

The Structure of Haplotype Blocks in the Human Genome

Stacey B. Gabriel,¹ Stephen F. Schaffner,¹ Huy Nguyen,¹
 Jamie M. Moore,¹ Jessica Roy,¹ Brendan Blumenstiel,¹
 John Higgins,¹ Matthew DeFelice,¹ Amy Lochner,¹
 Maura Faggart,¹ Shau Neen Liu-Cordero,^{1,2} Charles Rotimi,³
 Adebawale Adeyemo,⁴ Richard Cooper,⁵ Ryk Ward,⁶
 Eric S. Lander,^{1,2} Mark J. Daly,¹ David Altshuler^{1,7*}

Haplotype-based methods offer a powerful approach to disease gene mapping, based on the association between causal mutations and the ancestral haplotypes on which they arose. As part of The SNP Consortium Allele Frequency Projects, we characterized haplotype patterns across 51 autosomal regions (spanning 13 megabases of the human genome) in samples from Africa, Europe, and Asia. We show that the human genome can be parsed objectively into haplotype blocks: sizable regions over which there is little evidence for historical recombination and within which only a few common haplotypes are observed. The boundaries of blocks and specific haplotypes they contain are highly correlated across populations. We demonstrate that such haplotype frameworks provide substantial statistical power in association studies of common genetic variation across each region. Our results provide a foundation for the construction of a haplotype map of the human genome, facilitating comprehensive genetic association studies of human disease.

Variation in the human genome sequence plays a powerful but poorly understood role in the etiology of common medical conditions. Because the vast majority of heterozygosity in the human population is attributable to common variants and because the evolutionary history of common human diseases (which determined the allele spectrum for causal alleles) is not yet known, one promising approach is to comprehensively test common genetic variation for association to medical conditions (1–3). This approach is increasingly practical because 4 million (4, 5) of the estimated 10 million (6) common single nucleotide polymorphisms (SNPs) are already known.

In designing and interpreting association studies of genotype and phenotype, it is necessary to understand the structure of haplotypes in the human genome. Haplotypes are the particular combinations of alleles observed in a pop-

ulation. When a new mutation arises, it does so on a specific chromosomal haplotype. The association between each mutant allele and its ancestral haplotype is disrupted only by mutation and recombination in subsequent generations. Thus, it should be possible to track each variant allele in the population by identifying (through the use of anonymous genetic markers) the particular ancestral segment on which it arose. Haplotype methods have contributed to the identification of genes for Mendelian diseases (7–9) and, recently, disorders that are both common and complex in inheritance (10–12). However, the general properties of haplotypes in the human genome have remained unclear.

Many studies have examined allelic associations [also termed “linkage disequilibrium” (LD)] across one or a few gene regions. These studies have generally concluded that linkage disequilibrium is extremely variable within and among loci and populations [reviewed in (13–15)]. Recently, examination of a higher density of markers over contiguous regions (16–18) suggested a surprisingly simple pattern: blocks of variable length over which only a few common haplotypes are observed punctuated by sites at which recombination could be inferred in the history of the sample. In one segment of the major histocompatibility complex (MHC) on chromosome 6, it has been directly demonstrated that “hotspots” of meiotic recombination coincided with boundaries between such blocks (17). These studies suggested a model for human haplotype structure but left many questions unanswered. First, how much of the hu-

varicocele ($n = 3$), or after surgery for a testicular tumor ($n = 2$). In all cases, histological examination indicated normal spermatogenesis. We analyzed karyotypes and conducted STS-based assays for Yq microdeletions on eight individuals and detected no abnormalities. Seminiferous tubules were prepared for immunolocalization and fluorescence in situ hybridization (FISH) as described (5). We conducted standard immunostaining (23) with CREST antiserum to detect kinetochores and antibodies against MLH1 and either of two SC proteins, SCP1 and SCP3. Pachytene cells were identified on a Zeiss epifluorescence microscope, and images were captured with an Applied Imaging Quips Pathvision System. For subsequent FISH, we used paint probes (Vysis) to identify chromosomes 1, 16, 21, and 22. Cells with previously captured images were relocated on the microscope, and new images were captured. We attempted to analyze 50 cells per individual, scoring the number of MLH1 foci per autosomal bivalent and the total number of foci per autosomal complement. The XY bivalent was excluded, as it desynapses before the autosomes. In virtually all cells analyzed, at least one MLH1 focus was present on each autosome. Thus, special mechanisms to segregate achiasmate chromosomes, a feature of meiosis in organisms such as *Drosophila* (24), are unlikely to be an important component of human male meiosis.

