

and 15 edges, which Adleman used for his experiments.

Despite these limitations, the compelling aspect of this type of computational system is that an entire result that describes a complete solution is encoded in a single molecule. This coding enables an information representation density that is unheard of in conventional computers and also permits extremely energy-efficient computation; both of these issues are discussed by Adleman (2).

Evolution and NP-complete problems do not seem to have much in common with one another. Or do they? In the case of Adleman's method, he observed that adapter oligonucleotides could constrain a random process toward a solution, even though in a strict sense, constraints are not necessary with adequate screening of solution candidates. Might similar processes be at work in biological systems that evolve? In such a scenario, genetic plasticity would still be created by random events, but constraints might direct mutations to make desired outcomes more probable. If such constraints exist, understanding them would certainly be a major accomplishment.

A second lesson is that computation can take on forms that are not immediately rec-

ognizable to us as computation. We have seen how a process as ordinary as ligation can yield solutions to a hard computational problem. Possible computational applications of other common enzymatic reactions have yet to be fully explored, and thus, it is worthwhile keeping an open mind about the nature of computation in a cell. Transcriptional control and other gene regulation mechanisms certainly play a paramount role in the programming of cell behavior, but there may be other computational mechanisms lurking behind seemingly simple biological processes.

As shown in the figure, NP complete problems are not the most powerful computational systems known. This honor is held by so-called universal systems, which can simulate any computation that can be performed on a deterministic computer (3). If we were able to construct a universal machine out of biological macromolecular components, then we could perform any computation by means of biological techniques. There are certainly powerful practical motivations for this approach, including the information-encoding density offered by macromolecules and the high energy efficiency of enzyme systems.

At present, there is no known way of

creating a synthetic universal system based on macromolecules. Universal systems require the ability to store and retrieve information, and DNA is certainly up to the task if one could design appropriate molecular mechanisms to interpret and update the information in DNA. This ultimate goal remains elusive, but once solved, it will revolutionize the way we think about both computer science and molecular biology.

A great hope is that as we begin to understand how biological systems compute, we will identify a naturally occurring universal computational system. Understanding such a system would give us unprecedented insight into complex biological processes. Perhaps we will ultimately discover that developmental programs and other intricate biological behaviors are built from a common vocabulary of idioms which may be of value to both computer scientists and molecular biologists.

References

1. L. Adleman, *Science* **266**, 1021 (1994).
2. M. R. Garey and D. S. Johnson, *Computers and Intractability: A Guide to the Theory of NP-Completeness* (Freeman, New York, 1979).
3. M. L. Minsky, *Computation: Finite and Infinite Machines* (Prentice-Hall, Englewood Cliffs, NJ, 1967).

Neuroscience on the Net

Peter T. Fox and Jack L. Lancaster

Sure, the Internet provides low-cost entertainment: transcontinental trivia browsing by information junkies; late-night, on-line chat (Fig. 1); electronic junk mail; and even pornography. But a large and growing community of "wired" neuroscientists have found loftier ways to use the Net.

No one will deny that conversation is an important aspect of Net traffic. Discussions with colleagues further away than the next lab are usually by e-mail, instantaneous but buffered. Like traditional mail, you reply in your own time. Person-to-person data transmissions, once cumbersome, are now commonplace. Manuscripts, graphics, and massive data sets hurtle around the world guided by point-and-click interfaces, such as Mosaic, Gopher, and Fetch. This ease of access is complemented by a similar ease of creation. Thousands of laboratories have crafted WWW (World Wide Web) "Home Pages," which provide paths to research program information, preprints, public databases, software, and the like (1).

The authors are at the Research Imaging Center, University of Texas Health Science Center at San Antonio, San Antonio, TX 78784-6240, USA.

Publication, too, is being revolutionized by the Net. Submission by diskette, yesterday's leading edge, is being rendered obsolete by e-mail submission. Still more avant garde are Internet journals. There are now over 70 fully electronic, peer-reviewed, scholarly journals (2). *Psychology*, the most established electronic journal of neuroscience, uses the Net for every aspect of publication: submission, peer review, revision, and distribution. Although revolutionary, electronic publishing is probably not the Internet's most far-reaching restructuring of scientific communication.

Community databases open to all members of a scientific discipline offer the greatest potential for scientific exploitation of the Net. It takes but a moment to understand why. Envision this: On-line access to all relevant results produced by any laboratory in the world, before designing your next experiment. Alternatively, imagine similar access to aid in interpreting an unexpected result. Such is the goal. How do we get there?

