

The following resources related to this article are available online at www.sciencemag.org (this information is current as of August 4, 2009):

Updated information and services, including high-resolution figures, can be found in the online version of this article at:

<http://www.sciencemag.org/cgi/content/full/322/5903/881>

Supporting Online Material can be found at:

<http://www.sciencemag.org/cgi/content/full/322/5903/881/DC1>

A list of selected additional articles on the Science Web sites **related to this article** can be found at:

<http://www.sciencemag.org/cgi/content/full/322/5903/881#related-content>

This article **cites 137 articles**, 51 of which can be accessed for free:

<http://www.sciencemag.org/cgi/content/full/322/5903/881#otherarticles>

This article has been **cited by** 39 article(s) on the ISI Web of Science.

This article has been **cited by** 13 articles hosted by HighWire Press; see:

<http://www.sciencemag.org/cgi/content/full/322/5903/881#otherarticles>

This article appears in the following **subject collections**:

Genetics

<http://www.sciencemag.org/cgi/collection/genetics>

Information about obtaining **reprints** of this article or about obtaining **permission to reproduce this article** in whole or in part can be found at:

<http://www.sciencemag.org/about/permissions.dtl>

Genetic Mapping in Human Disease

David Altshuler,^{1,2,3,4,5*} Mark J. Daly,^{1,2,5*} Eric S. Lander^{1,6,7,8*}

Genetic mapping provides a powerful approach to identify genes and biological processes underlying any trait influenced by inheritance, including human diseases. We discuss the intellectual foundations of genetic mapping of Mendelian and complex traits in humans, examine lessons emerging from linkage analysis of Mendelian diseases and genome-wide association studies of common diseases, and discuss questions and challenges that lie ahead.

By the early 1900s, geneticists understood that Mendel's laws of inheritance underlie the transmission of genes in diploid organisms. They noted that some traits are inherited according to Mendel's ratios, as a result of alterations in single genes, and they developed methods to map the genes responsible. They also recognized that most naturally occurring trait variation, while showing strong correlation among relatives, involves the action of multiple genes and nongenetic factors.

Although it was clear that these insights applied to humans as much as to fruit flies, it took most of the century to turn these concepts into practical tools for discovering genes contributing to human diseases. Starting in the 1980s, the use of naturally occurring DNA variation as markers to trace inheritance in families led to the discovery of thousands of genes for rare Mendelian diseases. Despite great hopes, the approach proved unsuccessful for common forms of human diseases—such as diabetes, heart disease, and cancer—that show complex inheritance in the general population.

Over the past year, a new approach to genetic mapping has yielded the first general progress toward mapping loci that influence susceptibility to common human diseases. Still, most of the genes and mutations underlying these findings remain to be defined, let alone understood, and it remains unclear how much of the heritability of common disease they explain. Below, we discuss the intellectual foundations of genetic mapping, examine emerging lessons, and discuss questions and challenges that lie ahead.

Genetic Mapping by Linkage and Association

Genetic mapping is the localization of genes underlying phenotypes on the basis of correlation with DNA variation, without the need for prior hypotheses about biological function. The simplest form, called linkage analysis, was conceived

by Sturtevant for fruit flies in 1913 (1). Linkage analysis involves crosses between parents that vary at a Mendelian trait and at many polymorphic variants ("markers"); because of meiotic recombination, any marker showing correlated segregation ("linkage") with the trait must lie nearby in the genome.

In the 1970s, the ability to clone and sequence DNA made it possible to tie genetic linkage maps in model organisms to the underlying DNA sequence, and thereby to molecularly clone the genes responsible for any Mendelian trait solely on the basis of their genomic position (2, 3). Such studies typically involved three steps: (i) identifying the locus responsible through a genome-wide search; (ii) sequencing the region in cases and controls to define causal mutation(s); and (iii) studying the molecular and cellular functions of the genes discovered. So-called "positional cloning" became a mainstay of experimental genetics, identifying pathways that are crucial in development and physiology.

Linkage analysis in humans. For most of the 20th century, genome-wide linkage mapping was impractical in humans: Family sizes are small, crosses are not by design, and there were too few classical genetic markers to systematically trace inheritance. Progress in identifying the genes contributing to human traits was initially limited to studies of biological candidates such as blood-type antigens (4) and hemoglobin β protein in sickle-cell anemia (5).

In 1980, Botstein and colleagues, building on their use of DNA polymorphisms to study linkage in yeast (6) and the finding of DNA polymorphism at the globin locus in humans (7, 8), proposed the use of naturally occurring DNA sequence polymorphisms as generic markers to create a human genetic map and systematically trace the transmission of chromosomal regions in families (9). The feasibility of genetic mapping in humans was soon demonstrated with the localization of Huntington disease in 1983 (10). A rudimentary genetic linkage map with ~400 DNA markers was generated by 1987 (11) and was fleshed out to ~5000 markers by 1996 (12). Physical maps providing access to linked chromosomal regions were developed by 1995 (13). With these tools, positional cloning became possible in humans, and the number of disorders tied to a specific gene grew from ~100 in the late 1980s to >2200 today (14).

Several lessons emerged from studies of Mendelian disease genes: (i) The "candidate gene" approach was woefully inadequate; most disease genes were completely unsuspected on the basis

of previous knowledge. (ii) Disease-causing mutations often cause major changes in encoded proteins. (iii) Loci typically harbor many disease-causing alleles, mostly rare in the population. (iv) Mendelian diseases often revealed great complexity, such as locus heterogeneity, incomplete penetrance, and variable expressivity.

Geneticists were eager to apply genetic mapping to common diseases, which also show familial clustering. Mendelian subtypes of common diseases [such as breast cancer (15), hypertension (16), and diabetes (17)] were elucidated, but mutations in these genes explained few cases in the population. In common forms of common disease, risk to relatives is lower than in Mendelian cases, and linkage studies with excellent power to detect a single causal gene yielded equivocal results.

These features were consistent with, but did not prove, a polygenic model. The idea that commonly varying traits might be polygenic in nature was offered by East in 1910 (18). By 1920, linkage mapping was used to identify multiple unlinked factors influencing truncate wings in *Drosophila* (19), and Fisher had developed a mathematical framework for relating Mendelian factors and quantitative traits (20). In the late 1980s, linkage mapping of complex traits was made feasible for experimental organisms through the use of genetic mapping in large crosses (21). But there was little success in humans.

Genetic association in populations. A possible path forward emerged from population genetics and genomics. Instead of mapping disease genes by tracing transmission in families, one might localize them through association studies—that is, comparisons of frequencies of genetic variants among affected and unaffected individuals.

Genetic association studies were not a new idea. In the 1950s, such studies revealed correlations between blood-group antigens and peptic ulcer disease (4); in the 1960s and 1970s, common variation at the human leukocyte antigen (HLA) locus was associated with autoimmune and infectious diseases (22); and in the 1980s, apolipoprotein E was implicated in the etiology of Alzheimer's disease (23). Still, only about a dozen extensively reproduced associations of common variants (outside the HLA locus) were identified in the 20th century (24).

A central problem was that association studies of candidate genes were a shot in the dark: They were limited to specific variants in biological candidate genes, each with a tiny a priori probability of being disease-causing. Moreover, association studies were susceptible to false positives due to population structure, because there was no way to assess differences in the genetic background of cases and controls. Although many claims of associations were published, the statistical support tended to be weak and few were subsequently replicated (25).

In the mid-1990s, a systematic genome-wide approach to association studies was proposed (26–28): to develop a catalog of common human genetic variants and test the variants for associa-

¹Broad Institute of Harvard and MIT, Cambridge, MA 02142, USA. ²Center for Human Genetic Research and Department of Medicine, Massachusetts General Hospital, Boston, MA 02114, USA. ³Department of Molecular Biology, Massachusetts General Hospital, Boston, MA 02114, USA. ⁴Department of Genetics, Harvard Medical School, Boston, MA 02114, USA. ⁵Department of Medicine, Harvard Medical School, Boston, MA 02114, USA. ⁶Department of Systems Biology, Harvard Medical School, Boston, MA 02114, USA. ⁷Department of Biology, Massachusetts Institute of Technology, Cambridge, MA 02139, USA. ⁸Whitehead Institute for Biomedical Research, Cambridge, MA 02142, USA.

