



山東農業大學
SHANDONG AGRICULTURAL UNIVERSITY

生物大数据分析

第一章: 导论

2. 生物信息学大事记

桂松涛

songtaogui@163.com



生物信息学名词的提出



宝琳·霍格维
(Paulien Hogeweg)
(1943-)

荷兰乌得勒支大学教授

- 1978年提出Bioinformatics一词(官方认证的最早)



林华安 博士
(Hwa A. Lim)
(1957-)

马来西亚华裔, 佛州超算中心遗传学与生物物理学部门主任

- 1990年组织了世界第一个国际生物学信息学学术会议, 催生了`生物信息学`一词的出现。



Bioinformatics

- Biology
- Information
- Mathematics



No. 4356 April 25, 1953

NATURE

737

equipment, and to Dr. G. E. R. Deacon and the captain and officers of R.R.S. *Discovery II* for their part in making the observations.

- ¹ Young, F. B., Gerrard, H., and Jevons, W., *Phil. Mag.*, **40**, 149 (1920).
² Longuet-Higgins, M. S., *Mon. Not. Roy. Astro. Soc., Geophys. Supp.*, **8**, 255 (1949).
³ Von Arx, W. S., Woods Hole Papers in Phys. Oceanogr. Meteor., **11** (3) (1950).
⁴ Ekman, V. W., *Arkiv. Mat. Astron. Fysik. (Stockholm)*, **2** (11) (1905).

MOLECULAR STRUCTURE OF NUCLEIC ACIDS

A Structure for Deoxyribose Nucleic Acid

WE wish to suggest a structure for the salt of deoxyribose nucleic acid (D.N.A.). This structure has novel features which are of considerable biological interest.

A structure for nucleic acid has already been proposed by Pauling and Corey¹. They kindly made their manuscript available to us in advance of publication. Their model consists of two chains

is a residue on each chain every 3.4 Å. in the z-direction. We have assumed an angle of 36° between adjacent residues in the same chain, so that the structure repeats after 10 residues on each chain, that is, after 34 Å. The distance of a phosphorus atom from the fibre axis is 10 Å. As the phosphates are on the outside, cations have easy access to them.

The structure is an open one, and its water content is rather high. At lower water contents we would expect the bases to tilt so that the structure could become more compact.

The novel feature of the structure is the manner in which the two chains are held together by the purine and pyrimidine bases. The planes of the bases are perpendicular to the fibre axis. They are joined together in pairs, a single base from one chain being hydrogen-bonded to a single base from the other chain, so that the two lie side by side with identical z-co-ordinates. One of the pair must be a purine and the other a pyrimidine for bonding to occur. The hydrogen bonds are made as follows: purine position 1 to pyrimidine position 1; purine position 6 to pyrimidine position 6.

If it is assumed that the bases only occur in the structure in the most plausible tautomeric forms

738

NATURE

April 25, 1953 VOL. 171

King's College, London. One of us (J.D.W.) has been aided by a fellowship from the National Foundation for Infantile Paralysis.

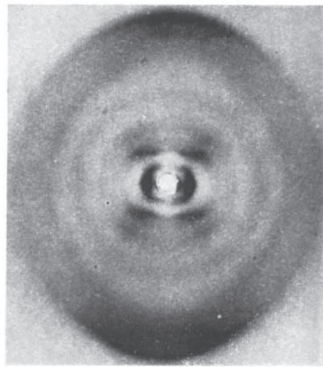
J. D. WATSON
F. H. C. CRICK

Medical Research Council Unit for the Study of the Molecular Structure of Biological Systems, Cavendish Laboratory, Cambridge, April 2.

- ¹ Pauling, L., and Corey, R. B., *Nature*, **171**, 346 (1953); *Proc. U.S. Nat. Acad. Sci.*, **38**, 81 (1953).
² Furlberg, S., *Acta Chem. Scand.*, **6**, 834 (1952).
³ Chargaff, E., for references see Zamenhof, S., Braverman, G., and Chargaff, E., *Biochim. et Biophys. Acta*, **9**, 402 (1952).
⁴ Wyatt, G. R., *J. Gen. Physiol.*, **36**, 201 (1952).
⁵ Astbury, W. T., *Symp. Soc. Exp. Biol.*, **1**, Nucleic Acid, 66 (Camb. Univ. Press, 1947).
⁶ Wilkins, M. H. F., and Randall, J. T., *Biochim. et Biophys. Acta*, **10**, 192 (1953).

Molecular Structure of Deoxypentose Nucleic Acids

WHILE the biological properties of deoxypentose



1953年，由沃森和克里克提出DNA双螺旋结构模型，开启了分子生物学时代

phosphates are on the outside and the bases on the inside, linked together by hydrogen bonds. This structure as described is rather ill-defined, and for this reason we shall not comment on it.



This figure is purely diagrammatic. The two ribbons symbolize the two phosphate-sugar chains, and the horizontal rods the pairs of bases holding the chains together. The vertical line marks the fibre axis

We wish to put forward a radically different structure for the salt of deoxyribose nucleic acid. This structure has two helical chains each coiled round the same axis (see diagram). We have made the usual chemical assumptions, namely, that each chain consists of phosphate diester groups joining β-D-deoxyribofuranose residues with 3',5' linkages. The two chains (but not their bases) are related by a dyad perpendicular to the fibre axis. Both chains follow right-handed helices, but owing to the dyad the sequences of the atoms in the two chains run in opposite directions. Each chain loosely resembles Furlberg's model No. 1; that is, the bases are on the inside of the helix and the phosphates on the outside. The configuration of the sugar and the atoms near it is close to Furlberg's 'standard configuration', the sugar being roughly perpendicular to the attached base. There



Wilkins, Dr. R. E. Franklin and their co-workers at

A straight line may be drawn approximately through

a group of the effects of the shape and size of the repeat unit of the nucleotide on the diffraction pattern. First, if the nucleotide consists of a unit having circular symmetry about an axis parallel to the helix axis, the whole diffraction pattern is modified by the form factor of the nucleotide. Second, if the nucleotide consists of a series of points on a radius at right-angles to the helix axis, the phases of radiation scattered by the helices of different diameter passing through each point are the same. Summation of the corresponding Bessel functions gives reinforcement for the inner-

leic acid ionization by as shown that the inter- ~ 34 Å. eat of a ain con- cation as than the s on or helical

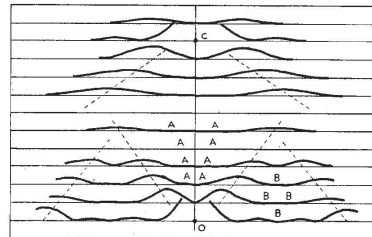


Fig. 2. Diffraction pattern of system of helices corresponding to structure of deoxypentose nucleic acid. The squares of Bessel functions are plotted about 0 on the equator and on the first, second, third and fifth layer lines for half of the nucleotide mass at 20 Å. diameter and remainder distributed along a radius, the mass at a given radius being proportional to the radius. About 10 on the tenth layer line similar functions are plotted for an outer diameter of 12 Å.



