barcelona Biomedical Research Park Building, 08003 Barcelona, Catalonia, Spain. ⁷⁷Broad Institute of MIT and Harvard, 7 Cambridge Center, Cambridge, Massachusetts 02142, USA. ⁷⁸Center for Statistical Genetics, Department of Biostatistics, SPH II, 1420 Washington Heights, Ann Arbor, Michigan 48109-2029, USA. ⁷⁹Department of Statistics, University of Oxford, Oxford OX1 3TG, UK. ⁸⁰Universitat Pompeu Fabra, c/o Dr. Aiguader 88, Barcelona Biomedical Research Park Building, 08003 Barcelona, Catalonia, Spain. †Present addresses: Department of Genome Sciences, University of Washington School of Medicine, Seattle, Washington 98195, USA (G.M.C.); Department of Biological Statistics & Computational Biology, Cornell University, Ithaca, New York 14853, USA (A.S.); Faculty of Life Sciences, University of Manchester, Michael Smith Building, Oxford Road, Manchester, M13 9PT, UK (S.W.); SwitchGear Genomics, 1455 Adams Drive, Suite 2015, Menlo Park, California 94025, USA (N.D.T.; S.F.A.).