7. D. A. Laurie, M. A. Hulten, *Ann. Hum. Genet.* **49**, 189 (1985).
8. J. Yu et al., *Am. J. Hum. Genet.* **59**, 1186 (1996).
9. B. F. J. Manly, *Randomization, Bootstrap and Monte Carlo Methods in Biology* (Chapman and Hall/CRC, New York, 1997).
10. M. Hattori et al., *Nature* **405**, 311 (2000).
11. I. Dunham et al., *Nature* **402**, 489 (1999).
12. A. Lynn, C. Kashuk, A. Chakravarti, unpublished data.
13. N. E. Morton, *Proc. Natl. Acad. Sci. U.S.A.* **88**, 7474 (1991).
14. MicroMeasure 3.3 (25) was used to measure SC lengths of individual chromosomes. Because of cell-cell variation, we compared only SCs measured in the same cell. SCs were considered equal if their lengths (in micrometers) were within 10% of each other; they were considered unequal if the difference exceeded 10%. For SCs of chromosomes 21 and 22, $22 > 21$ in 50 cells and $21 > 22$ in 1 cell. For SCs of chromosomes 16 and 19, $16 > 19$ in 9 cells, $16 = 19$ in 35 cells, and $19 > 16$ in 8 cells.
15. K. E. Koehler, J. P. Cherry, A. Lynn, P. A. Hunt, T. J. Hassold, unpublished data.
16. S. Stack, *J. Cell Sci.* **71**, 159 (1984).
17. M. Sym, G. S. Roeder, *Cell* **79**, 283 (1994).
18. S. L. Page, R. S. Hawley, *Genes Dev.* **15**, 3130 (2001).
19. A. T. C. Carpenter, in *Genetic Recombination*, R. Kucherlapati, G. R. Smith, Eds. (American Society for Microbiology, Washington, DC, 1988), pp. 529–548.
20. S. Agarwal, G. S. Roeder, *Cell* **102**, 245 (2000).
21. S. K. Mahadevaiah et al., *Nature Genet.* **27**, 271 (2001).
22. C. Tease et al., *Am. J. Hum. Genet.*, in press.
23. L. M. Woods et al., *J. Cell Biol.* **145**, 1395 (1999).
24. R. S. Hawley, W. E. Theurkauf, *Trends Genet.* **9**, 310 (1993).
25. A. Reeves, J. Tear, MicroMeasure for Windows, version 3.3. (2000). <http://www.colostate.edu/Depts/Biology/MicroMeasure>.
26. A. Solari, *Chromosoma* **81**, 315 (1980).
27. We thank T. Ashley, H. Willard, and G. Matera for helpful suggestions. This work was funded by National Institutes of Health (NIH) grants HD21341 (to T.J.H.) and HD37502 (to P.A.H.). A.L. is supported by NIH training grant HD07518, and K.E.K. is supported by an American Cancer Society fellowship.

Supporting Online Material

www.sciencemag.org/cgi/content/full/1071220/DC1
 SOM Text

Figs. S1 and S2
 Tables S1 to S3

25 February 2002; accepted 17 May 2002

Published online 6 June 2002;

10.1126/science.1071220

Include this information when citing this paper.

¹Whitehead/MIT Center for Genome Research, Cambridge, MA 02139, USA. ²Department of Biology, Massachusetts Institute of Technology, Cambridge, MA 02142, USA. ³National Human Genome Center, Howard University, Washington, DC 20059, USA. ⁴Department of Pediatrics, College of Medicine, University of Ibadan, Ibadan, Nigeria. ⁵Department of Preventive Medicine and Epidemiology, Loyola University Medical School, Maywood, IL 60143, USA. ⁶Institute of Biological Anthropology, University of Oxford, Oxford, England OX2 6QS. ⁷Departments of Genetics and Medicine, Harvard Medical School; Department of Molecular Biology and Diabetes Unit, Massachusetts General Hospital, Boston, MA 02114, USA.

*To whom correspondence should be addressed. E-mail: altshuler@molbio.mgh.harvard.edu

man genome exists in such blocks, and what are the size and diversity of haplotypes within blocks? Second, to what extent do these characteristics vary across population samples? Third, can haplotype patterns be parsed using only common SNPs sampled from the population, or will the pattern only emerge after complete resequencing (19)? Fourth, how completely does such a haplotype framework capture common sequence variation within each region?

To determine the general structure of human haplotypes, we selected 54 autosomal regions, each with an average size of 250,000 base pairs (bp), spanning 13.4 megabases (Mb) ($\approx 0.4\%$) of the human genome. Regions were selected to fit two criteria: that they be evenly spaced throughout the genome and that they contain an average density [in a core region of 150 kilobases (kb)] of one candidate SNP discovered by The SNP Consortium (TSC) every 2 kb (20) (table S1). Genotyping was performed by primer extension of multiplex products with detection by matrix-assisted laser desorption ioniza-

tion-time of flight (MALDI-TOF) mass spectroscopy (20, 21). Each SNP was genotyped in 275 individuals (400 independent chromosomes) sampled from four population groups: 30 parent-offspring trios (90 individuals) from Nigeria (Yoruba), 93 members of 12 multigenerational pedigrees of European ancestry, 42 unrelated individuals of Japanese and Chinese origin, and 50 unrelated African Americans.