The Amide Groups of Insulin

By F. SANGER,* E. O. P. THOMPSON† AND RUTH KITAI
Department of Biochemistry, University of Cambridge

(Received 6 September 1954)

1955年, Sanger用二硝基氟苯 (FDNB) 法, 首次成功地完成了第一个蛋白质-牛胰岛素的序列分析

† Present address: Wool Textile Institute,
343 Royal Parade Parkville, Victoria, Australia

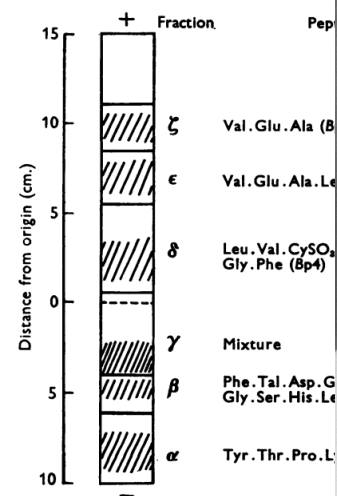
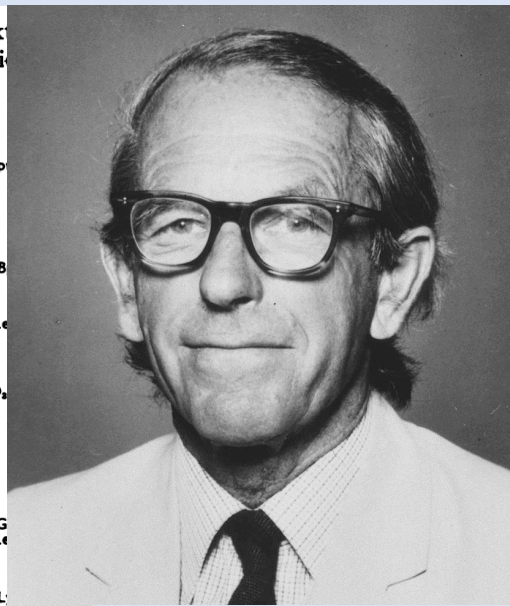
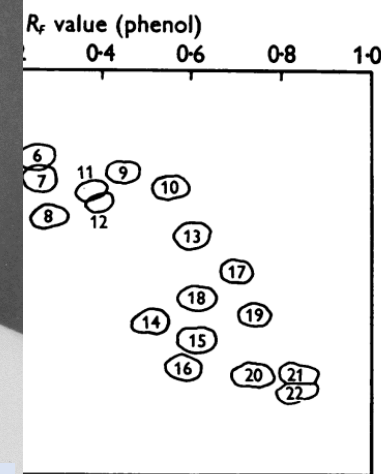


Fig. 1. Ionophoresis of peptic hydrolysate of insulin in 0.05 M ammonium acetate. 2.25 V, 20 min.



Frederick Sanger

hydrolyde Chibnall & Rees (1952)
the three aspartic acids, one of which is



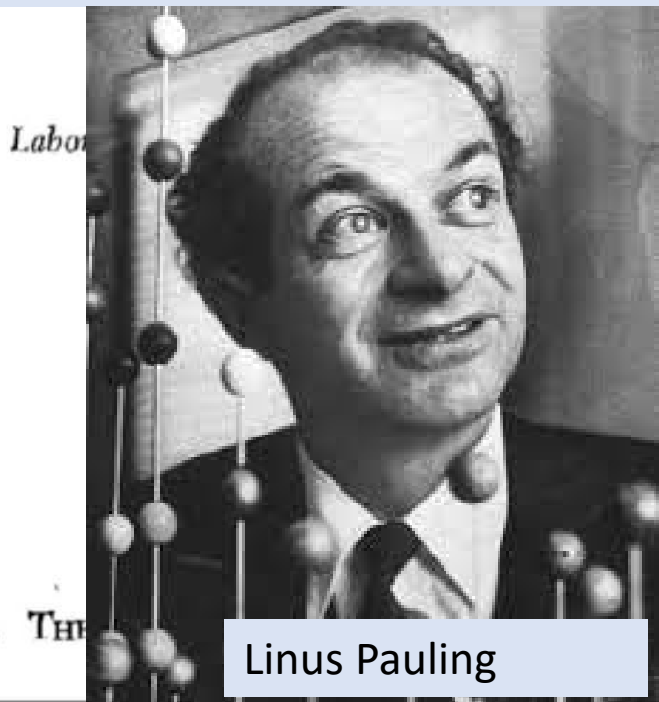
of mould protease hydrolysate of insulin (Expt. 4m) (see Table 1).



Evolving genes and proteins, 1965

Evolutionary Divergence and Convergence in Proteins

1965年，祖卡坎德尔和鲍林提出的“分子钟”理论



Linus Pauling



1966年，我国第一次人工合成了胰岛素





Proc. Natl. Acad. Sci. USA
Vol. 74, No. 12, pp. 5463–5467, December 1977
Biochemistry

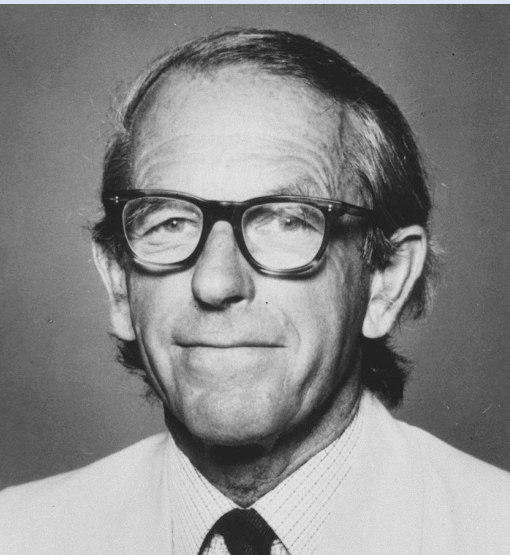
1977年，桑格等发表双脱氧链末端终止法，测定 ϕ X174序列。

Medical Research Council Laboratory of Molecular Biology, Cambridge, England

Contributed by F. Sanger, October 3, 1977

ABSTRACT A new method for determining the sequence of DNA is described. It is a "plus and minus" method [Sanger, F. & Coulson, 1975, *Nature* 251, 441–448] but makes use of the 2',3'-dideoxynucleoside analogues of the normal deoxynucleosides which act as specific chain-terminating agents for DNA polymerase. The technique has been applied to the sequencing of bacteriophage ϕ X174 and is more rapid than either the plus or the minus method.