The genome community has databased via the Net for roughly a decade. Before

proceeding too far, prospective developers of neuroscience databases should absorb the collective wisdom of these network pioneers. There are dozens of public genetics databases. Most begin as in-house compilations and, when successful, evolve into "collaborative" databases (that is, allowing off-site access by formal agreement). The



Fig. 1. The Internet as entertainment. [Doonesbury © 1993 G. B. Trudeau. Reprinted with permission of Universal Press Syndicate. All rights reserved]

best among them emerge as “community” databases. A community database is an open collaboration. Via Internet, a scientific community at-large both queries and contributes. The accepted structure is client-server. Ideally, the client application program runs on a personal computer to query a large, centralized database served from a high-end workstation or super-computer by a commercial database management system (DBMS). A sophisticated client application will allow a scientist with no knowledge of the DBMS query language to create complex query statements via a point-and-click interface. Through the client, the scientist constructs the unique view of the collective data that best addresses his research question. Remote entry is similarly streamlined. The community database is the most highly evolved form of scientific sharing ... so far. The next step in the evolutionary process—database federation—is envisioned but not yet enacted (3).

Federation is the cooperative creation of multiple, independent databases (for example, each serving a scientific subcommunity) with sufficient commonality of syntax and semantics so that any number of databases can be viewed simultaneously. Although one might naïvely suppose that many syntactic frameworks might be capable of such interoperability, the genome informatics community recommends only one: the Structured Query Language or SQL (3). While there are many SQL database “engines” commercially available, all use the same query-language and, hence, are syntactically compatible. Object-oriented databases offer only product-specific query languages or none at all (3). “Built-from-scratch” databases (that is, not using a commercial DBMS) follow no standard, offering little hope of federation. Thus, the guidelines for syntactic compatibility are established (3). They need only be followed.

Semantic compatibility is by far the greater challenge. At the lowest level, this means that terms must be used in precisely the same way. The genome community learned too late that casual semantics plunge an emerging federation into civil disorder. “Our inability to produce a single definition for ‘gene’ has no adverse effect upon bench research, but it poses real challenges for the development of federated databases” (3). Still more problematic are fundamental differences in the data objects (often differences in physical scale) studied by different scientific subcommunities. Although relationships between base pairs, genes, proteins, and inherited traits are of obvious interest, these data objects are stored in different databases with few semantic bridges. Analogously, in neuroscience molecules, membranes, neurons, cir-

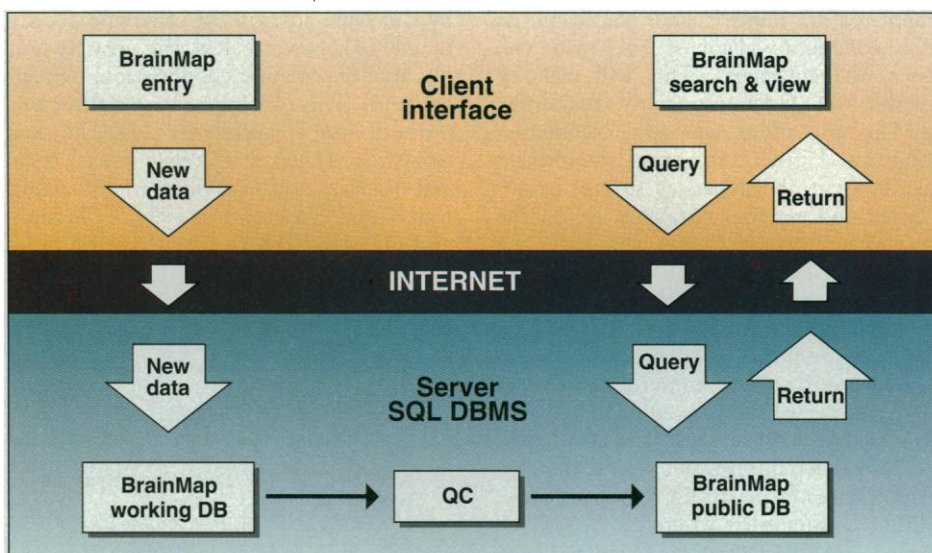


Fig. 2. A classic community database. In BrainMap (7), a centralized SQL (structured query language) database is the server for clients who view and enter data via the Internet through a graphical user interface.

cuits, and systems each has a unique scientific lexicon. If the genome community, despite such a strong tradition of databasing, has made so little progress in federation, where does neuroscience stand?