We designed assays to 4532 candidate SNPs of which 3738 (82%) were successfully genotyped (20, 22). Three of the 54 regions were withheld from further analysis due to inconsistencies in genome assembly and/or evidence for a closely related paralogous region [making locus-specific polymerase chain reaction (PCR) difficult]. In the remaining 51 regions, accuracy of genotype calls was empirically assessed as $\approx 99.6\%$ (20). We note that a very low rate of genotyping error is absolutely necessary for studies of multi-marker haplotypes; even a modest error rate creates the appearance of "rare variant"

haplotypes that do not exist in nature.

Of candidate TSC SNPs successfully assayed, 89% were verified to be polymorphic in one or more populations. The proportion polymorphic in each sample varied from 70% (Asian) to 86% (African American) (Fig. 1A). Although the majority of SNPs (59%) were observed in all four populations, there are dramatic differences in the allele frequencies of individual SNPs across samples (fig. S1) consistent with prior estimates of population differentiation and origin (23).

If haplotype blocks represent regions inherited without substantial recombination in the ancestors of the current population, then a biological basis for defining haplotype blocks is to examine patterns of recombination across each region. The history of recombination between a pair of SNPs can be estimated with the use of the normalized measure of allelic association, D' (16, 24). Because D' values are known to fluctuate upward when a small number of samples or rare alleles are examined, we relied on confidence bounds on D' rather than point estimates (20). We define pairs to be in "strong LD" if the one-sided upper 95% confidence bound on D' is >0.98 (that is, consistent with no historical recombination) and the lower bound is above 0.7 (25). Conversely, we term "strong evidence for historical recombination" pairs for which the upper confidence bound on D' is less than 0.9. On average, 87% of all pairs of markers with minor allele frequency >0.2 fell into one of these two categories (and were thus termed "informative" marker pairs). This method should be robust to study-specific differences in the frequencies of SNPs and sample sizes examined, because it relies on those pairs for which narrow confidence intervals (that is, precise estimates) have been obtained.

When this definition is applied to pairs of markers separated by less than 1000 bp, a small fraction of informative pairs show strong evidence of historical recombination (Fig. 1B): 14 to 18% in the Yoruban (African) and African-American samples and 3 to 6% in the European and Asian samples. In the Yoruban and African-American samples, the proportion of pairs displaying evidence for historical recombination rises rapidly with distance, increasing to 50% at a separation of ≈ 8 kb. In the European and Asian samples, by contrast, the fraction of pairs showing strong evidence for recombination rises to 50% at 22 kb. These differences in LD among populations are likely attributable to differences in demographic history (26), because the biological determinants of LD (rates of recombination, mutation, gene conversion) are expected to be constant across groups. The data show that LD extends to a similar and long extent in Asian as well as European samples and that African-American samples show very similar patterns to those observed in the Yoruban population.

The spatial distribution of D' values across

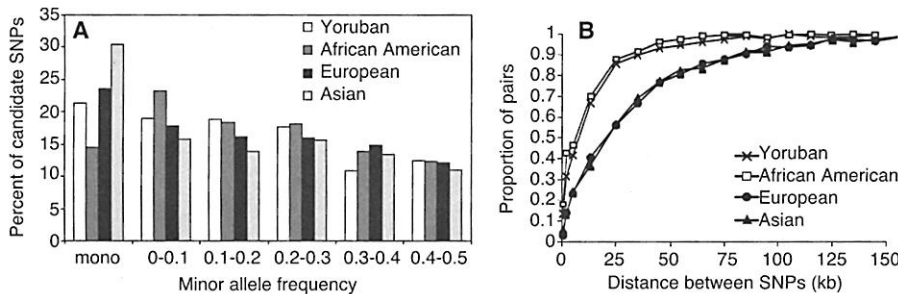


Fig. 1. (A) Normalized allele frequency of candidate SNPs. The distribution is normalized to a constant number of chromosomes ($n = 64$, randomly sampled) from the European, African-American, Asian, and Yoruban samples. Of candidate SNPs assayed in all four populations, both predicted alleles were observed in 89% of cases. **(B)** Assessment of pairwise linkage disequilibrium across populations. The proportion of informative SNP pairs that display strong evidence for recombination is plotted at various intermarker distances. Between 9,860 and 13,980 SNP pairs were examined in each sample.