*To whom correspondence should be addressed. E-mail: altshuler@molbio.mgh.harvard.edu (D.A.); mj Daly@chgr.mgh.harvard.edu (M.J.D.); lander@broad.mit.edu (E.S.L.)

tion to disease risk. The focus on common variants as a mapping tool was a matter of practicality, grounded in population genetics. The human population has recently grown exponentially from a small size. As predicted by classical theory (29), humans have limited genetic variation: The heterozygosity rate for single-nucleotide polymorphisms (SNPs) is ~1 in 1000 bases (30–32). Moreover, perhaps 90% of heterozygous sites in each individual are common variants, typically shared among continental populations (33).
If most genetic variation in an individual is common, then why are mutations responsible for Mendelian diseases typically rare? One

answer is natural selection: Mutations that cause strongly deleterious phenotypes—as most Mendelian diseases appear to be—are lost to purifying selection. But if deleterious mutations are typically rare, how could common variants play a role in disease? Common diseases often have late onset, with modest or no obvious impact on reproductive fitness. Mildly deleterious alleles can rise to moderate frequency, particularly in populations that have undergone recent expansion (34). Moreover, some alleles that were advantageous or neutral during human evolution might now confer susceptibility to disease because of changes in living conditions accompanying civi-

lization. Finally, disease-causing alleles could be maintained at high frequency if they were under balancing selection, with disease burden offset by a beneficial phenotype (as in sickle-cell disease and malaria resistance).
These lines of reasoning led to the so-called “common disease–common variant” (CD-CV) hypothesis: the proposal that common polymorphisms (classically defined as having a minor allele frequency of >1%) might contribute to susceptibility to common diseases (26–28). If so, genome-wide association studies (GWASs) of common variants might be used to map loci contributing to common diseases. The concept

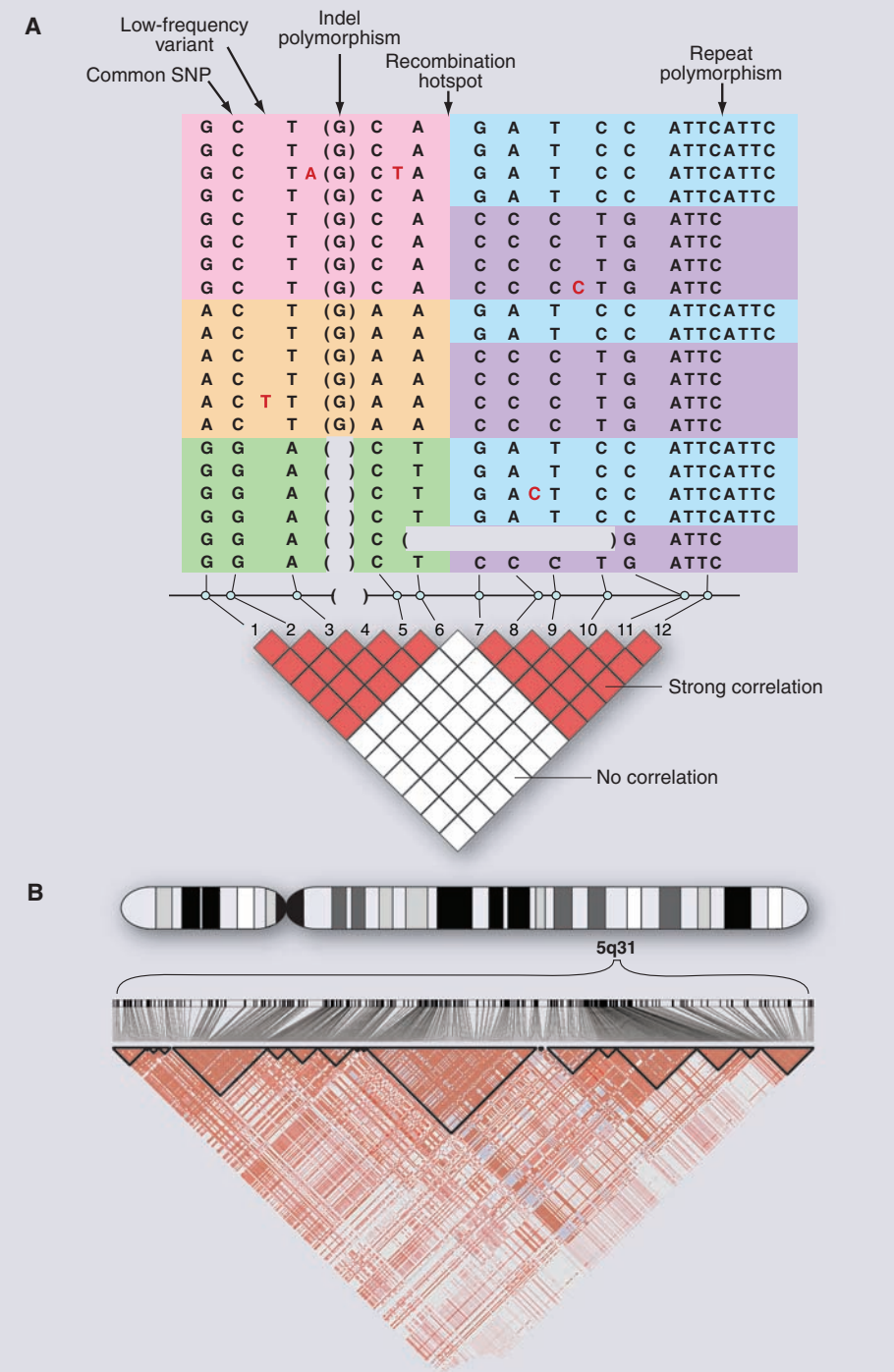


Fig. 1. DNA sequence variation in the human genome. **(A)** Common and rare genetic variation in 10 individuals, carrying 20 distinct copies of the human genome. The amount of variation shown here is typical for a 5-kb stretch of genome and is centered on a strong recombination hotspot. The 12 common variations include 10 SNPs, an insertion-deletion polymorphism (indel), and a tetranucleotide repeat polymorphism. The six common polymorphisms on the left side are strongly correlated. Although these six polymorphisms could theoretically occur in 2^6 possible patterns, only three patterns are observed (indicated by pink, orange, and green). These patterns are called haplotypes. Similarly, the six common polymorphisms on the right side are strongly correlated and reside on only two haplotypes (indicated by blue and purple). The haplotypes occur because there has not been much genetic recombination between the sites. By contrast, there is little correlation between the two groups of polymorphisms, because a hotspot of genetic recombination lies between them. The pairwise correlation between the common sites is shown by the red and white boxes below, with red indicating strong correlation and white indicating weak correlation. In addition to the common polymorphisms, lower-frequency polymorphisms also occur in the human genome. Five rare SNPs are shown, with the variant nucleotide marked in red and the reference nucleotide not shown. In addition, on the second to last chromosome, a larger deletion variant is observed that removes several kilobases of DNA. Such larger deletion or duplication events (i.e., CNVs) may be common and segregate as other DNA variants. **(B)** Small regions such as in (A) are often embedded in genomic regions with much greater extents of LD. The diagram shows actual data from the International HapMap Project, showing 420 genetic variants in a region of 500 kb on human chromosome 5q31. Positions of the variants are shown below. Blocks of strong correlation are indicated by the black outlines. Longer-range patterns are often more complex than shown in (A) because weaker recombination hotspots may reduce, but not completely eliminate, marker-to-marker correlation.

was not that all causal mutations at these genes should be common (to the contrary, a full spectrum of alleles is expected), only that some common variants exist and could be used to pinpoint loci for detailed study.

It took a decade to develop the tools and methods required to test the CD-CV hypothesis: (i) catalogs of millions of common variants in the human population, (ii) techniques to genotype these variants in studies with thousands of patients, and (iii) an analytical framework to distinguish true associations from noise and artifacts.

Cataloging SNPs and linkage disequilibrium. Pilot projects in the late 1990s showed that it was possible to identify thousands of SNPs and to perform highly multiplexed genotyping by means of DNA microarrays (35). A public-private partnership, the SNP Consortium, built an initial map of 1.4 million SNPs (32); this has grown to more than 10 million SNPs (36) and is estimated to contain 80% of all SNPs with frequencies of >10% (37).

As the SNP catalog grew, a critical question loomed: Would GWASs require directly testing each of the ~10 million common variants for association to disease? That is, if only 5% of variants were tested, would 95% of associations be missed? Or could a subset serve as reliable proxies for their neighbors? Experience from Mendelian diseases suggested that substantial efficiencies might be possible. Each disease-causing mutation arises on a particular copy of the human genome and bears a specific set of common alleles in cis at nearby loci, termed a haplotype. Because the recombination rate is low [~1 crossover per 100 megabases (Mb) per generation], disease alleles in the population typically show association with nearby marker alleles for many generations, a phenomenon termed linkage disequilibrium (LD) (Fig. 1).