Neuroscience databasing is in its infancy. A few community databases are operational. CHILDES, a database of speech transcripts, has been sharing linguistic data and analysis tools via the Net for better than a decade (4). GENESIS, an environment for building neural simulations, is distributed via file transfer protocol (FTP) but is not served on-line (5). GENESIS users contribute experimentally derived modeling objects, which are then incorporated in GENESIS and redistributed. BrainMap, a database of human functional neuroanatomy (6), follows the idealized community database model quite closely (Fig. 2). A fully graphical client interface queries a centralized SQL (Oracle) database running on a high-speed UNIX workstation (SUN Sparc 20). Users query a repository of functional-anatomical associations derived from PET (positron emission tomography), fMRI (functional MRI), ERPs (event-related potentials), and ERFs (event-related magnetic fields) to create metanalyses that transcend paper, laboratory, or imaging modality (Fig. 3). Submissions are similarly managed, with a fully graphical application. The latest releases of software tools for analyzing PET and fMRI

brain-mapping experiments not only express results in BrainMap-compatible semantics, they even output files already formatted for remote submission (7, 8). The bottom line, however, is that these three databases (CHILDES, GENESIS, and BrainMap) could not be mutually federated. They operate on fundamentally different data objects and have no common syntax. The Human Brain Project (HBP) (9), a multiagency funding initiative for neuroscience informatics, is scarcely a year old but already funds a family of databases moving toward federation. For example, a database of structural variability (SPMap), designed with the explicit goal of federating with BrainMap, is being created by an International Consortium for Brain Mapping (ICBM), led by John Mazziotta. Information about connectivity patterns among cortical areas in the macaque monkey, already in a local database (10), is being ex-

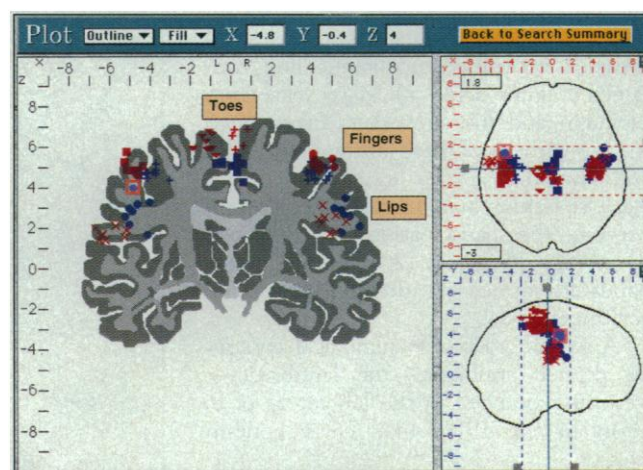


Fig. 3. Transcending paper, laboratory, and imaging modality. BrainMap (7) creates metanalyses of multiple experiments. Two studies (red and blue) of sensory somatotopy are plotted.

trapolated to humans and placed in an SQL community database by David Van Essen and colleagues. This will make it suitable for federation with BrainMap, SPMaP, and other emerging community databases. Further promoting interproject coordination, developers of several neuroscience and genetics databases sit on the BrainMap Advisory Group. Similarly, the developers of BrainMap (P.T.F. and J.L.L.) collaborate on several emerging HBP-funded neuroscience databases.

Thus, despite its youth, the neuroscience informatics community can boast a vision of federation, agreement on a syntax that will support federation, and well-advanced interproject coordination. This is already in strong contrast to the genome community, where no two of the major community databases (Genome Data Base, GenBank, PIR-International, PDB) "are funded by the same program, advised

by the same advisors, or otherwise coordinated" (3). Nevertheless, the greatest hurdle, that of semantic coordination, remains before us. This is being addressed through a series of annual workshops (11). The next workshop, *Database Development in Brain and Behavior*, will take place in San Antonio, 4 and 5 December 1994. Federation and plans to achieve it will be the order of the day.

References and Notes

1. B. R. Schatz and J. B. Hardin, *Science* **265**, 895 (1994).
2. J. Leslie, *Wired* **10**, 68 (1994).
3. R. J. Robbins, *J. Comput. Biol.* **1**, 173 (1994).
4. B. MacWhinney, *The CHILDES Project: Tools for Analyzing Talk* (Lawrence Erlbaum Association, Hillsdale, NJ, 1991).
5. J. M. Bower and D. B. Beeman, *The Book of GENESIS: A Workbook of Tutorials for the General Neural Simulation System* (Springer-Verlag, New York, 1994).
6. P. T. Fox, S. Mikiten, G. Davis, J. L. Lancaster, in

Functional Neuroimaging: Technical Foundations, R. W. Thatcher, M. Hallett, T. Zeffiro, E. R. John, M. Huerta, Eds. (Academic Press, San Diego, CA, 1994), pp. 95–105.