The "plus and minus" method (1) is a simple technique that has made possible the sequencing of the genome of bacteriophage ϕ X174. It depends on the use of DNA polymerase and the synthesis of complementary regions of the DNA under controlled conditions. This method is considerably more rapid than the available techniques, neither the "plus" nor the "minus" method is completely accurate, and in some cases sequence data both must be used together, and in some cases sequence data are necessary. W. M. Barnes has recently developed a third method, involving the substitution, which has certain advantages over the plus and minus



Frederick Sanger, again!

of ribose in which the 3'-hydroxyl group is oriented in the opposite position with respect to the 2'-hydroxyl group. The 2',3'-dideoxynucleoside (ara) nucleotides act as chain terminating inhibitors for *Escherichia coli* DNA polymerase I in a manner similar to that of ddT (4), although synthesized chains ending in ddT are further extended by some mammalian DNA polymerases (5). In order to obtain a suitable pattern of bands for sequencing an extensive sequence can be read it is necessary to use a ratio of terminating triphosphate to normal triphosphate of only partial incorporation of the terminator. In the case of the dideoxy derivatives this ratio is about 100, and in the case of the arabinoside derivatives about 5000.

METHODS

Preparation of the Triphosphate Analogues. The preparation of the triphosphate analogues has been described (6, 7), and the material is commercially available. ddA has been prepared by the method of Tener (8). We essentially followed their procedure and used the methods of Tener (9) and of Hoard and Ott (10) to convert it to the triphosphate, which was then purified on



1988年，人类基因组计划提出

Perspective

A Turning Point in Cancer Research: Sequencing the Human Genome

RENATO DULBECCO

ONE OF THE GOALS OF CANCER RESEARCH IS TO ASCERTAIN the mechanisms of cancer. Efforts in this direction have been made by using model systems of limited complexity, such as cancer cells in vitro and oncogenic viruses. The use of cell cultures avoided the complexity of the whole animal but not the complexity of the animal genome. The use of oncogenic viruses seemed to circumvent this complexity by replacing it with the extraordinary simplicity of the viral genome. This simplicity made the study of viruses very productive. The persistence of the transformed state in a cell clone could be explained by the persistence of the viral genome in cells (1); genetic and molecular results showed that transformation is the consequence of the expression of one or a few viral genes. Finally, the viral transforming genes, or "oncogenes," and the proteins they specify were identified. The crowning development was the demonstration that in retroviruses the oncogenes are picked up from the cellular genome during the viruses' most recent history (2). As a result of these studies, cancer seemed to be locked to the expression of some viral gene; the possibility of a "hit-and-run" mechanism, in which the virus alters the cell and then vanishes, seemed excluded. Two types of oncogenes were identified: some which immortalize cells, and others which make them tumorigenic (3). In most cases oncogenes of both types are needed to cause a continuously growing tumor.

Subsequent work, however, blurred the distinction between immortalizing and transforming oncogenes by showing that their effects differ in primary cultures or permanent lines and in cells of different species (4). These findings suggested that the state of the cellular genes is important for the effect of oncogenes, in agreement with the great differences in cancer incidence and in the effects of chemical or viral carcinogens in different species.

These studies dealt with the initial cancer events. But natural cancers evolve slowly toward malignancy through many definable stages in a process called "progression" (5), which is the least understood but probably the most crucial phase in the generation of malignancy. Progression generates the marked heterogeneity of cancers (6) and their many chromosomal abnormalities (7); it must be differentiated from the initial action of oncogenes (8). Progression is observed in cells transformed by virus. This is the case, for instance, of bursal lymphomas induced by avian leukosis viruses (9), of viral T-cell lymphomas in mice (10), and of leukemogenesis by Friend leukemia virus in cultures of mouse bone marrow cells (11). Stepwise transformation is observed also with DNA viruses (12). Fibroblastic cells from a variety of organs of a transgenic mouse

containing *myc* and simian virus 40 (SV40) sequences, although expressing SV40 T antigen, were normal but became gradually transformed upon cultivation (13). In all these cases cellular changes occurring during culture growth determined full transformation. The "hit-and-run" hypothesis of viral transformation must be reconsidered.

A clue as to what these changes are is obtained by examining the heterogeneity of chemically induced rat mammary carcinomas with respect to several well-characterized markers. The expression of the markers is altered in different ways in different parts of the same cancer; the alterations seem to be clonal, being uniform in small parts of a tumor but different in adjacent parts (14). The closeness of the parts makes it unlikely that the differences are due to the environment; it is more likely that they are caused by structural changes of the genes, as is also suggested by the chromosomal rearrangements observed in cancers (15) and by the finding that each chemically or radiation-induced mouse sarcoma expresses a different class I major histocompatibility antigen, probably produced through gene rearrangement (16).

A major gap in our understanding of cancer is how the activity of an oncogene is related to the events of progression. But the first task is to ascertain whether the DNA of an advanced cancer is as heterogeneous as the phenotype of its cells. If it is so, a new field of cancer research opens up, possibly leading to the discovery of the genes whose activity or inactivity is responsible for infiltration and metastasis.

We are at a turning point in the study of tumor virology and cancer in general. If we wish to learn more about cancer, we must now concentrate on the cellular genome. We are back to where cancer research started, but the situation is drastically different because we have new knowledge and crucial tools, such as DNA cloning. We have two options: either to try to discover the genes important in malignancy by a piecemeal approach, or to sequence the whole genome of a selected animal species. The former approach seems less formidable, but it will still require a vast investment of research, especially if the important genes differ in cancers of different organs and if they encode regulatory proteins. A major difficulty for conventional approaches is the heterogeneity of tumors and the lack of cultures representative of the various cell types present in a cancer. I think that it will be far more useful to begin by sequencing the cellular genome. The sequence will make it possible to prepare probes for all the genes and to classify them for their expression in various cell types at the level of individual cells by means of cytological hybridization. The classification of the genes will facilitate the identification of those involved in progression.

In which species should this effort be made? If we wish to understand human cancer, it should be made in humans because the genetic control of cancer seems to be different in different species. Research on human cancer would receive a major boost from the detailed knowledge of DNA. Humans would become the preferred experimental species for cancer research with cells in culture or in immunodeficient mice. Because cancer could be defined in molecular terms, the agents capable of inducing cancer in humans could be identified by the combination of in vitro and epidemiological studies. Knowledge of the genes involved in progression would open new therapeutic approaches, which might lead to a general cancer cure if progression has common features in all cancers.

Knowledge of the genome and availability of probes for any gene would also be crucial for progress in human physiology and pathology outside cancer; for instance, for learning about the regulation of individual genes in various cell types. Many fields of

The author is in the Monoclonal Antibody Laboratory of the Armand Hammer Cancer Center, the Salk Institute, La Jolla, CA 92037.