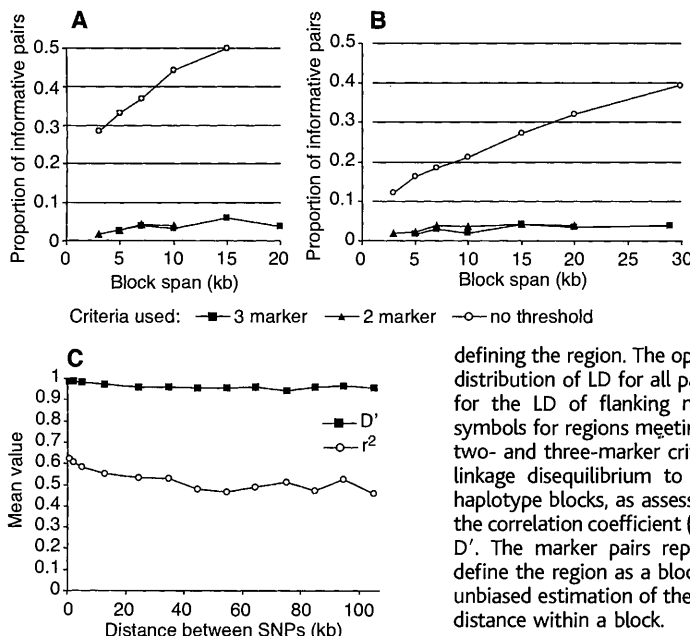


Fig. 2. Scaffold analysis of Yoruban and African-American **(A)** and European and Asian **(B)** samples. The y axis indicates the fraction of independent, informative marker pairs (within each region) displaying strong evidence for recombination. The x axis indicates the distance between the outermost marker pair defining the region. The open symbols represent the distribution of LD for all pairs (without any filtering for the LD of flanking markers), and the closed symbols for regions meeting the empirically derived two- and three-marker criteria (20). **(C)** Relation of linkage disequilibrium to physical distance within haplotype blocks, as assessed by the mean value of the correlation coefficient (r^2) and the mean value of D' . The marker pairs reported were not used to define the region as a block and, thus, represent an unbiased estimation of the relation between LD and distance within a block.

REPORTS

each region (fig. S2) demonstrated clusters of markers over which strong evidence of historical recombination was minimal. We defined a haplotype block as a region over which a very small proportion (<5%) of comparisons among informative SNP pairs show strong evidence of historical recombination. [We allow for 5% because many forces other than recombination (both biological and artifactual) can disrupt haplotype patterns, such as recurrent mutation, gene conversion, or errors of genome assembly or genotyping.] We implemented this definition in two ways (20). Where many markers were sampled, we simply counted the proportion of pairs with strong evidence of historical recombination. However, over much of our survey, we observed regions in which all of the informative markers showed strong evidence of linkage disequilibrium but the number of comparisons was insufficient to confidently conclude (simply by counting) that the proportion of such pairs was >95%. By systematically sampling the entire dataset, we found that information from as few as two or three markers was sufficient to identify regions as blocks (Fig. 2, A and B). These criteria (20) allowed us to define blocks even where the marker coverage is less complete.

Armed with these criteria, we systematically examined the data set for haplotype blocks, identifying a total of 928 blocks in the four populations samples. Within blocks, independent measures of pairwise linkage disequilibrium did not decline substantially with distance (Fig. 2C). The minimum span of the blocks (measured as the interval between the flanking markers used to define them) averaged 9 kb in the Yoruban and African-American samples and 18 kb in the European and Asian samples. However, the size of each block varied dramatically, from <1 to 94 kb in the African-American and Yoruban samples and <1 to 173 kb in the European and Asian samples. Though most of the blocks were small (Fig. 3A), most of the sequence spanned by blocks was in large blocks (Fig. 3B).

Our survey consisted of randomly spaced markers (based on the public map), averaging one marker (with frequency >0.1) every 7.8 kb across the regions surveyed. The partial information leads to two biases in block detection. First, in regions in which we had few markers, we are less likely to detect small blocks. Conversely, identified blocks will typically extend some distance beyond the randomly spaced markers that happen to fall within their boundaries. To estimate the true distribution of block sizes, we performed computer simulations in which block sizes were exponentially distributed (with a specified average size) and markers were randomly spaced (with a mean spacing of one every 7.8 kb). These simulations provided a good fit to the observed data when the mean size of blocks was estimated to be 11 kb in the Yoruban and African-American samples and 22 kb in the European and Asian samples (Table 1)

(27). This corresponds to an N50 size of 22 kb in the Yoruban and African-American samples and 44 kb in the European and Asian populations. (The N50 size is defined as the length x such that 50% of the genome lies in blocks of x or longer.) In addition, the model predicts that the proportion of the human genome spanned by blocks of 10 kb or larger is 65% in the Yoruban and African-American samples and 85% in the European and Asian samples.