Early studies had demonstrated LD of nearby polymorphisms at the globin locus (38), which proved useful in tracking sickle-cell mutation. In the mid-1980s, it was proposed that a genome-wide search might be performed in genetically isolated populations, scanning the genome for a haplotype shared among unrelated patients carrying the same founder mutation (39). Such "LD mapping" in essence treated the entire population as a very large and very old extended family. This method soon proved useful in fine-mapping the founder $\Delta 508$ mutation in the transmembrane conductance regulator CFTR as a cause of cystic fibrosis (40) and in screening the entire genome in isolated populations such as Finland (41).

The key question was whether the same approach could be used more generally to study common alleles in large human populations, where recombination had more time to whittle down haplotypes. A simulation study suggested that LD might typically be too short to be useful, with a SNP every 5 kb (500,000 SNPs across the genome) providing very weak LD (average correlation $r^2 = 0.1$) (42). Studies of individual loci showed great heterogeneity in local LD (43).

As denser genetic maps became available, a clear picture emerged. Nearby variants were

observed to form a block-like structure consisting of regions characterized by little evidence for historical recombination and limited haplotype diversity (44, 45). Within such regions, which soon proved general (46), genotypes of common SNPs could be inferred from knowledge of only a few empirically determined tag SNPs (45–47). These patterns were shaped by hot and cold spots of recombination in the human genome (48–50), as well as historical population bottlenecks (51).

The International HapMap Project was launched in 2002, with the goal of characterizing SNP frequencies and local LD patterns across the human genome in 270 samples from Europe, Asia, and West Africa. The project genotyped ~1 million SNPs by 2005 (37) and more than 3 million by 2007 (52). Sequence data collected by the project confirmed that the vast majority of common SNPs are strongly correlated to one or more nearby proxies: 500,000 SNPs provide excellent power to test >90% of common SNP variation in out-of-Africa populations, with roughly twice that number required in African populations (37).

Massively parallel genotyping. SNP genotyping was initially performed one SNP at a time, at a cost of ~\$1 per measurement. Multiplex genotyping of hundreds of SNPs on DNA microarrays was demonstrated in 1998 (35), and capacity per array grew from 10,000 to 100,000 SNPs in 2002 to 500,000 to 1 million SNPs in 2007. In parallel, cost fell to \$0.001 per genotype, or less than \$1000 per sample for a whole-genome analysis. By 2006, several technologies could simultaneously genotype hundreds of thousands of SNPs at >99% completeness and >99% accuracy.

Copy-number variation. SNPs are only one type of genetic variation (Fig. 1). Using microarray technology, two groups in 2004 observed that individual copies of the human genome contain large regions (tens to hundreds of kilobases in size) that are deleted, duplicated, or inverted relative to the reference sequence (53, 54). Structural variants had been previously associated with developmental disorders and were often assumed to be pathogenic; the presence of so many segregating copy-number variations (CNVs) in the general population was surprising. The generality of CNVs was soon established (55–59). Many CNVs display tight LD with nearby SNPs (56, 57) and thus can be proxied by nearby SNPs in GWASs. Others occur in regions that are difficult to follow with SNPs, are highly mutable, or are rare (58, 59). Hybrid genotyping platforms have recently been developed to genotype SNPs and CNVs simultaneously (60).

Statistical analysis. Recognizing causal loci amid a genome's worth of random fluctuation required advances in statistical design, analysis, and interpretation. The risk of false negatives was illustrated by a study of type 2 diabetes (T2D) and the Pro¹² → Ala polymorphism in peroxisome proliferator-activated receptor γ . Whereas an initial positive report (61) had not been confirmed in four modest-sized replication studies, larger studies produced strong and consistent evidence of increased risk by a factor of 1.2 (62, 63). The negative studies were actually consistent with

the level of increased risk, but simply lacked adequate power to detect it.

Conversely, stringent thresholds for statistical significance are needed to avoid false positives due to multiple hypothesis testing. Simulations indicated that a dense genome-wide scan of common variants involves the equivalent of ~1 million independent hypotheses (64). A significance level of $P = 5 \times 10^{-8}$ thus represents a finding expected by chance once in 20 GWASs. Large sample sizes would be needed to reach such a stringent threshold (Fig. 2).

Systematic biases could also cause false positives. Differences in ancestry between cases and controls would yield spurious associations (65), suggesting the need for family-based controls (66). It was later recognized that genome-wide studies provide their own internal control: Mismatched ancestry is readily detectable because it produces frequency differences at thousands of SNPs, which could not all reflect causal associations. Methods were developed to detect and adjust for such biases (67–69) as well as unexpected relatedness between subjects. Technical artifacts, which are particularly problematic if cases and controls are not genotyped in parallel (70), were overcome by improved genotyping methods, quality control, and stringent filtering. To maximize efficiency and power, several groups developed methods of selecting tag SNPs (47, 71–73) from empirical LD data and using them to impute genotypes at other SNPs not genotyped in clinical samples (74) on the basis of LD relationships in the HapMap.

Genome-Wide Associations: Lessons

By early 2006, the tools were in place and studies were under way in many laboratories to resolve the hotly debated issue (75, 76) of whether genetic mapping of common SNPs would shed light on common disease. Since then, scores of publications have reported the localization of common SNPs associated with a wide range of common diseases and clinical conditions (age-related macular degeneration, type 1 and type 2 diabetes, obesity, inflammatory bowel disease, prostate cancer, breast cancer, colorectal cancer, rheumatoid arthritis, systemic lupus erythematosus, celiac disease, multiple sclerosis, atrial fibrillation, coronary disease, glaucoma, gallstones, asthma, and restless leg syndrome) as well as various individual traits (height, hair color, eye color, freckles, and HIV viral set point). Figure 3 illustrates data from a paradigmatic genome-wide association study of Crohn's disease performed by the Wellcome Trust Case Control Consortium.

Various lessons have already emerged about genetic mapping by GWAS:

1) GWASs work. Before 2006, only about two dozen reproducible associations outside the HLA locus had been discovered (25). By early 2008, more than 150 relationships were identified between common SNPs and disease traits (table S1). In most diseases studied, GWASs have revealed multiple independent loci, although some traits have not yet yielded associations that meet stringent thresholds (e.g., hypertension). It is not clear whether this

reflects inadequate sample size, phenotypic definition, or a different genetic architecture.

2) Effect sizes for common variants are typically modest. In a few cases, common variants with effects of a factor of ≥ 2 per allele have been found: APOE4 in Alzheimer's disease (23), CFH in age-related macular degeneration (77–79), and LOXL1 in exfoliative glaucoma (80). In the vast majority of cases, however, the estimated effects are much smaller—mostly increases in risk by a factor of 1.1 to 1.5 per associated allele.

3) The power to detect associations has been low. Given the effect sizes now known to exist, and the need to exceed stringent statistical thresholds, the first wave of GWASs provided low power

and height (87–90). Across these four traits and diseases, individual GWASs together documented 29 associations. Increasing the power by pooling the samples to perform meta-analysis and replication genotyping has increased this yield to more than 100 replicated loci for these four conditions.

4) Association signals have identified small regions for study but have not yet identified causal genes and mutations. Genetic mapping is a double-edged sword: Local correlation of genetic variants facilitates the initial identification of a region but makes it difficult to distinguish causal mutation(s). Luckily, whereas family-based linkage methods typically yield regions of 2 to 10 Mb in span, GWASs typically yield more manageable regions of 10 to 100 kb.

already identified seven independent alleles at 8q24 for prostate cancer (92), three at complement factor H (CFH) for age-related macular degeneration (93, 94), three at IRF5 for systemic lupus erythematosus (95), and two at IL23R for Crohn's disease (96). Multiple distinct alleles with different frequencies and risk ratios may well be the rule.

6) A single locus can harbor both common variants of weak effect and rare variants of large effect. In recent GWASs, studies of common SNPs enabled the identification of 19 loci as influencing low- or high-density lipoprotein (LDL, HDL) or triglycerides (84, 85). Nine of these 19 were already known to carry rare Mendelian mutations with large effects, such as the loci for the LDL receptor (LDLR) and familial hypercholesterolemia (FH). Similarly, the genes encoding Kir6.2, WFS1, and TCF2 are all known to cause Mendelian syndromes including T2D, as well as common SNPs with modest effects.