Downloaded from www.sciencemag.org on March 11, 2014

人类基因组计划

开始于1988年

由美国能源部和国家医学研究院发起

美国、英国、法国、德国、日本和中国参加

经费三十亿美元

主要目标：

1、人类基因组测序

2、各种模式生物基因组测序

3、推动全基因组水平的高通量技术的发展



RESEARCH ARTICLE

Whole-Genome Random Sequencing and Assembly of *Haemophilus influenzae* Rd

Robert D. Fleischmann, Mark D. Adams, Owen White, Rebecca A. Clayton,
Ewen F. Kirkness, Anthony R. Kerlavage, Carol J. Bult, Jean-Francois Tomb,
Brian A. Dougherty, Joseph M. Marwick, Keith McKenney, Granger Sutton

natural host is human. Six *H. influenzae* serotype strains (a through f) have been identified on the basis of immunologically distinct capsular polysaccharide antigens. Non-typeable strains also exist and are distinguished by their lack of detectable capsular polysaccharide. They are commensal residents of the upper respiratory mucosa of children and adults and cause otitis media and respiratory tract infections, mostly in

1995年，H. influenza (流感嗜血杆菌)基因组：第一个测序成功的基因组。

Hamilton O. Smith, J. Craig Venter†

An approach for genome analysis based on sequencing and assembly of unselected pieces of DNA from the whole chromosome has been applied to obtain the complete nucleotide sequence (1,830,137 base pairs) of the genome from the bacterium *Hae-mophilus influenzae* Rd. This approach eliminates the need for initial mapping efforts and is therefore applicable to the vast array of microbial species for which genome maps are unavailable. The *H. influenzae* Rd genome sequence (Genome Sequence DataBase accession number L42023) represents the only complete genome sequence from a free-living organism.

ly reduced the incidence of the disease in Europe and North America.

Genome sequencing. The strategy for a shotgun approach to whole genome sequencing is outlined in Table 1. The theory follows from the Lander and Waterman (14) application of the equation for the Poisson distribution. The probability that a base is not sequenced is $P_o = e^{-m}$, where m is the sequence coverage. Thus after 1.83 Mb of sequence has been randomly generated for the *H. influenzae* genome ($m = 1, 1 \times \text{coverage}$), $P_o = e^{-1} = 0.37$ and approximately 37 percent of the genome is unsequenced. Fivefold coverage (approximately 9500 clones sequenced from both insert ends and an average sequence read length

A prerequisite to understanding the complete biology of an organism is the determination of its entire genome sequence. Several viral and organellar genomes have

Homo sapiens (11). These projects, as well as viral genome sequencing, have been based primarily on the sequencing of clones usually derived from extensively mapped restriction

articles

Initial sequencing and analysis of the human genome

International Human Genome Sequencing Consortium*

* A partial list of authors appears on the opposite page. Affiliations are listed at the end of the paper.

The human genome holds an extraordinary trove of information about human development, physiology, medicine and evolution. Here we report the results of an international collaboration to produce and make freely available a draft sequence of the human genome. We also present an initial analysis of the data, describing some of the insights that can be gleaned from the sequence.

The rediscovery of Mendel's laws of heredity in the open 20th century¹⁻³ sparked a scientific quest to unravel nature and content of genetic information that biology for the last hundred years. The scientific pursuit falls naturally into four main phases, corresponding to four quarters of the century. The first established the cellular basis of heredity: the chromosomes. The second defined the molecular basis of heredity: the DNA double helix. The third unlocked the informational basis of heredity, with the discovery of the biological mechanism by which cells read the information contained in genes and with the invention of the recombinant DNA technologies of cloning and sequencing by which scientists can do the same.

The last quarter of a century has been marked by a relentless drive to decipher first genes and then entire genomes, spawning the field of genomics. The fruits of this work already include the genome sequences of 599 viruses and 31 eubacteria, 205 naturally occurring plasmids, 185 organelles, 31 eubacteria, seven archaea, one fungus, two animals and one plant.

Here we report the results of a collaboration involving 20 groups from the United States, the United Kingdom, Japan, France, Germany and China to produce a draft sequence of the human genome. The draft genome sequence was generated from a physical map covering more than 96% of the euchromatic part of the human genome and, together with additional sequence in public databases, it covers about 94% of the human genome. The sequence was produced over a relatively short period, with coverage rising from about 10% to more than 90% over roughly fifteen months. The sequence data have been made available without restriction and updated daily throughout the project. The task ahead is to produce a finished sequence, by closing all gaps and resolving all ambiguities. Already about one billion bases are in final form and the task of bringing the vast majority of the sequence to this standard is now straightforward and should proceed rapidly.

The sequence of the human genome is of interest in several respects. It is the largest genome to be extensively sequenced so far, being 25 times as large as any previously sequenced genome and eight times as large as the sum of all such genomes. It is the first vertebrate genome to be extensively sequenced. And, uniquely, it is the genome of our own species.

Much work remains to be done to produce a complete finished sequence, but the vast trove of information that has become available through this collaborative effort allows a global perspective on the human genome. Although the details will change as the sequence is finished, many points are already clear.

- The genomic landscape shows marked variation in the distribution of a number of features, including genes, transposable elements, GC content, CpG islands and recombination rate. This gives us important clues about function. For example, the developmentally important HOX gene clusters are the most repeat-poor regions of the human genome, probably reflecting the very complex

2001年，人类基因组草图公布。

• The full set of proteins (the proteome) encoded by the human genome is more complex than those of invertebrates. This is due in part to the presence of vertebrate-specific protein domains and motifs (an estimated 7% of the total), but more to the fact that vertebrates appear to have arranged pre-existing components into a richer collection of domain architectures.

- Hundreds of human genes appear likely to have resulted from horizontal transfer from bacteria at some point in the vertebrate lineage. Dozens of genes appear to have been derived from transposable elements.

● Although about half of the human genome derives from transposable elements, there has been a marked decline in the overall activity of such elements in the hominid lineage. DNA transposons appear to have become completely inactive and long-terminal repeat (LTR) retroposons may also have done so.

- The pericentromeric and subtelomeric regions of chromosomes are filled with large recent segmental duplications of sequence from elsewhere in the genome. Segmental duplication is much more frequent in humans than in yeast, fly or worm.

● Analysis of the organization of Alu elements explains the long-standing mystery of their surprising genomic distribution, and suggests that there may be strong selection in favour of preferential retention of Alu elements in GC-rich regions and that these 'selfish' elements may benefit their human hosts.

- The mutation rate is about twice as high in male as in female meiosis, showing that most mutation occurs in males.

● Cytogenetic analysis of the sequenced clones confirms suggestions that large GC-poor regions are strongly correlated with 'dark G-bands' in karyotypes.

- Recombination rates tend to be much higher in distal regions (around 20 megabases (Mb)) of chromosomes and on shorter chromosome arms in general, in a pattern that promotes the occurrence of at least one crossover per chromosome arm in each meiosis.

● More than 1.4 million single nucleotide polymorphisms (SNPs) in the human genome have been identified. This collection should allow the initiation of genome-wide linkage disequilibrium mapping of the genes in the human population.