We next examined haplotype diversity within blocks. We note that our block definition, unlike one previously proposed (18), is based on recombination and, thus, does not require low haplotype diversity. Nevertheless, within regions with scant evidence for historical re-

combination, we observe only three to five common (>5%) haplotypes in each of the population samples (Fig. 3C). As few as six to eight randomly chosen common markers are sufficient to identify these common haplotypes: the number of haplotypes reached a plateau at six to eight common markers, with little evidence for the discovery of additional common haplotypes if up to 17 markers are included (Fig. 3C). Thus, low haplotype diversity is not simply an artifact of examining only a small number of markers, but it is a true feature of regions with low rates of historical recombination. Haplotype diversity was greatest in the Yoruban and African-American samples, with an average of 5.0 common haplotypes observed. Lower di-

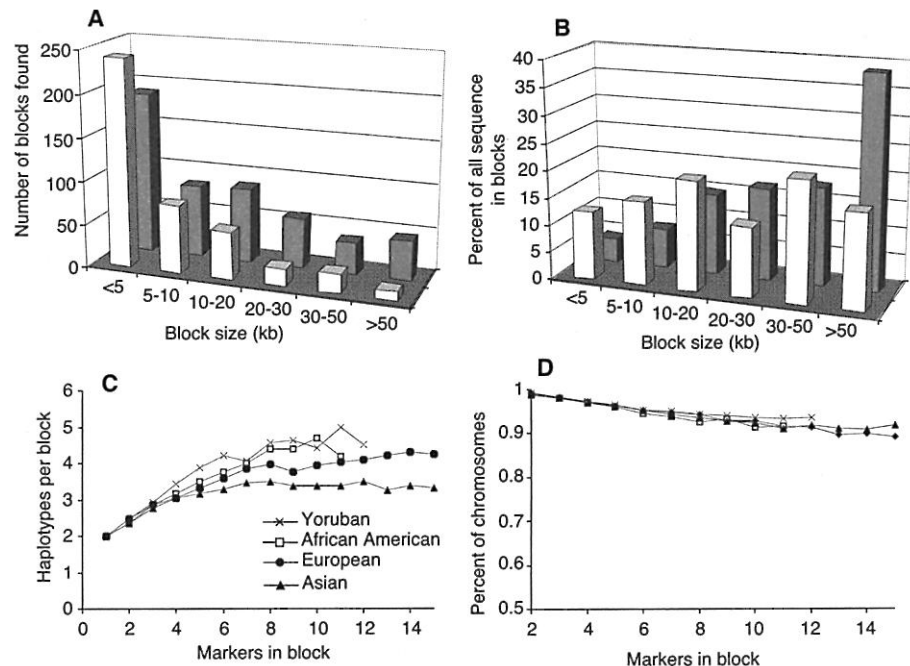


Fig. 3. Block characteristics across populations. (A) Size (in kb) distribution of all haplotype blocks found in the analysis. (B) Proportion of all genome sequence spanned by blocks, binned according to the size of each block. (C and D) Summary of haplotype diversity across all blocks. The number of common ($\geq 5\%$) haplotypes per block (C) and fraction of all chromosomes representing a perfect match to one of these common haplotypes (D) are plotted as a function of the number of markers typed in each block.

Table 1. Observed and predicted proportion of sequence found in haplotype blocks. Model is based on the best fit to the observed data and assumes randomly spaced markers with an average density of one every 7.8 kb. Block span is based on an exponentially distributed random variable with a mean size of 22 kb in the European sample and 11 kb in the Yoruban sample. In the model, block boundaries of 2 kb in length are assumed (17). Although the observed and predicted values were not statistically significantly different (data not shown), both models show a trend toward underestimating the incidence of short blocks (0 to 5 kb). Obs., observed percentage of spanned sequence; pred., predicted percentage of spanned sequence.

Block size	Yoruban sample		European sample	
	Obs.	Pred.	Obs.	Pred.
0 to 5	12.4	6.3	4.4	1.8
5 to 10	15.3	15.1	7.4	5.2
10 to 20	20	31.5	14.9	15.2
20 to 30	12.8	21.8	16.6	16.6
30 to 50	22.2	19.1	18	26.9
>50	17.4	6.3	38.7	34.2

versity was observed in the European samples (4.2 common haplotypes), and the smallest number of common haplotypes (3.5) was observed in the Asian samples. Even where many markers are examined, these few common haplotypes explained the vast majority ($\approx 90\%$) of all chromosomes in each population sample (Fig. 3D) (28, 29).

To compare block boundaries across different populations, we examined adjacent pairs of SNPs successfully assayed in at least two populations. In each population, we asked whether the pair was assigned to a single block or showed strong evidence of historical recombination. A SNP pair was termed concordant if the assignment was the same in both populations and discordant if the assignments disagreed (30). We found the great majority of SNP pairs (77 to 95%, depending on the population comparison) were concordant across population samples (Fig. 4, A to D). Where discordance across populations was observed, it was nearly always due to pairs displaying strong evidence of historical recombination in the Yoruban and African-American samples, but not in the European and Asian samples (Fig. 4, A to D).

We compared the specific haplotypes observed across the European, Asian, and Yoruban samples. To ensure that haplotype diversity was well-defined in this comparison, we considered only those blocks in which six or more polymorphic markers were obtained (31). Each single population sample contained 3.1 to 4.9 haplotypes with a frequency $>5\%$. The union of these sets, however, contained only 5.3 haplotypes (Fig. 4E). That is, the specific haplotypes observed in each group were remarkably similar: 51% (2.7) were identified in all three populations and 72% in two of the three groups.