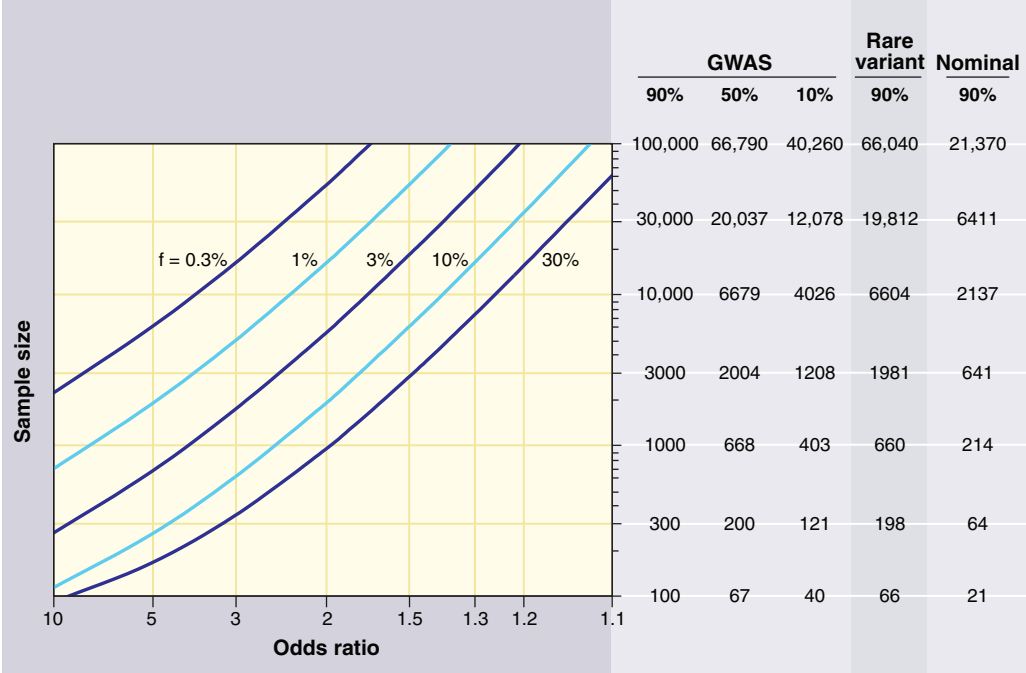


Fig. 2. Sample sizes required for genetic association studies. The graphs show the total number N of samples (consisting of $N/2$ cases and $N/2$ controls) required to map a genetic variant as a function of the increased risk due to the disease-causing allele (x axis) and the frequency of the disease-causing allele (various curves). The required sample size is shown in the table on the right for various different kinds of association studies. The first three columns pertain to GWASs using common variants across the entire genome; the columns correspond to different levels of statistical power to achieve a significant result at $P < 10^{-8}$. The fourth column pertains to a search for rare variants where the frequency listed is the collective frequency of rare variants in controls, and the odds ratio is the excess in cases as compared to controls. Sample sizes assume correction for a genome-wide search of $\sim 20,000$ protein-coding genes in the genome (aiming to achieve $P < 10^{-5}$ with one test performed per gene). The fifth column pertains to a test of a single hypothesis (e.g., testing association with a single SNP). For example, in a GWAS, 1000 samples provide 90% statistical power to detect a 30% allele with a factor of 2 effect. In a genome-wide search via exon sequencing, 660 samples provide 90% power to detect a gene in which rare variants have aggregate population frequency 1% and convey a factor of ~ 8 increase in risk. Note that the sample size to test essentially all common SNPs in the human genome is only 5 times the sample size to test a single SNP.

to discover disease-causing loci (81, 82). For example, achieving 90% power to detect an allele with 20% frequency and a factor of 1.2 effect at a statistical significance of 10^{-8} requires 8600 samples (Fig. 2). Thus, although it is unlikely that common alleles of large effect have been missed, GWASs of hundreds to several thousand cases have necessarily identified only a fraction of the loci that can be found with larger sample sizes. This prediction has been empirically confirmed in T2D (83), serum lipids (84, 85), Crohn's disease (86),

These regions have yet to be scrutinized by fine-mapping and resequencing to identify the specific gene and variants responsible. Even when a locus is identified by SNP association, the causal mutation itself need not be a SNP. For example, the *IRGM* gene was associated with Crohn's disease on the basis of GWAS. Subsequent study suggests that the causal mutation is a deletion upstream of the promoter affecting tissue-specific expression (91).

5) A single locus can contain multiple independent common risk variants. Intensive study has

GWAS to carry common genetic variation influencing LDL levels (84, 85). Similarly, SNPs in the β -cell zinc transporter encoded by *SLC30A8* were associated with risk of T2D (97).

2) Most associations do not involve previous candidate genes. In some cases, GWAS results immediately suggest new biological hypotheses—for example, the role of complement factor H in age-related macular degeneration (77–79), *FGFR2* in breast cancer (98), and *CDKN2A* and *CDKN2B* in T2D (99–101). In many other cases, such as

LOC387715/HTRA1 with age-related macular degeneration (102), nearby genes have no known function.

3) Many associations implicate non-protein-coding regions. Although some associated non-coding SNPs may ultimately prove attributable to LD with nearby coding mutations, many are sufficiently far from nearby exons to make this outcome unlikely. Examples include the region at 8q24 associated with prostate, breast, and colon cancer, 300 kb from the nearest gene (103, 104), and the region at 9q21 associated with myocardial infarction and T2D, 150 kb from the nearest genes encoding CDKN2A and CDKN2B (99–101, 105–107).

A role for noncoding sequence in disease risk is not surprising: Comparative genome analysis has shown that 5% of the human genome is evolutionarily conserved and thus functional; less than one-third of this 5% consists of genes that encode proteins (108). Noncoding mutations with roles in disease susceptibility will likely open new doors to understanding genome biology and gene regulation. Regulatory variation also suggests different therapeutic strategies: Modulating levels of gene expression may prove more tractable than replacing a fully defective protein or turning off a gain-of-function allele.

4) Some regions contain expected associations across diseases and traits. Crohn's disease, psoriasis, and ankylosing spondylitis have long been recognized to share clinical features; the association of the same common polymorphisms in IL23R in all three diseases points to a shared molecular cause (96, 109, 110). SNPs in STAT4 (signal transducer and activator of transcription 4) are associated with rheumatoid arthritis and systemic lupus, two diseases that share clinical features. Multiple variants associated with T2D are associated with insulin secretion defects in nondiabetic individuals (101, 111–116), highlighting the role of β -cell failure in the pathogenesis of T2D.

5) Some regions reveal surprising associations. For example, unexpected connections have emerged among T2D, inflammatory diseases (two loci), and cancer (four loci). A single intron of CDKAL1 was found to contain a SNP associated with T2D and insulin secretion defects (99–101, 116), and another with Crohn's disease and psoriasis (117). A coding variant in glucokinase regulatory protein is associated with triglyceride levels and fasting glucose (101) but also with C reactive protein levels (118, 119) and Crohn's disease (86). A SNP in TCF2 is associated with protection from T2D, as expected on the basis of Mendelian mutations at the same gene (120). Unexpectedly, the same association signal turned up in a GWAS for prostate cancer (121). Similarly, JAZF1 was identified as containing SNPs associated with T2D (83) and prostate cancer (122), and TCF7L2 with T2D (123) and colon cancer (124, 125).

From Common SNPs to the Full Allelic Spectrum

The current HapMap provides reliable proxies for the vast majority of SNPs at frequencies above

5%, but its coverage declines rapidly for lower-frequency alleles (37). Such lower-frequency alleles may be particularly important: Alleles with strong deleterious effects are constrained by natural selection from becoming too common. We divide these alleles into two conceptually distinct classes:

1) Common variants with frequencies below 5%. By "common," we refer to variants that occur at sufficient frequency to be cataloged in studies of the general population and measured (directly, or indirectly through LD) in association studies. In practice, this class may include allele frequencies in the range of 0.5% and above. A GWAS of 2000 cases and 2000 controls provides good power for a 1% allele causing a factor of 4 increase in risk (even at $P < 10^{-8}$) (Fig. 2).