In this paper, we start by presenting background information on the project and describing the generation, assembly and evaluation of the draft genome sequence. We then focus on an initial analysis of the sequence itself: the broad chromosomal landscape; the repeat elements and the rich palaeontological record of evolutionary and biological processes that they provide; the human genes and proteins and their differences and similarities with those of other

THE HUMAN GENOME

The Sequence of the Human Genome

J. Craig Venter,¹ Mark D. Adams,¹ Eugene W. Myers,¹ Peter W. Li,¹ Richard J. Mural,¹ Granger G. Sutton,¹ Hamilton O. Smith,¹ Mark Yandell,¹ Cheryl A. Evans,¹ Robert A. Holt,¹ Jeannine D. Gocayne,¹ Peter Amanatides,¹ Richard M. Ballew,¹ Daniel H. Huson,¹ Jennifer Russo Wortman,¹ Qing Zhang,¹ Chinnappa D. Kodira,¹ Xiangqun H. Zheng,¹ Lin Chen,¹ Marian Skuski,¹ Gangadharan Subramanian,¹ Paul D. Adams,¹ Jinghui Zhang,¹ George L. Gabor Miklos,² Catherine Nelson,³ Samuel Broder,¹ Andrew G. Clark,⁴ Joe Nadeau,⁵ Victor A. McKusick,⁶ Norton Zinder,⁷ Arnold J. Levine,⁷ Richard J. Roberts,⁸ Mel Simon,⁹ J. Michael Smith,¹⁰ Richard A. Young,¹¹ Richard A. Young,¹² Arthur Delcher,¹ Ian Dew,¹ Daniel Fasulo,¹ Abu Hannehalli,¹ Saul Kravitz,¹ Samuel Levy,¹ Haru-Threidh,¹ Ellen Beasley,¹ Kendra Biddick,¹ Iwar Chandramouliswaran,¹ Rosane Charlab,¹ rancesco,¹ Patrick Dunn,¹ Karen Eliseck,¹ Wangmao Ge,¹ Fangcheng Gong,¹ Zhiping Gu,¹ Ping Guan,¹ Thomas J. Heiman,¹ Maureen E. Higgins,¹ Rui-Ru Ji,¹ Zhaoxi Ke,¹ Karen A. Ketchum,¹ Zhongwu Lai,¹ Yiding Lei,¹ Zhenya Li,¹ Jiyin Li,¹ Yong Liang,¹ Xiaoying Lin,¹ Fu Lu,¹ Gennady V. Merkulov,¹ Natalia Milshina,¹ Helen M. Moore,¹ Ashwinkumar K Naik,¹ Vaibhav A. Narayan,¹ Beena Neelam,¹ Deborah Nusskern,¹ Douglas B. Rusch,¹ Steven Salzberg,¹² Wei Shao,¹ Bixiong Shue,¹ Jingtao Sun,¹ Zhen Yuan Wang,¹ Aihui Wang,¹ Xin Wang,¹ Jian Wang,¹ Ming-Hui Wei,¹ Ron Wildes,¹³ Chunlin Xiao,¹ Chunhua Yan,¹ Alison Yao,¹ Jane Ye,¹ Ming Zhan,¹ Weiqing Zhang,¹ Hongyu Zhang,¹ Qi Zhao,¹ Liansheng Zheng,¹ Fei Zhong,¹ Wenyuan Zhong,¹ Shiaoqing C. Zhu,¹ Shaying Zhao,¹² Dennis Gilbert,¹ Suzanna Baumhueter,¹ Gene Spier,¹ Christine Carter,¹ Anibal Cravchik,¹ Trevor Woodage,¹ Feroze Ali,¹ Huijin An,¹ Aderonke Awe,¹ Danita Baldwin,¹ Holly Baden,¹ Mary Barnstead,¹ Ian Barrow,¹ Karen Beeson,¹ Dana Busam,¹ Amy Carver,¹ Angela Center,¹ Ming Lai Cheng,¹ Liz Curry,¹ Steve Danaher,¹ Lionel Davenport,¹ Raymond Deslites,¹ Susanne Dietz,¹ Kristina Dodson,¹ Lisa Doup,¹ Steven Ferreira,¹ Neha Garg,¹ Anders Gluecksmann,¹ Brit Hart,¹ Jason Haynes,¹ Charles Haynes,¹ Cheryl Heiner,¹ Suzanne Hladun,¹ Damon Hostin,¹ Jarrett Houck,¹ Timothy Howland,¹ Chinyere Ibegwan,¹ Jeffery Johnson,¹ Francis Kalush,¹ Lesley Kline,¹ Shashi Koduru,¹ Amy Love,¹ Felecia Mann,¹ David May,¹ Steven McCawley,¹ Tina McIntosh,¹ Ivy McMullen,¹ Mee Moy,¹ Linda Moy,¹ Brian Murphy,¹ Keith Nelson,¹ Cynthia Pfannkoch,¹ Eric Pratts,¹ Vinita Puri,¹ Hina Qureshi,¹ Matthew Reardon,¹ Robert Rodriguez,¹ Yu-Hui Rogers,¹ Deanna Romblad,¹ Bob Ruhfel,¹ Richard Scott,¹ Cynthia Sitter,¹ Michelle Smallwood,¹ Erin Stewart,¹ Renee Strong,¹ Ellen Shu,¹ Reginald Thomas,¹ Ni Ni Tint,¹ Sukyee Tse,¹ Claire Vech,¹ Gary Wang,¹ Jeremy Wetter,¹ Sherita Williams,¹ Monica Williams,¹ Sandra Windsor,¹ Emily Winn-Deen,¹ Kerielien Wolfe,¹ Jaysree Zaveri,¹ Karena Zaveri,¹ Josep F. Abril,¹⁴ Roderic Guigó,¹⁴ Michael J. Campbell,¹ Kimmen V. Sjolander,¹ Brian Karlak,¹ Anish Kejariwal,¹ Huaiyu Mi,¹ Betty Lazareva,¹ Thomas Hatton,¹ Apurva Narechania,¹ Karen Diemer,¹ Anushya Muruganujan,¹ Nan Guo,¹ Shinji Sato,¹ Vineet Bafna,¹ Sorin Istrail,¹ Ross Lippert,¹ Russell Schwartz,¹ Brian Walenz,¹ Shibu Yoseph,¹ David Allen,¹ Anand Basu,¹ James Baxendale,¹ Dahl Blink,¹ Anne Deslattes Caminha,¹ John Carnes-Stine,¹ Parris Caulk,¹ Yen-Hui Chiang,¹ My Coyne,¹ Carl Dahlke,¹ Marco Delacotte Mayhew,¹ Maria Dombroski,¹ Michael Donnelly,¹ Dale Ely,¹ Shiva Esparham,¹ Carl Foster,¹ Harold Gire,¹ Stephen Glanowski,¹ Kenneth Glasser,¹ Anna Glodek,¹ Mark Gorokhov,¹ Ken Graham,¹ Barry Gropman,¹ Michael Harris,¹ Jeremy Hell,¹ Scott Henderson,¹ Jeffrey Hoover,¹ Donald Jennings,¹ Catherine Jordan,¹ James Jordan,¹ John Kasha,¹ Leonid Kagan,¹ Cheryl Kraft,¹ Alexander Levitsky,¹ Mark Lewis,¹ Xiangjun Liu,¹ John Lopez,¹ Daniel Ma,¹ William Majoros,¹ Joe McDaniel,¹ Sean Murphy,¹ Matthew Newman,¹ Trung Nguyen,¹ Ngoc Nguyen,¹ Marc Nodell,¹ Sue Pan,¹ Jim Peck,¹ Marshall Peterson,¹ William Rowe,¹ Robert Sanders,¹ John Scott,¹ Michael Simpson,¹ Thomas Smith,¹ Arlan Sprague,¹ Timothy Stockwell,¹ Russell Turner,¹ Eli Venter,¹ Mei Wang,¹ Meiyan Wen,¹ David Wu,¹ Mitchell Wu,¹ Ashley Xia,¹ Ali Zandieh,¹ Xiaohong Zhu,¹