Of the 28% of haplotypes found in only one population sample, nearly all (90%) were found in the Yoruban sample. The similarity in haplotype identities across the European and Asian samples is striking, with an average of only 0.1 haplotypes per block that were unique to either population sample.

The comparison across populations of SNP polymorphism (Fig. 1A; fig. S1, A to D) recombinant sites (Fig. 4, A to D) and haplotypes (Fig. 4E) is supportive of a single "out of Africa" origin (32, 33) for both the European and Asian samples. The data suggest a considerable bottleneck in the ancestry of these samples, with only a subset of the diversity (of SNPs, haplotypes, and recombinant chromosomes) in Africa found in the two non-African populations (26, 34–37). Because bottlenecks preferentially affect lower-frequency alleles, this model predicts that the alleles (haplotypes and recombinant chromosomes) present only in the African samples would have lower allele frequencies in Africa than pan-ethnic alleles, and our data support this hypothesis (38).

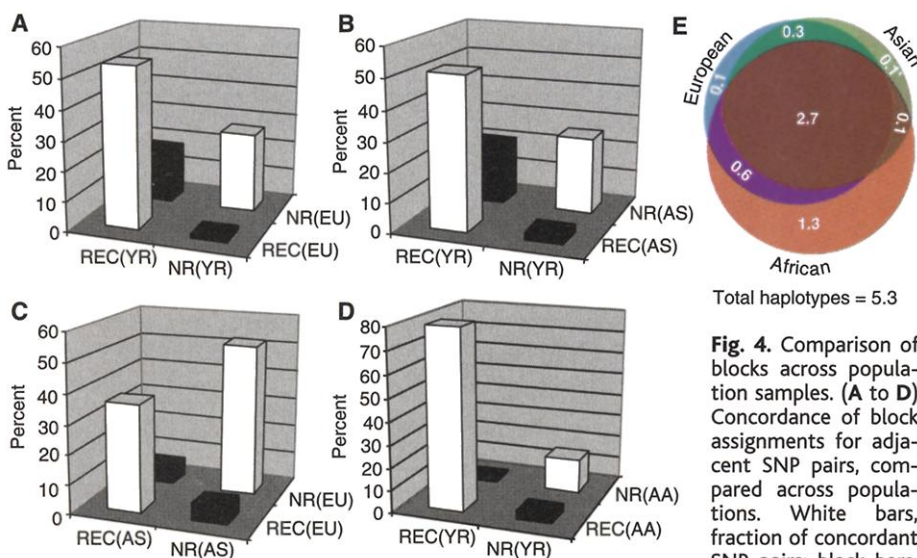
The major attraction of haplotype methods is the idea that common haplotypes capture most of the genetic variation across sizable regions and that these haplotypes (and the undiscovered variants they contain) can be tested with the use of a small number of haplotype tag SNPs ("htSNPs") (16, 18, 19, 39). A number of reports (39–41), however, have suggested that many SNPs fail to conform to the underlying haplotype structure and would be missed by haplotype-based approaches.

To examine this question empirically, we defined a framework of haplotype blocks using a randomly selected subset of our data (requiring a minimum of six markers per block) and

examined the correlation coefficient (r^2) between these haplotypes and an additional set of SNPs (not used to define the blocks) within their span. These additional SNPs were meant to model the undiscovered variation in each region that one would hope to track by using a haplotype approach. We found that the average maximal r^2 value between each additional SNP and the haplotype framework was high, ranging from 0.67 to 0.87 in the four population samples. That is, for the average untested marker, only a small increase in sample size (15 to 50%) would be needed for the use of a haplotype-based (as compared to direct) association study. Moreover, we find that within blocks, a large majority (77 to 93%) of all untested markers showed strong correlation to the haplotype framework (r^2 values greater than 0.5) (42). These results demonstrate that for the vast majority of common alleles it will not be necessary to discover and test each variant individually; a haplotype map could be used with little loss of statistical power.

Our results show that haplotype blocks can be reliably identified by genotyping a sample of common markers within their span, that is, without complete resequencing. However, to ensure that a region is a block, one must type a high density of polymorphic markers in a sufficiently large sample to confidently parse the patterns of historical recombination across the region. Our data provide strong evidence that most of the human genome is contained in blocks of substantial size; we estimate that half of the human genome exists in blocks of 22 kb or larger in African and African-American samples and in blocks of 44 kb or larger in European and Asian samples. Within each block, a very small number of common haplotypes (three to five) typically capture $\approx 90\%$ of all chromosomes in each population. Both the boundaries of blocks and the specific haplotypes observed are shared to a remarkable extent across populations. The main variation is a subset of alleles (haplotypes and recombinant forms) that are observed only in samples with more recent African ancestry. Lastly, blocks defined with a small number of common markers capture quite well the common variation across each locus.