The value of lower-frequency common variants is illustrated by PCSK9 (proprotein convertase subtilisin/kexin type 9). The gene encoding PCSK9 contains very rare mutations causing autosomal dominant hypercholesterolemia (discovered by linkage analysis), as well as high-frequency common variants with modest effects. The former are too rare and the latter too weak to enable effective clinical study of PCSK9 with respect to coronary artery disease risk. Hobbs and Cohen sequenced the gene (126, 127) and identified low-frequency common variants (0.5 to 1%), which allowed epidemiological research documenting a protective effect on myocardial infarction (128).

2) Rare variants. Most Mendelian diseases involve rare mutations that are essentially never observed in the general population. Rare mutations likely also play an important role in common diseases. Because they are numerous and individually rare, it is not possible to create a complete catalog in the general population. Instead, they must be identified by sequencing in cases and controls in each study. Moreover, because each variant is too rare to prove statistical evidence of association, the mutations must be aggregated as a class to compare the overall frequency of cases versus controls.

A few examples are known through candidate gene studies. Rare nonsynonymous mutations in MC4R are found in patients with extreme early-onset obesity (129). Rare nonsynonymous mutations in ABCA1 are more common in patients with extremely low HDL than in those with high HDL (130). An excess of rare mutations in renal salt-handling genes has been associated with lower blood pressure and protection against hypertension (131).

The sample size required to perform a genome-wide search based on coding mutations depends on the background frequency (μ) of mutations that confer disease risk and the level (ω) of increased risk for each such mutation. ABCA1 is a favorable case because μ and ω are high (the gene has an unusually large coding region of ~ 7 kb, and mutations confer a factor of ~ 6 increase in risk). Achieving genome-wide significance will likely require resequencing studies of thousands of cases and controls, similar to GWASs (Fig. 2).

GWASs of rare variants are already under way for large structural variants through the use of microarray analysis. A recent GWAS of autism revealed

that a highly penetrant, recurrent microdeletion and microduplication of a 593-kb region in 16p11.2 explains 1% of cases (132). Moreover, several recent studies report that patients with autism and schizophrenia may have an excess of rare deletions across the genome relative to unaffected controls (133, 134). Although these studies did not identify specific loci (none of the novel loci were observed more than once), they suggest that the universe of rare structural changes contributing to each disease may be as large and diverse as that of common SNPs.

The Genetic Architecture of Common Disease

Variants so far identified by GWASs together explain only a small fraction of the overall inherited risk of each disease (for example, $\sim 10\%$ of the variance for Crohn's and $\sim 5\%$ for T2D). Where is the remaining genetic variance to be found? There are several answers:

1) At disease loci already identified by GWAS, the locus-attributable risk will often be higher than currently estimated. This is because marker SNPs used in GWASs will typically be imperfect proxies for the actual causal mutation that led to the association signal. The causal gene will often contain additional mutations not tagged by the initial marker SNPs, both common and rare. Determining the contribution of each gene will require intensive studies of variants at each locus.

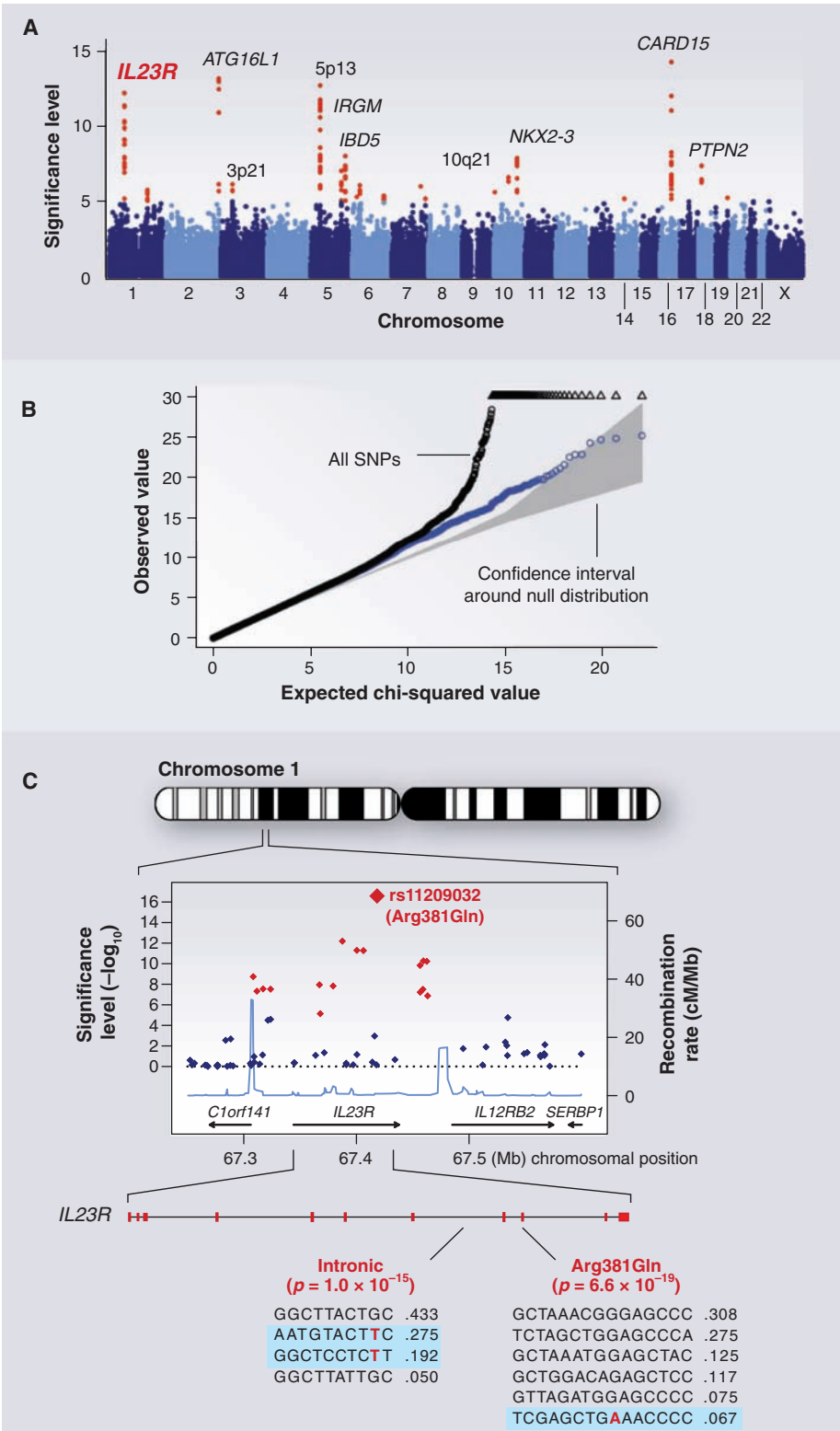
2) Many more disease loci remain to be identified by GWAS. As noted above, GWASs to date have had low statistical power and thus necessarily missed many loci with common variants of similar and smaller effects. The first studies did not have proxies for common structural variants and have failed to capture lower-frequency common variants (0.5 to 5%). Moreover, the vast majority of studies have been performed only in samples of European ancestry. Larger, more comprehensive, and more diverse GWASs will reveal many more loci.

3) Some disease loci will contain only rare variants. Such loci (if not already found by Mendelian genetics) cannot be identified by study of common variants alone. They will require systematic resequencing of all genes in large samples (Fig. 2).

4) Current estimates of the variance explained are based on simplifying assumptions. Because the genotype-phenotype correlation has yet to be well characterized, the estimates assume that the variants interact in a simple additive manner. Yet gene-gene and gene-environment interactions play important roles in disease risk. Although searches have not yet found much evidence for epistasis [e.g., (93, 94, 135)], this may simply reflect limited power to assess the many possible modes of interaction, including pairwise interactions and threshold effects. Once patterns of association and interaction are understood, effects of specific gene and environmental exposures on each phenotype may be larger.