nature

Vol 437|15 September 2005|doi:10.1038/nature03959

ARTICLES

Genome sequencing in microfabricated high-density picolitre reactors

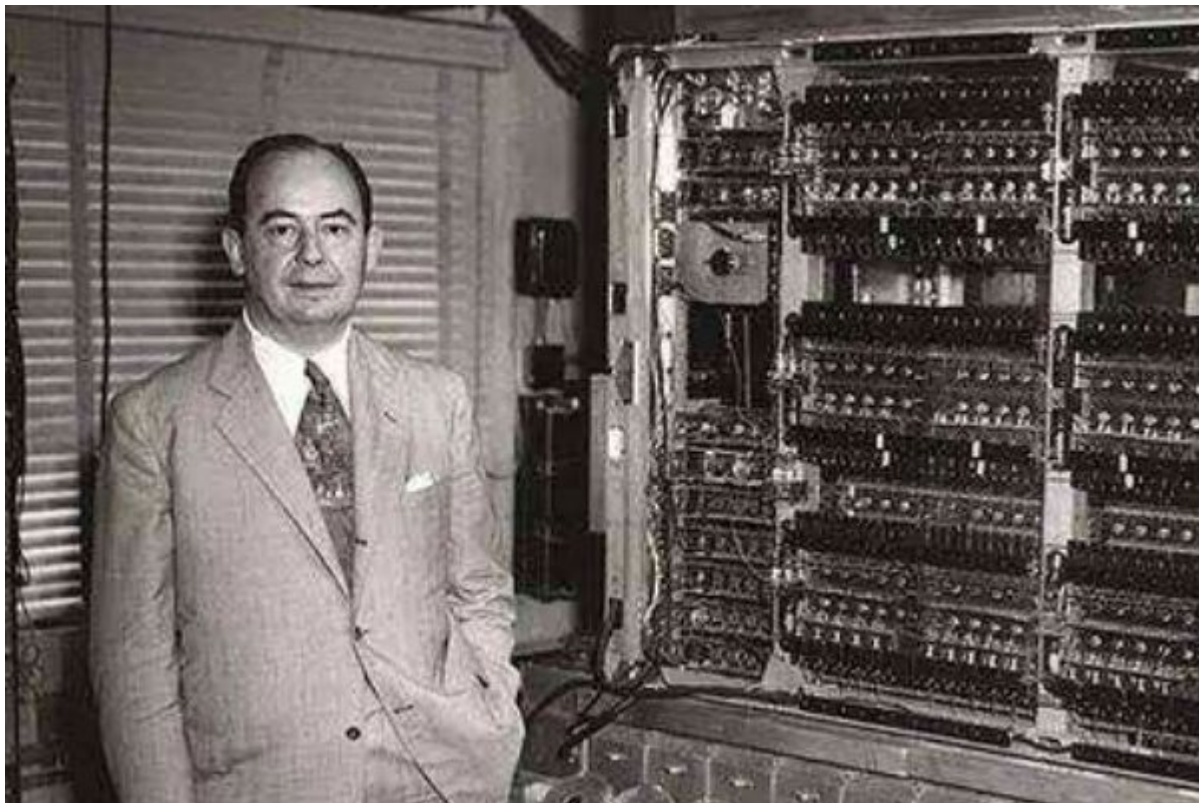
Marcel Margulies^{1*}, Michael Egholm^{1*}, William E. Altman¹, Said Attiya¹, Joel S. Bader¹, Lisa A. Bemben¹, Jan Berka¹, Michael S. Braverman¹, Yi-Ju Chen¹, Zhoutao Chen¹, Scott B. Dewell¹, Lei Du¹, Joseph M. Fierro¹, Xavier V. Gomyk¹, Szilveszter C. J. Kim¹, James R. Knight¹, John R. Lander¹, John H. Leamon¹, Steven M. Lockman¹, Ming Lu¹, Kenton L. Lohman¹, Hong Lu¹, Vinod B. Makhijani¹, Keith E. McDade¹, Michael P. McKenna¹, Eugene W. Myers², Elizabeth Nickerson¹, John R. Nobile¹, Ramona Plant¹, Bernard P. Puc¹, Michael T. Ronan¹, George T. Roth¹, Gary J. Sarkis¹, Jan Fredrik Simons¹, John W. Simpson¹, Maithreyan Srinivasan¹, Karrie R. Tartaro¹, Alexander Tomasz³, Kari A. Vogt¹, Greg A. Volkmer¹, Shally H. Wang¹, Yong Wang¹, Michael P. Weiner⁴, Pengguang Yu¹, Richard F. Begley¹ & Jonathan M. Rothberg¹

2005年，新一代测序技术出现。

The proliferation of large-scale DNA-sequencing projects in recent years has driven a search for alternative methods to reduce time and cost. Here we describe a scalable, highly parallel sequencing system with raw throughput significantly greater than that of state-of-the-art capillary electrophoresis instruments. The apparatus uses a novel fibre-optic slide of individual wells and is able to sequence 25 million bases, at 99% or better accuracy, in one four-hour run. To achieve an approximately 100-fold increase in throughput over current Sanger sequencing technology, we have developed an emulsion method for DNA amplification and an instrument for sequencing by synthesis using a pyrosequencing protocol