Our results provide a methodological and quantitative foundation for the construction of a haplotype map of the human genome with the use of common SNP markers. Although the patterns are simpler and haplotypes longer than some had predicted, our results suggest that very dense SNP coverage will be needed to complete such a map. With an average block size of 11 to 22 kb and three to five haplotypes per block, our data suggest that fully powered haplotype association studies could ultimately require as many as 300,000 to 1,000,000 well-chosen htSNPs (in non-African and African samples, respectively). However, this number represents an upper limit;



cordant SNP pairs. Population samples are abbreviated as follows: EU, European sample; AS, Asian sample; AA, African-American sample; YR, Yoruban sample. (E) Distribution of haplotypes across populations.

there is often substantial linkage disequilibrium between adjacent blocks (data not shown), allowing fewer markers to be used without loss of power. It will likely be productive to perform initial haplotype mapping in populations whose history contains one or more bottlenecks, because longer-range LD may make initial localization more efficient and favorable. Conversely, populations with shorter-range LD and greater haplotype diversity may offer advantages for fine mapping. In suggesting that block boundaries and common haplotypes are largely shared across populations, our data suggest that many common disease alleles can be studied—and likely will be broadly relevant—across human populations. In the future, comprehensive analysis of human haplotype structure promises insights into the origin of human populations, the forces that shape genetic diversity, and the population basis of disease.

References and Notes

- E. S. Lander, *Science* **274**, 536 (1996).
- F. S. Collins, M. S. Guyer, A. Chakravarti, *Science* **278**, 1580 (1997).
- N. Risch, K. Merikangas, *Science* **273**, 1516 (1996).
- R. Sachidanandam et al., *Nature* **409**, 928 (2001).
- J. C. Venter et al., *Science* **291**, 1304 (2001).
- L. Kruglyak, D. A. Nickerson, *Nature Genet.* **27**, 234 (2001).
- E. G. Puffenberger et al., *Cell* **79**, 1257 (1994).
- B. Kerem et al., *Science* **245**, 1073 (1989).
- J. Hastbacka et al., *Nature Genet.* **2**, 204 (1992).
- J. D. Rioux et al., *Nature Genet.* **29**, 223 (2001).
- J. P. Hugot et al., *Nature* **411**, 599 (2001).
- Y. Ogura et al., *Nature* **411**, 603 (2001).
- J. K. Pritchard, M. Przeworski, *Am. J. Hum. Genet.* **69**, 1 (2001).
- L. B. Jorde, *Genome Res.* **10**, 1435 (2000).
- M. Boehnke, *Nature Genet.* **25**, 246 (2000).
- M. J. Daly, J. D. Rioux, S. F. Schaffner, T. J. Hudson, E. S. Lander, *Nature Genet.* **29**, 229 (2001).
- A. J. Jeffreys, L. Kauppi, R. Neumann, *Nature Genet.* **29**, 217 (2001).
- N. Patil et al., *Science* **294**, 1719 (2001).
- G. C. Johnson et al., *Nature Genet.* **29**, 233 (2001).
- Materials and methods are available as supporting material on Science Online.
- K. Tang et al., *Proc. Natl. Acad. Sci. U.S.A.* **96**, 10016 (1999).
- When 82% of assays were successful in at least one population, genotyping success rates in each population range from 72 to 79%. The difference between these numbers is due to a low rate of laboratory failure in each attempt.
- L. L. Cavalli-Sforza, P. Menozzi, A. Piazza, *The History and Geography of Human Genes* (Princeton University Press, Princeton, NJ, 1994).
- R. C. Lewontin, *Genetics* **49**, 49 (1964).
- An upper confidence bound of 0.98 was used instead of 1.0 because even a single observation of a fourth haplotype makes it mathematically impossible for D' to be consistent with a value of 1.0, though the confidence interval could be arbitrarily close to 1.0.
- D. E. Reich et al., *Nature* **411**, 199 (2001).
- As a further test of the model, we simulated the proportion of pairs at a fixed distance (5 kb) that should show evidence of crossing block boundaries (that is, show strong evidence of historical recombination). The model predicts these proportions to be 47% (Yoruban and African-American samples) and 27% (European and Asian samples). In the empirical data, we observe 42 and 23%, similar to these predictions.
- A low rate of genotyping error is critical to obtaining an accurate measure of haplotype diversity and the proportion in common haplotypes. Even a modest (1 to 2%) genotyping error will create a substantial number of false rare haplotypes; for example, with a 10-marker haplotype and a 2% error rate, 18% of chromosomes will contain at least one error and, thus, not match the few common haplotypes.
- Within blocks, the common haplotypes showed little evidence for historical recombination. For example, we performed the four gamete tests using SNPs drawn only from haplotypes with frequency 5% or higher in each block. One or more violations of the four gamete tests were observed in only 5% of the blocks.
- To maximize power, these comparisons were made only for SNP pairs spaced 5 to 10 kb apart. At shorter distances, nearly all SNP pairs are in a single block; at greater distances, most SNP pairs are in different blocks.
- Blocks and haplotypes were identified separately in each population sample, and the results were compared for those blocks that were physically overlapping in all three samples.
- R. L. Cann, W. M. Brown, A. C. Wilson, *Genetics* **106**, 479 (1984).
- C. B. Stringer, P. Andrews, *Science* **239**, 1263 (1988).
- D. E. Reich, D. B. Goldstein, *Proc. Natl. Acad. Sci. U.S.A.* **95**, 8119 (1998).
- M. Ingman, H. Kaessmann, S. Paabo, U. Gyllenstein, *Nature* **408**, 708 (2000).
- S. A. Tishkoff et al., *Science* **271**, 1380 (1996).
- S. A. Tishkoff et al., *Am. J. Hum. Genet.* **67**, 907 (2000).
- We examined SNP pairs that were in different blocks in the Yoruban samples but in a single block in the European sample. Such pairs had higher D' values in the Yoruban sample ($D' = 0.46$) than pairs found in different blocks in both population samples ($D' = 0.28$). The average frequency of all haplotypes in the Yoruban population was 0.21, whereas those that were found only in the Yoruban sample (but not in the European and Asian samples) had a mean frequency of 0.16.
- A. G. Clark et al., *Am. J. Hum. Genet.* **63**, 595 (1998).
- A. R. Templeton et al., *Am. J. Hum. Genet.* **66**, 69 (2000).
- S. M. Fullerton et al., *Am. J. Hum. Genet.* **67**, 881 (2000).
- The small fraction of SNPs that show r^2 values < 0.5 could be attributable to a range of causes: branches of the gene tree not defined with the number of markers used, gene conversion events, or recurrent mutations. We note that errors in genotyping or map position decrease (but cannot increase) the value of r^2 .
- This work was supported by a grant to D.A. from The SNP Consortium. We thank members of the Program in Medical and Population Genetics at the Whitehead/MIT Center for Genome Research for helpful discussion, particularly J. Hirschhorn, D. Reich, and N. Patterson. The authors also thank the following colleagues for sharing data and analyses before publication: D. Bentley (Sanger Institute); L. Cardon (Oxford); and D. Cutler, M. Zwick, and A. Chakravarti (Johns Hopkins). D.A. is a Charles E. Culpeper Scholar of the Rockefeller Brothers Fund and a Burroughs Wellcome Fund Clinical Scholar in Translational Research.