For these reasons, it is premature to make inferences about the overall genetic architecture of common disease. Only by systematically exploring each of these directions over the coming years will a general picture emerge—with the likely

Fig. 3. GWAS for Crohn’s disease. The panels show data from the study of Crohn’s disease by the Wellcome Trust Case Control Consortium. **(A)** Significance level (P value on $\log_{10}$ scale) for each of the 500,000 SNPs tested across the genome. SNP locations reflect their positions across the 23 human chromosomes. SNPs with significance levels exceeding 10^{-5} (corresponding to 5 on the y axis) are colored red; the remaining SNPs are in blue. Ten regions with multiple significant SNPs are shown, labeled by their location or by the likely disease-related gene (e.g., *IL23R* on chromosome 1). **(B)** The fact that the SNPs in red are extreme outliers is made clear from a so-called Q-Q plot. A Q-Q plot is made as follows: The SNPs are ordered (from 1 to n) according to their observed P values; observed and expected P values are plotted for each SNP. Under the null distribution, the expected P value for the i th SNP is i/n . If there are no significant associations, the Q-Q plot will lie along the 45° line; the gray region corresponds to a 95% confidence region around this null expectation. Black points correspond to all 500,000 SNPs studied that passed strict quality control; they diverge strongly from the null expectation. Blue points reflect the P values that remain when the SNPs in the 10 most significant regions are removed; there is still some excess of significant P values, indicating the presence of additional loci of more modest effect. **(C)** Close-up of the region around the *IL23R* locus on chromosome 1. The first part shows the significance levels for SNPs in a region of ~400 kb, with colors as in (A). The highest significance level occurs at a SNP in the coding region of the *IL23R* gene (causing an Arg³⁸¹ → Gln change). The light blue curve shows the inferred local rate of recombination across the region. There are two clear hotspots of recombination, with SNPs lying between these hotspots being strongly correlated in a few haplotypes. The second part shows that the *IL23R* locus harbors at least two independent, highly significant disease-associated alleles. The first site is the Arg³⁸¹ → Gln polymorphism, which has a single disease-associated haplotype (shaded in blue) with frequency of 6.7%. The second site is in the intron between exons 7 and 8; it tags two disease-associated haplotypes with frequencies of 27.5% and 19.2%.



outcome being that different diseases will each be characterized by a different balance of allele frequencies, interactions, and types. Although the proportion of genetic variance explained is certain to grow in the coming years, it is unlikely to approach 100% because of practical limitations, such as the difficulty of detecting common variants with extremely small effects, genes harboring rare var-

iants at very low frequency, and complex interactions among genes and with the environment.

Disease Risk Versus Disease Mechanism

The primary value of genetic mapping is not risk prediction, but providing novel insights about mechanisms of disease. Knowledge of disease pathways (not limited to the causal genes and mutations) can

suggest strategies for prevention, diagnosis, and therapy. From this perspective, the frequency of a genetic variant is not related to the magnitude of its effect, nor to the potential clinical value that may be obtained.

The classic example is Brown and Goldstein’s studies of FH, which affects ~0.2% of the population and accounts for a tiny fraction of the heritability of LDL and myocardial infarction. Studies of FH led

to the discovery of the LDL receptor and supported the development of HMGCR inhibitors (statins) for lowering LDL, the use of which is not limited to FH carriers.

More recently, GWASs have shown that common genetic variation in LDLR and HMGCR influences LDL levels (84, 85). Although SNPs in HMGCR have only a small effect (~5%) on LDL levels, drugs targeting the encoded protein decrease LDL levels by a much greater extent (~30%). This is because the effect of an inherited variant is limited by natural selection and pleiotropy, whereas the effect of a drug treatment is not.

The Path Ahead

Given the long-standing success of genetic mapping in providing new insights into biology and disease etiology, and the recent proof that systematic association studies can identify novel loci, our aim should be nothing less than identifying all pathways at which genetic variation contributes to common diseases. We sketch key steps in achieving this goal.

Expanding clinical studies. Current studies are underpowered for the types of SNP alleles that we now know exist, and available evidence indicates that increasing sample size will yield substantial returns. A study of 1000 cases and 1000 controls provides only 1% power to detect a 20% variant that increases risk by a factor of 1.3, but a study of 5000 cases and 5000 controls provides 98% power (Fig. 2). Moreover, early data on rare single-nucleotide (130, 131) and structural variants (133, 134, 136) indicate that similarly large samples will be needed to achieve the levels of statistical significance required to detect rare events in a genome-wide search.

Nearly all GWASs to date have been performed in populations of European ancestry. Even if a variant has the same effect in all ancestry groups, it may be more readily detected in one population simply because it happens to have higher frequency. Genetic effects will likely vary across groups because of modification by environment and behavior, which may vary more across groups than does genotype.

Many important diseases remain to be studied by GWAS. Disease-related intermediate traits can also offer substantial insight, particularly in conjunction with clinical endpoints. For example, newly described variants on chromosome 1 (near SORT1) are associated both with levels of LDL cholesterol (84, 85) and with risk of myocardial infarction (106); this provides not only increased statistical confidence, but also a biomarker for gene function and pathophysiological insight. Genetic variants that influence gene expression [e.g., (137)] hold promise for elucidating regulatory pathways. Mapping of modifiers of Mendelian mutations—for example, genes that influence the age of onset in carriers of BRCA1 and BRCA2 mutations—may suggest ways to reverse high risk due to mutations.

Correlations between genetic variants and phenotypes are limited by the accuracy with which

each is measured. The ability to measure genotype now far exceeds our ability to measure phenotype. Continuous ambulatory monitoring, imaging methods, and comprehensive (“-omic”) approaches to biological samples all have promise in improving the accuracy of phenotype measurement.

Environmental exposures play a larger role in human phenotypic variation than does genetic variation, but environmental exposures are fundamentally more difficult to measure. DNA is stable throughout life, with a single physical chemistry that enables generic approaches to measurement. Environmental exposures are heterogeneous and may be fleeting. Improved methods for measuring environmental exposures, perhaps based on epigenetic marks they leave, are sorely needed.

Expanding the range of genetic variation. The lowest-hanging fruit will be to resequence loci that have been definitively implicated in disease by Mendelian genetics or by GWAS. Because the prior probability of a true association is higher, such regions will be the best setting to develop methods for understanding the statistical significance and biological importance of rare mutations. Initially, resequencing of coding exons will be easiest to interpret. Rare coding mutations with large effect will be especially valuable, because physiological studies of mutation carriers can help illuminate the biological basis of the disease, and because coding mutations of large effect are more straightforwardly transferred to cellular and animal models for mechanistic studies.

Extending GWASs to include structural variants and lower-frequency common variants will require comprehensive catalogs of genomic variation, as well as characterization of LD relationships. With new massively parallel sequencing technologies, an accurate map of all 1% alleles (both single-nucleotide and structural) should be achievable. A “1000 Genomes Project” was recently launched toward this end (138).

Some loci may harbor neither common variants nor rare structural variants, and thus will be missed by array and LD-based approaches. Discovering such genes will require sequencing in thousands of cases and controls. Initial studies will likely focus on exons, where functional mutations are enriched to the greatest extent. Highly parallel methods to capture hundreds of thousands of exons, and other targets of interest, are under development (139).

Multiple instances of de novo coding mutations at a locus (by comparing affected individuals with parents) could provide particularly powerful association information, because the human mutation rate is so low (in the range of 10^{-8}). But identifying de novo mutations without being overwhelmed by false positives will require extraordinary sequencing accuracy (far better than finished genome sequence). Because such studies will be expensive at first, priority should go to disorders with high heritability, where there is an unmet medical need, and for which other approaches

have met with limited success. Psychiatric disorders might represent one such target.

Eventually, it will become practical to resequence entire genomes from thousands of cases and controls. The problem of interpretation will be much harder for noncoding functional elements, because it is unclear either how to aggregate elements to achieve a large enough target size, or to develop ways to recognize function-altering changes.

Routine genome sequencing of deeply phenotyped cohorts will fundamentally change the nature of genetic mapping: from the current serial process (in which initial localization by linkage or GWAS is followed by scrutiny of DNA variation and phenotypes) to a joint estimation procedure combining variation information of all types, frequencies, and phenotypes to discover and characterize genotype-phenotype correlations. New statistical methods will be required to combine evidence from rare and common alleles at a locus and across multiple loci, phenotypes, and non-genetic exposures. A particular challenge will be to identify mutations in regions without known function or evolutionary conservation.

There may be inherent limits to our ability to relate phenotypic variation and genotypic variation. To the extent that disease is influenced by tiny effects at hundreds of loci or highly heterogeneous rare mutations, it may be impractical to assemble sufficiently large samples to give a complete accounting.

Implications for Biology, Medicine, and Society

Genetic mapping is only a first step toward biological understanding and clinical application. Useful tools will include maps of evolutionary conservation (108) and chromatin state (140), as well as databases of cell-state signatures, such as genome-wide expression patterns, that may integrate aspects of cell biology under resting and provoked conditions (141). Creation of disease models, both in human cell culture and nonhuman animals, will be key. Physiological studies in patients classified by genotype may inform disease processes and lead to useful nongenetic biomarkers. Given the limits of human clinical research, rare alleles of strong effect may be more useful than common alleles of weak effect.