7. K. J. Friston, C. D. Liddle, R. S. J. Frackowiak, *J. Cereb. Blood Flow Metab.* **11**, 690 (1991).
8. P. T. Fox and M. A. Mintun, *J. Nucl. Med.* **30**, 141 (1990).
9. M.F. Huerta, S. H. Koslow, A. I. Leshner, *Trends Neurosci.* **16**, 436 (1993).
10. D. J. Fellman and D. C. Van Essen, *Cerebral Cortex* **1**, 1 (1991).
11. A. Gibbons, *Science* **258**, 1872 (1992).
12. Uniform Resource Locators are as follows: BrainMap - <http://biad38.uthsa.edu/brainmap/brainmap94.html>; ICBM/SPMaP - <http://www.loni.ucla.edu>; Genesis - <http://www.bbb.caltech.edu/GENESIS>; Childes - <ftp://poppy.psy.cmu.edu/www/childes.html>
13. We thank R. Lucier for many conversations on database theory and practice; R. Robbins, P. Gilna, and C. Fields for conversations on database federation; L. Parsons, S. Harnad, and S. Mikiten for insights into the Internet culture; and J. Sergent for unwavering support of BrainMap and the community database concept. This work was supported by a grant from the EJLB Foundation.

Conversion of L- to D-Amino Acids: A Posttranslational Reaction

Günther Kreil

Living organisms synthesize proteins composed of L-amino acids. But on page 1065 of this issue, Heck and colleagues describe an enzymatic activity that converts an L-amino acid to its D form in a peptide from spider venom. Is this a bizarre exception? Probably not. It now seems prudent to consider that D-amino acids can be present in the sequences of secreted peptides of diverse origin.

D-amino acids do occur in bacterial peptides. More than 50 years ago, Lipmann—who was then collaborating with Hotchkiss and Dubos, two pioneers of the early antibiotic era—demonstrated that the small peptides tyrocidine and gramicidine contained D-amino acids (1). Many years later, Lipmann and his co-workers investigated the biosynthesis of these antibiotic peptides and showed that they are assembled in a stepwise fashion by multienzyme complexes without the participation of messenger RNAs and ribosomes (2, 3). The peptide bonds are formed via intermediate aminoacylthioesters.

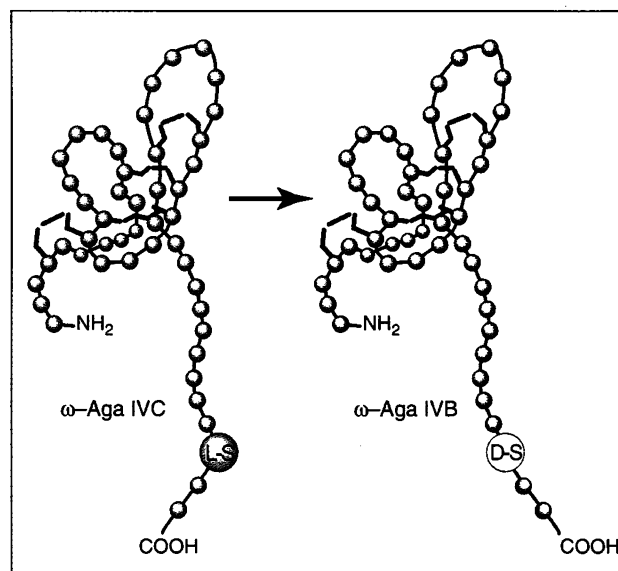
A second group of D-amino acid-containing peptide antibiotics—the lantibiotics—were first analyzed in the laboratory of E. Gross in the 1970s (4). These contain

many unusual amino acids, including lanthionine. At least some of these peptides are derived from larger precursors assembled on ribosomes. A multitude of modifications of the primary translation product then yields the final products. Some of the serine and threonine residues are converted to the corresponding dehydro amino acid and, upon subsequent addition of thiol groups from cysteine residues in

the same sequence, the chirality of the α -carbon changes from the L- to the D-configuration (5).

The report by Heck *et al.* (6) in this issue is the latest addition to a different story, which began in 1981. At that time a group of Italian scientists described the sequence of an opioid peptide isolated from skin of a South American tree frog, *Phyllomedusa sauvagei* (7). This heptapeptide—dermorphin—has the amino-terminal sequence Tyr-D-Ala-Phe. It has a high affinity and selectivity for the μ -type of opiate receptors, and upon injection into the brains of rats and mice acts about a thousand times more effectively than morphine producing long lasting, deep analgesia (7). The D-amino acid is essential for the biological activity of dermorphin. Several additional

peptides containing a D-alanine, D-methionine, or D-leucine as the second amino acid have since been isolated from the skin of *Phyllomedusa* species (8, 9). Some of these resemble dermorphin in their biological activity, while another group of peptides, the deltorphins, are highly selective agonists for δ receptors. More recently, a family of antimicrobial peptides termed bombinins H (10) was isolated from the skin of another frog species, *Bombina variegata*. Some of the members of this family contain D-alloisoleucine instead of isoleucine. Several peptides containing a D-amino acid have also been



Schematic representation of two spider venom toxins. Peptide IVB is generated from IVC by conversion of an L-serine to the D-isomer.

The author is in the Institute of Molecular Biology, Austrian Academy of Sciences, Billrothstrasse 11, A-5020 Salzburg, Austria.