Supporting Online Material

www.sciencemag.org/cgi/content/full/1069424/DC1
Materials and Methods

Figs. S1 to S3
Table S1

28 December 2001; accepted 13 May 2002

Published online 23 May 2002;

10.1126/science.1069424

Include this information when citing this paper.

Pseudomonas-Candida Interactions: An Ecological Role for Virulence Factors

Deborah A. Hogan and Roberto Kolter*

Bacterial-fungal interactions have great environmental, medical, and economic importance, yet few have been well characterized at the molecular level. Here, we describe a pathogenic interaction between *Pseudomonas aeruginosa* and *Candida albicans*, two opportunistic pathogens. *P. aeruginosa* forms a dense biofilm on *C. albicans* filaments and kills the fungus. In contrast, *P. aeruginosa* neither binds to nor kills yeast-form *C. albicans*. Several *P. aeruginosa* virulence factors that are important in disease are involved in the killing of *C. albicans* filaments. We propose that many virulence factors studied in the context of human infection may also have a role in bacterial-fungal interactions.

Interactions between prokaryotes and eukaryotes are ubiquitous. Although the pathogenic and symbiotic relationships bacteria have with plants and animals have garnered the most attention, the prokaryote-eukaryote encounters that occur among microbes are likely far more common. Many of the virulence factors that we study in the context of human disease may also have an ecological role within microbial communities.

Bacteria and unicellular eukaryotes, such as yeasts and filamentous fungi, are found together in a myriad of environments and exhibit both synergistic and antagonistic interactions (1, 2). Here, we describe a pathogenic relationship between a fungus, *Candida albicans*, and a bacterium, *Pseudomonas aeruginosa*, that involves genes important for bacterial virulence in mammals. *P. aeruginosa* is prevalent in soils and is often found on the skin and mucosa of healthy individuals (3). In compromised hosts, however, *P. aeruginosa* uses an arsenal of virulence factors to cause serious infections associated with burns, catheters, and implants. *C. albicans* is also a benign member of the skin and mucosal flora. When host defenses falter, however, *C.*

Department of Microbiology and Molecular Genetics, Harvard Medical School, 200 Longwood Avenue, Boston, MA 02115, USA.

*To whom correspondence should be addressed. E-mail: rkolter@hms.harvard.edu