约翰·冯诺依曼和计算机

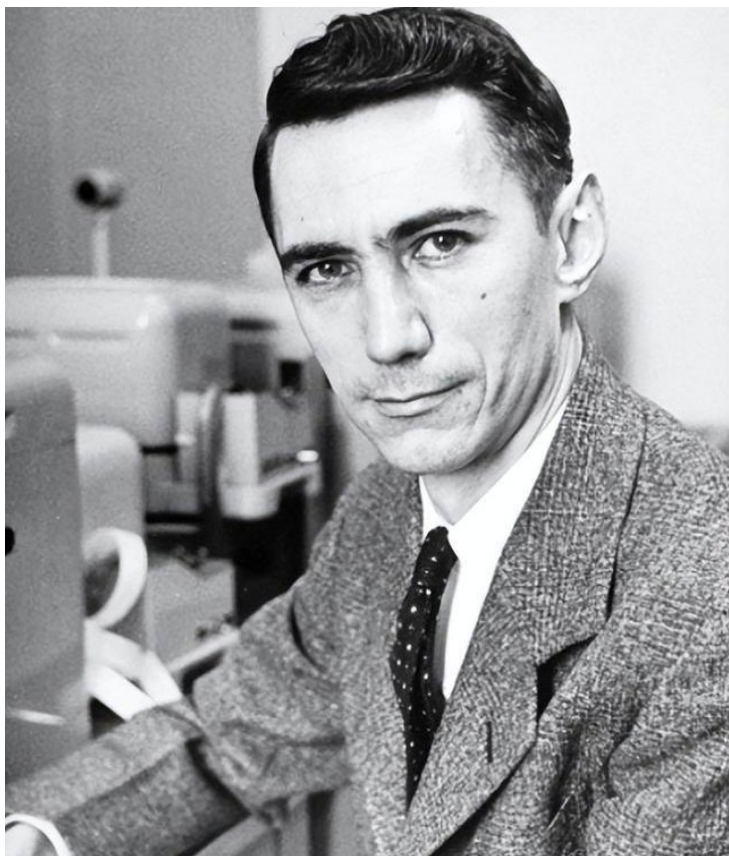


- 算符环理论
- 博弈论
- 蒙特卡洛方法
- 冯诺依曼体系

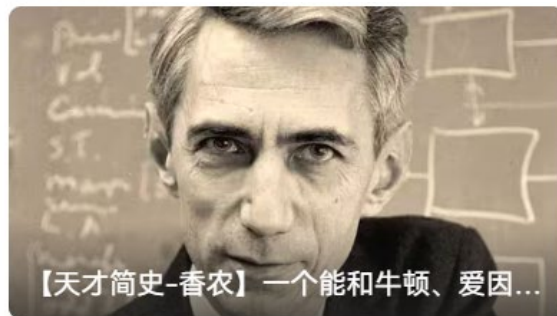
计算机制造的三个基本原则，即采用二进制逻辑、程序存储执行以及计算机由五个部分组成(运算器、控制器、存储器、输入设备、输出设备)



信息论之父: 克劳德·艾尔伍德·香农



$$H(X) = - \sum_{x \in \mathcal{N}} p(x) \log p(x)$$



【天才简史-香农】一个能和牛顿、爱因...