The high failure rate of clinical trials testifies to the limited predictive value of current approaches. By focusing attention on genes and processes, human genetics has the potential to yield productive targets and predictive animal models. In clinical trials, the ability to stratify patients by genotype or biological pathway may reveal differences in therapeutic response. Genetics may also increase the efficiency of outcome trials by focusing on patients at higher-than-average risk.

The extent to which genetic information will figure in “personalized medicine” will depend on whether predictive accuracy beyond conventional measures can be attained, and whether there are interventions whose effectiveness is improved by knowledge of a genetic test. Knowledge of a common variant that increases T2D

risk by 20% may eventually lead to new understanding and therapeutic strategies, but whether an increase in absolute risk (from 8% to 10%) is useful for patients remains to be seen. Although it is tempting to think that knowledge of individual risk might promote greater adherence to a healthy lifestyle, human behavior is complex and risk estimates are challenging to interpret. Even where genotype can predict response to a drug with a narrow therapeutic window, it cannot be assumed that genetic testing will necessarily lead to improved clinical outcomes.

Our understanding of complex disease will be in constant flux over the coming years. The pace of discovery, while scientifically exhilarating, poses daunting challenges. Direct-to-consumer marketing of genetic information is already under way. It will be a challenge for the public to understand the difference between relative and absolute risk, and to figure in their thinking the larger component of genetic and environmental factors not yet captured by today's technologies. Rigorous assessment of health benefit and cost are needed, including costs of testing and treatment that may flow from an altered sense of risk. As genetic information is shown to be useful, equitable access will be critical.

Finally, we must ensure that the promise of research on genetic factors in complex disease does not encourage a mistaken sense of genetic determinism. This is especially important for behavioral traits, which are especially prone to misinterpretation and misguided policy. We must constantly remind the public—and ourselves—that although genes play a role (and can lead us to new biological insight), our traits are powerfully shaped by the environment, and the solutions to important problems will often lie outside our genes.