【纪录片】香农传.The.Bit.Player.2018



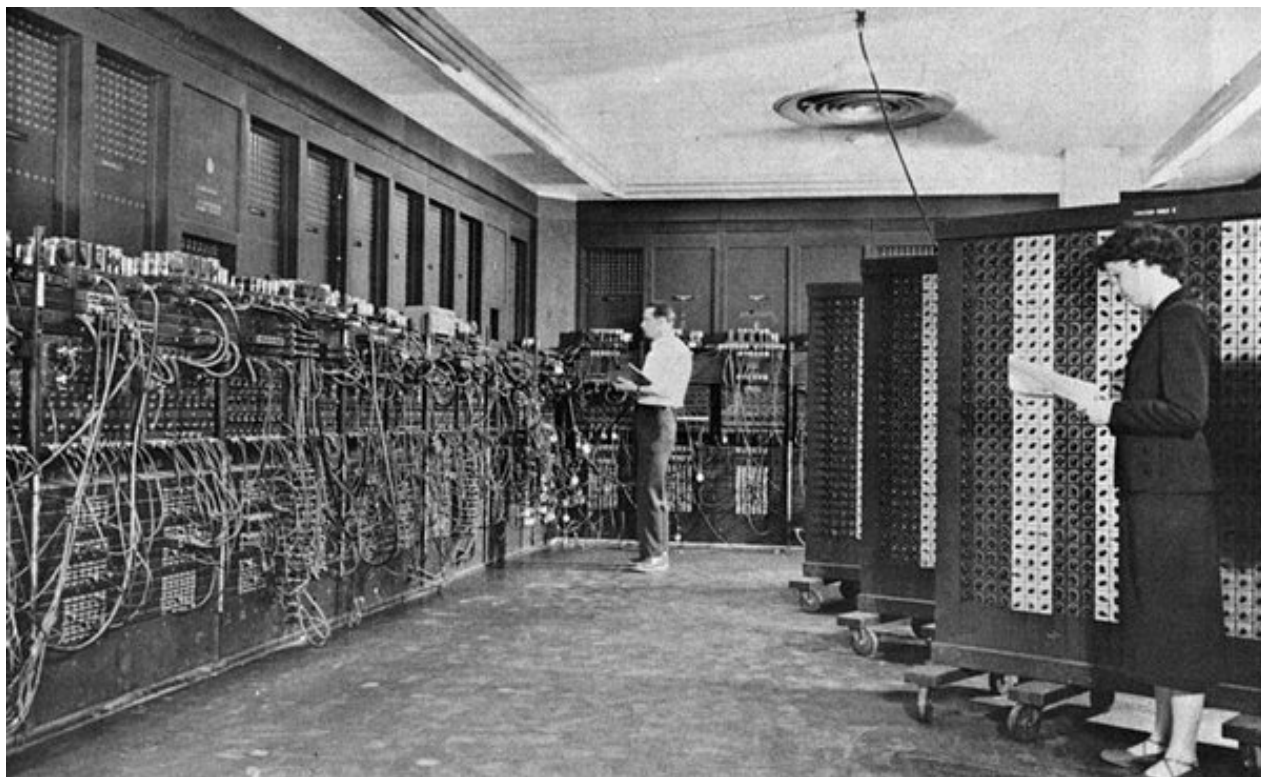


计算机科学与人工智能之父: 艾伦·麦席森·图灵





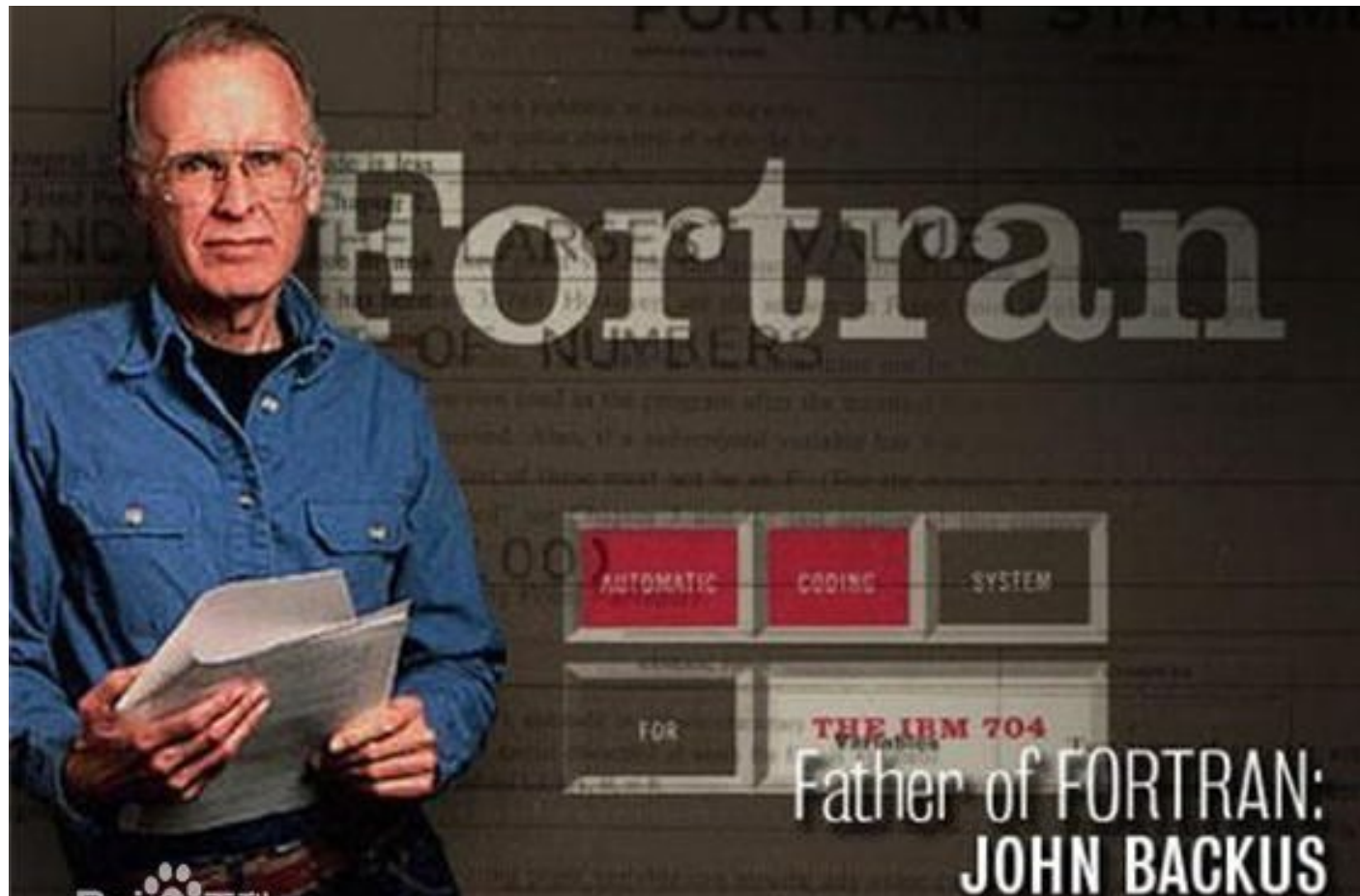
1946年，第一台计算机，ENIAC（埃尼亚克）



ENIAC：长30.48米，宽1米，占地面积约63平方米，30个操作台，约相当于10间普通房间的大小，重达30吨，耗电量150千瓦，造价48万美元。它包含了17,468 真空管7,200水晶 二极管, 1,500 中转, 70,000 电阻器, 10,000 电容器, 1500 继电器, 6000多个开关，每秒执行5000次加法或400次乘法，是继电器计算机的1000倍、手工计算的20万倍。



1951年，IBM公司的约翰·贝克斯在纽约正式对外发布Fortran语言，称为FORTRAN I。





1969-1991: WWW的诞生

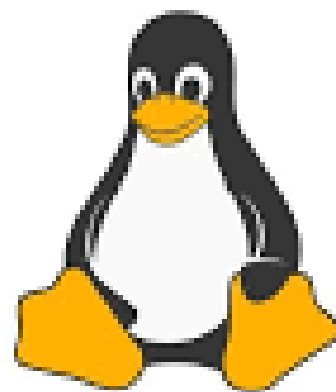
- 1969年，ARPANet建立，对计算机网络技术发展做出重要贡献。
- 1970年后，出现了E-mail、Ethernet、TCP协议。
- 1980年后，以IBM为代表的个人计算机开始普及。
- 1991年，World Wide Web协议被建立。



1991: Linux操作系统



Linus Torvalds



Linux

Shell
Scripting





1952年，肯德鲁用计算机程序来解析蛋白结构。（acta Cryst, 1952）

1962年，**Dayhoff**开发序列分析软件COMPROTEIN。

1965年，**Dayhoff**出版一个蛋白质数据库Atlas（第一年65条序列），发展为1983年的PIR。

1966年，**Dayhoff**对蛋白质家族进化深入研究。（ Science, 1966 ）

1967年，Fitch发表系统发育树（ Science, 1967 ）

1970年，Hesper提出“Bioinformatics”单词，生物信息学概念被定义。

1970年，Needleman和Wunsch提出全局比对算法。（J. Mol. Biol., 1970）

1977年，Protein Data Bank (PDB) 数据库建立。

1978年，**Dayhoff**提出氨基酸序列比对的PAM矩阵。



Dr. Margaret Oakley Dayhoff The Mother of Bioinformatics





1981年，Smith和Waterman发表**局部比对算法**。（J. Mol. Biol., 1981）

1982年，建立核酸序列数据库**Genbank**（最初606条序列）。

1983年，建立Protein information Resource (PIR) 蛋白数据库。2003年（4200万）

1988年，Lipman和Pearson发表FastA算法。（PNAS, 1988）

1990年，Altschul发表**Blast算法**。（J. Mol. Biol., 1990）

1997年，Altschul发表Gapped BLAST和PSI-BLAST算法。（Nucleic Acids Research）

1997年，Chris Burge等发明了GENSCAN算法。（J. Mol. Biol., 1997）

2002年，Kent建立BLAT算法。（Genome Research, 2002）

2003年，Ouzounis对前期的生物信息学发展进行了总结。（BIOINFORMATICS, 2003）



- (1) 萌芽期(60-70年代): 以Dayhoff的替换矩阵和Neelleman-Wunsch算法为代表, 它们实际组成了生物信息学的一个最基本的内容和思路: 序列比较。它们的出现, 代表了生物信息学的诞生(虽然“生物信息学”一词很晚才出现);
- (2) 形成期(80年代): 以分子数据库和FASTA等相似性搜索程序为代表。在这一阶段, 生物信息学作为一个新兴学科已经形成, 并确立了自身学科的特征和地位;
- (3) 基因组测序时代(90年代-至今): 以模式基因组测序与BLAST为代表;
- (4) 高通量测序时代(2005 -至今): 以第二和三代测序技术和基因组重测序为代表。



1994年，中国终于获准加入互联网，并在同年5月完成全部中国联网工作。

1998年，中国人类基因组学南方中心和北方中心分别在上海和北京成立。

1999年，华大基因在北京成立。

2003年，中科院北京基因组所成立。

2008年，魏丽萍老师发表了中国生物信息学发展情况。（PLoS Computational Biology, 2008）



我国生物信息学的开拓者



郝柏林 院士
理论物理所 进化发育分析



陈润生 院士
生物物理所 ncRNA



张春霆 院士
天津大学 Z曲线DNA分析



李衍达 院士
清华大学 基因表达调控



孙之荣 教授
清华大学 分子网络分析



罗辽复 教授
内蒙古大学 基因组进化



- 快速正反馈
- 终身学习
- 头脑风暴
- 学好数理化
- 科学问题驱动



程序员有三种美德: 懒惰, 急躁和傲慢

---- Larry Wall