References

1. A. Sturtevant, *J. Exp. Zool.* **14**, 43 (1913).
2. L. Clarke, J. Carbon, *Proc. Natl. Acad. Sci. U.S.A.* **77**, 2173 (1980).
3. W. Bender *et al.*, *Science* **221**, 23 (1983).
4. I. Aird, H. H. Bental, J. A. Mehigan, J. A. Roberts, *Br. Med. J.* **2**, 315 (1954).
5. V. M. Ingram, *Nature* **178**, 792 (1956).
6. T. D. Petes, D. Botstein, *Proc. Natl. Acad. Sci. U.S.A.* **74**, 5091 (1977).
7. A. J. Jeffreys, *Cell* **18**, 1 (1979).
8. Y. W. Kan, A. M. Dozy, *Proc. Natl. Acad. Sci. U.S.A.* **75**, 5631 (1978).
9. D. Botstein, R. L. White, M. Skolnick, R. W. Davis, *Am. J. Hum. Genet.* **32**, 314 (1980).
10. J. F. Gusella *et al.*, *Nature* **306**, 234 (1983).
11. H. Donis-Keller *et al.*, *Cell* **51**, 319 (1987).
12. C. Dib *et al.*, *Nature* **380**, 152 (1996).
13. T. J. Hudson *et al.*, *Science* **270**, 1945 (1995).
14. Online Mendelian Inheritance in Man (www.ncbi.nlm.nih.gov/sites/entrez?db=omim).
15. P. L. Welch, M. C. King, *Hum. Mol. Genet.* **10**, 705 (2001).
16. R. P. Lifton, *Harvey Lect.* **100**, 71 (2004).
17. G. I. Bell, K. S. Polonsky, *Nature* **414**, 788 (2001).
18. E. East, *Am. Nat.* **44**, 65 (1910).
19. E. Altenburg, H. J. Muller, *Genetics* **5**, 1 (1920).
20. R. A. Fisher, *Trans. R. Soc. Edinburgh* **52**, 399 (1918).
21. A. H. Paterson *et al.*, *Nature* **335**, 721 (1988).
22. J. Klein, A. Sato, *N. Engl. J. Med.* **343**, 782 (2000).
23. W. J. Strittmatter, A. D. Roses, *Annu. Rev. Neurosci.* **19**, 53 (1996).
24. J. N. Hirschhorn, K. Lohmueller, E. Byrne, K. Hirschhorn, *Genet. Med.* **4**, 45 (2002).
25. K. E. Lohmueller, C. L. Pearce, M. Pike, E. S. Lander, J. N. Hirschhorn, *Nat. Genet.* **33**, 177 (2003).
26. F. S. Collins, M. S. Guyer, A. Chakravarti, *Science* **278**, 1580 (1997).
27. E. S. Lander, *Science* **274**, 536 (1996).
28. N. Risch, K. Merikangas, *Science* **273**, 1516 (1996).
29. M. Kimura, T. Ota, *Genetics* **75**, 199 (1973).
30. H. Harris, *Proc. R. Soc. London Ser. B* **164**, 298 (1966).
31. W. H. Li, L. A. Sadler, *Genetics* **129**, 513 (1991).
32. R. Sachidanandam *et al.*, *Nature* **409**, 928 (2001).
33. R. Lewontin, in *Evolutionary Biology* **6**, T. Dobzhansky, M. K. Hecht, W. C. Steere, Eds. (Appleton-Century-Crofts, New York, 1972), pp. 391–398.
34. D. E. Reich, E. S. Lander, *Trends Genet.* **17**, 502 (2001).
35. D. G. Wang *et al.*, *Science* **280**, 1077 (1998).
36. Entrez SNP (www.ncbi.nlm.nih.gov/sites/entrez?db=snp).
37. International HapMap Consortium, *Nature* **437**, 1299 (2005).
38. S. E. Antonarakis, C. D. Boehm, P. J. Giardina, H. H. Kazazian Jr., *Proc. Natl. Acad. Sci. U.S.A.* **79**, 137 (1982).
39. E. S. Lander, D. Botstein, *Cold Spring Harb. Symp. Quant. Biol.* **51**, 49 (1986).
40. B. Kerem *et al.*, *Science* **245**, 1073 (1989).
41. J. Hastbacka *et al.*, *Nat. Genet.* **2**, 204 (1992).
42. L. Kruglyak, *Nat. Genet.* **22**, 139 (1999).
43. K. G. Ardlie, L. Kruglyak, M. Seielstad, *Nat. Rev. Genet.* **3**, 299 (2002).
44. M. J. Daly, J. D. Rioux, S. F. Schaffner, T. J. Hudson, E. S. Lander, *Nat. Genet.* **29**, 229 (2001).
45. N. Patil *et al.*, *Science* **294**, 1719 (2001).
46. S. B. Gabriel *et al.*, *Science* **296**, 2225 (2002); published online 23 May 2002 (10.1126/science.1069424).
47. G. C. Johnson *et al.*, *Nat. Genet.* **29**, 233 (2001).
48. D. E. Reich *et al.*, *Nat. Genet.* **32**, 135 (2002).
49. D. C. Crawford *et al.*, *Nat. Genet.* **36**, 700 (2004).
50. G. A. T. McVean *et al.*, *Science* **304**, 581 (2004).
51. S. A. Tishkoff, B. C. Verrelli, *Annu. Rev. Genomics Hum. Genet.* **4**, 293 (2003).
52. International HapMap Consortium, *Nature* **449**, 851 (2007).
53. J. Sebat *et al.*, *Science* **305**, 525 (2004).
54. A. J. lafrate *et al.*, *Nat. Genet.* **36**, 949 (2004).
55. E. Tuzun *et al.*, *Nat. Genet.* **37**, 727 (2005).
56. D. A. Hinds, A. P. Kloek, M. Jen, X. Chen, K. A. Frazer, *Nat. Genet.* **38**, 82 (2006).
57. S. A. McCarroll *et al.*, *Nat. Genet.* **38**, 86 (2006).
58. D. P. Locke *et al.*, *Am. J. Hum. Genet.* **79**, 275 (2006).
59. R. Redon *et al.*, *Nature* **444**, 444 (2006).
60. S. A. McCarroll *et al.*, *Nat. Genet.* **40**, 1166 (2008).
61. S. E. Deeb *et al.*, *Nat. Genet.* **20**, 284 (1998).
62. D. Altshuler *et al.*, *Nat. Genet.* **26**, 76 (2000).
63. J. C. Florez, J. N. Hirschhorn, D. Altshuler, *Annu. Rev. Genomics Hum. Genet.* **4**, 257 (2003).
64. I. Pe'er, R. Yelensky, D. Altshuler, M. J. Daly, *Genet. Epidemiol.* **32**, 381 (2008).
65. W. C. Knowler, R. C. Williams, D. J. Pettitt, A. G. Steinberg, *Am. J. Hum. Genet.* **43**, 520 (1988).
66. R. S. Spielman, R. E. McGinnis, W. J. Ewens, *Am. J. Hum. Genet.* **52**, 506 (1993).
67. B. Devlin, K. Roeder, *Biometrics* **55**, 997 (1999).
68. J. K. Pritchard, N. A. Rosenberg, *Am. J. Hum. Genet.* **65**, 220 (1999).
69. A. L. Price *et al.*, *Nat. Genet.* **38**, 904 (2006).
70. D. G. Clayton *et al.*, *Nat. Genet.* **37**, 1243 (2005).
71. J. M. Chapman, J. D. Cooper, J. A. Todd, D. G. Clayton, *Hum. Hered.* **56**, 18 (2003).
72. D. A. Hinds *et al.*, *Science* **307**, 1072 (2005).
73. P. I. W. de Bakker *et al.*, *Nat. Genet.* **37**, 1217 (2005).
74. J. Marchini, B. Howie, S. Myers, G. McVean, P. Donnelly, *Nat. Genet.* **39**, 906 (2007).
75. K. M. Weiss, J. D. Terwilliger, *Nat. Genet.* **26**, 151 (2000).
76. J. Couzin, *Science* **296**, 1391 (2002).
77. R. J. Klein *et al.*, *Science* **308**, 385 (2005); published online 10 March 2005 (10.1126/science.1109557).
78. A. O. Edwards *et al.*, *Science* **308**, 421 (2005); published online 10 March 2005 (10.1126/science.1110189).
79. J. L. Haines *et al.*, *Science* **308**, 419 (2005); published online 10 March 2005 (10.1126/science.1110359).
80. G. Thorleifsson *et al.*, *Science* **317**, 1397 (2007); published online 9 August 2007 (10.1126/science.1146554).
81. Wellcome Trust Case Control Consortium, *Nature* **447**, 661 (2007).
82. D. Altshuler, M. Daly, *Nat. Genet.* **39**, 813 (2007).
83. E. Zeggini *et al.*, *Nat. Genet.* **40**, 638 (2008).
84. S. Kathiresan *et al.*, *Nat. Genet.* **40**, 189 (2008).
85. C. J. Willer *et al.*, *Nat. Genet.* **40**, 161 (2008).
86. J. C. Barrett *et al.*, *Nat. Genet.* **40**, 955 (2008).
87. G. Lettre *et al.*, *Nat. Genet.* **40**, 584 (2008).
88. M. N. Weedon *et al.*, *Nat. Genet.* **40**, 575 (2008).
89. S. Sanna *et al.*, *Nat. Genet.* **40**, 198 (2008).
90. D. F. Gudbjartsson *et al.*, *Nat. Genet.* **40**, 609 (2008).
91. S. A. McCarroll *et al.*, *Nat. Genet.* **40**, 1107 (2008).
92. C. A. Haiman *et al.*, *Nat. Genet.* **39**, 638 (2007).
93. J. Maller *et al.*, *Nat. Genet.* **38**, 1055 (2006).
94. M. Li *et al.*, *Nat. Genet.* **38**, 1049 (2006).
95. R. R. Graham *et al.*, *Proc. Natl. Acad. Sci. U.S.A.* **104**, 6758 (2007).
96. R. H. Duerr *et al.*, *Science* **314**, 1461 (2006); published online 26 October 2006 (10.1126/science.1135245).
97. R. Sladek *et al.*, *Nature* **445**, 881 (2007).
98. D. F. Easton *et al.*, *Nature* **447**, 1087 (2007).
99. E. Zeggini *et al.*, *Science* **316**, 1336 (2007); published online 25 April 2007 (10.1126/science.1142364).
100. L. J. Scott *et al.*, *Science* **316**, 1341 (2007); published online 25 April 2007 (10.1126/science.1142382).
101. Diabetes Genetics Initiative of Broad Institute of Harvard and MIT, Lund University, and Novartis Institutes for BioMedical Research, *Science* **316**, 1331 (2007); published online 26 April 2007 (10.1126/science.1142358).
102. A. Rivera *et al.*, *Hum. Mol. Genet.* **14**, 3227 (2005).
103. L. T. Amundadottir *et al.*, *Nat. Genet.* **38**, 652 (2006).
104. M. L. Freedman *et al.*, *Proc. Natl. Acad. Sci. U.S.A.* **103**, 14068 (2006).
105. R. McPherson *et al.*, *Science* **316**, 1488 (2007); published online 2 May 2007 (10.1126/science.1142447).
106. N. J. Samani *et al.*, *N. Engl. J. Med.* **357**, 443 (2007).
107. A. Helgadóttir *et al.*, *Science* **316**, 1491 (2007); published online 2 May 2007 (10.1126/science.1142842).
108. R. H. Waterston *et al.*, *Nature* **420**, 520 (2002).
109. P. R. Burton *et al.*, *Nat. Genet.* **39**, 1329 (2007).
110. M. Cargill *et al.*, *Am. J. Hum. Genet.* **80**, 273 (2007).
111. J. C. Florez *et al.*, *N. Engl. J. Med.* **355**, 241 (2006).
112. N. Grarup *et al.*, *Diabetes* **56**, 3105 (2007).
113. L. Pascoe *et al.*, *Diabetes* **56**, 3101 (2007).
114. R. Saxena *et al.*, *Diabetes* **55**, 2890 (2006).
115. H. Staiger *et al.*, *PLoS One* **2**, e832 (2007).
116. V. Steinthorsdóttir *et al.*, *Nat. Genet.* **39**, 770 (2007).
117. N. Wolf *et al.*, *J. Med. Genet.* **45**, 114 (2008).
118. A. P. Reiner *et al.*, *Am. J. Hum. Genet.* **82**, 1193 (2008).
119. P. M. Ridker *et al.*, *Am. J. Hum. Genet.* **82**, 1185 (2008).
120. W. Winckler *et al.*, *Diabetes* **56**, 685 (2007).
121. J. Gudmundsson *et al.*, *Nat. Genet.* **39**, 977 (2007).
122. G. Thomas *et al.*, *Nat. Genet.* **40**, 310 (2008).
123. S. F. A. Grant *et al.*, *Nat. Genet.* **38**, 320 (2006).
124. A. Hazra *et al.*, *Cancer Causes Control* **19**, 975 (2008).
125. A. R. Folsom *et al.*, *Diabetes Care* **31**, 905 (2008).
126. I. K. Kotowski *et al.*, *Am. J. Hum. Genet.* **78**, 410 (2006).
127. J. Cohen *et al.*, *Nat. Genet.* **37**, 161 (2005).
128. J. C. Cohen, E. Boerwinkle, T. H. Mosley Jr., H. H. Hobbs, *N. Engl. J. Med.* **354**, 1264 (2006).
129. J. N. Hirschhorn, D. Altshuler, *J. Clin. Endocrinol. Metab.* **87**, 4438 (2002).
130. J. C. Cohen *et al.*, *Science* **305**, 869 (2004).
131. W. Ji *et al.*, *Nat. Genet.* **40**, 592 (2008).
132. L. A. Weiss *et al.*, *N. Engl. J. Med.* **358**, 667 (2008).
133. T. Walsh *et al.*, *Science* **320**, 539 (2008); published online 27 March 2008 (10.1126/science.1155174).
134. J. Sebat *et al.*, *Science* **316**, 445 (2007); published online 14 March 2007 (10.1126/science.1138659).
135. J. D. Rioux *et al.*, *Nat. Genet.* **39**, 596 (2007).
136. L. A. Weiss *et al.*, *N. Engl. J. Med.* **358**, 667 (2008).
137. V. Emilsson *et al.*, *Nature* **452**, 423 (2008).
138. 1000 Genomes (www.1000genomes.org).
139. T. J. Albert *et al.*, *Nat. Methods* **4**, 903 (2007).
140. T. S. Mikkelsen *et al.*, *Nature* **448**, 553 (2007).
141. J. Lamb *et al.*, *Science* **313**, 1929 (2006).

Supporting Online Material

www.sciencemag.org/cgi/content/full/322/5903/881/DC1
Table S1

10.1126/science.1156409