

ขั้นตอนง่ายๆ ในการสร้างไฟโลจีเนติก

สมชาย แสงอำนาจเดช

Simple Steps in Reconstruction of Phylogenetic Trees

Somchai Saengamnatdej

ภาควิชาจุลชีววิทยาและปรสิตวิทยา คณะวิทยาศาสตร์การแพทย์ มหาวิทยาลัยนเรศวร อ.เมือง จ.พิษณุโลก 65000

Department of Microbiology and Parasitology, Faculty of Medical Science, Naresuan University, Phitsanulok 65000, Thailand.

Corresponding author. E-mail address: somchais@nu.ac.th (S. Saengamnatdej)

Received 24 July 2007; accepted 26 February 2008

บทสรุป

การวิเคราะห์ไฟโลจีเนติก นอกจากจะใช้เพื่อศึกษาวิวัฒนาการและความหลากหลายของสิ่งมีชีวิตแล้ว ยังใช้เป็นเครื่องมือที่มีประโยชน์มากในด้านอื่นๆ เช่น ระบาดวิทยา การควบคุมโรค การออกแบบวัคซีน ยา และเวชภัณฑ์ การค้นหาสปีชีส์ทดแทนที่มีขึ้นทดแทน การค้นหาสารออกฤทธิ์ เป็นต้น บทความปริทัศน์นี้เสนอขั้นตอนง่ายๆ ที่นำไปสู่ความเข้าใจในการสร้าง การอ่าน การทดสอบ และการนำเสนอ โดยใช้ชุดโปรแกรมไฟลิปเป็นแนวทาง พร้อมทั้งแนะนำโปรแกรมชีวสารสนเทศที่เกี่ยวข้อง บทความนี้เหมาะสำหรับผู้ต้องการทำความเข้าใจกับไฟโลจีเนติกเป็นครั้งแรก

คำสำคัญ: ไฟโลจีเนติก; ไฟลิป; การสร้าง; ชีวสารสนเทศ

Summary

Not only is analysis of phylogenetic trees used to study the evolution and diversity of organisms, but also it makes subtle contribution to many other fields or applications, such as epidemiology, control of infectious diseases, design of vaccines, drugs and vectors, search of similar natural products from closely related species, and bioprospecting of bioactive compounds, etc. This review presents simple steps guiding to the understanding of the reconstruction, reading, test, and presentation of phylogenetic trees by using some programs in PHYLIP package as guidance. Other relevant bioinformatic programs are also introduced. This article is suitable for the novice readers who encounter phylogenetic trees for the first time.

Keywords: Phylogenetic tree; PHYLIP; Tree reconstruction; Bioinformatics

บทนำ

การสร้างไฟโลจีเนติก นอกจากจะใช้เพื่อศึกษาวิวัฒนาการและความหลากหลายของสิ่งมีชีวิตแล้ว ยังมีประโยชน์ในด้านอื่น เช่น (1) ประโยชน์ทางระบาดวิทยา การควบคุมโรคติดเชื้อและการออกแบบวัคซีน เช่น การวิเคราะห์ไฟโลจีเนติกของแบคทีเรีย enterotoxigenic *Escherichia coli* (ETEC) ซึ่งเป็นสาเหตุของโรคท้องร่วงสายพันธุ์ต่างๆ ที่มาจากหลายแหล่งภูมิศาสตร์และชนิดของโฮสต์ ช่วยให้เข้าใจการได้รับยีนของสารพิษที่ก่อความรุนแรงในการเกิดอาการท้องเสียจากแบคทีเรียชนิดนี้ (Turner et al., 2006) หรือการวิเคราะห์ไฟโลจีเนติกของไวรัสโปลิโอ (poliovirus) และไวรัสค็อกซซากี (coxsackie virus) ที่จัดรวมอยู่ในเอ็นเทอร์โอไวรัส คลัสเตอร์ C ของคน (HEV-C) ช่วยให้สามารถทำนายความเป็นไปได้ที่ไวรัสโปลิโออาจจะกำเนิดมาจากการกลายพันธุ์ของไวรัสใกล้เคียงที่ก่อโรครุนแรง ทำให้ได้ข้อมูลที่เป็นประโยชน์ในโครงการรณรงค์ทำลายล้างไวรัสโปลิโอทั่วโลก (Jiang et al., 2007) (2) ใช้เป็นแผนที่นำทางเพื่อช่วยทำนายคุณสมบัติของสปีชีส์ที่ยังไม่มีการศึกษา เพื่อค้นหาสารที่มีฤทธิ์ทางเภสัชวิทยาเพื่อใช้เป็นยา หรือเครื่องมือวิจัยทางสรีรวิทยา เช่น การใช้ไฟโลจีเนติกช่วยทำนายพบว่า ปลามากกว่า 1,200 ชนิด ที่อาจมีพิษและจะนำไปสู่การ

ศึกษาเพื่อใช้ประโยชน์หรือการพบสปีชีส์ใกล้เคียง อาจใช้เป็นแหล่งของสารตั้งต้นทดแทนสารที่ต้องการ เช่น การใช้ใบของ *Taxus baccata* เป็น แหล่งของสารแทนการใช้เปลือกของ *T. brevifolia* ในการผลิตยาต้านมะเร็ง paclitaxel (Smith & Wheeler, 2006) (3) นำไปสู่ความรู้ความหลากหลายในโครงสร้างสำหรับเป็นแหล่งศึกษาโครงสร้างและหน้าที่ และการออกแบบยารักษา เช่น การวิเคราะห์ไฟโลจีเนติกของยีนของแฟคเตอร์ eIF4E (Joshi et al., 2005) หรือ การวิเคราะห์ที่พบว่า ภูและกิ่งก้านมีระบบพิษที่วิวัฒนาการจากจุดเริ่มต้นเดียวกันที่จะนำไปสู่การวิจัยทางชีวการแพทย์และการออกแบบยา (Fry et al., 2006) (4) เป็นพื้นฐานสำหรับการออกแบบเวกเตอร์อะดีโนไวรัสให้มีความจำเพาะต่อเนื้อเยื่อ (Madisch et al., 2007)

บทความนี้ขยายความและเสริมส่วนที่ผู้เขียนเห็นว่าสำคัญต่อความเข้าใจที่ขาดไปจากบทความที่ตีพิมพ์เผยแพร่ก่อนหน้านี้ (Baldauf, 2003; Harrison & Langdale, 2006) จะช่วยให้ผู้ที่ยังไม่เคยสร้างหรือเข้าใจการอ่านไฟโลจีเนติก และ ผู้ที่อยู่ในสาขาอื่นที่ไม่ได้ศึกษาด้านวิวัฒนาการและความหลากหลายของสิ่งมีชีวิต แต่สนใจจะนำไปใช้เป็นเครื่องมือในสาขาของตน สามารถเข้าใจการสร้างการอ่าน ระบุโปรแกรมที่สำคัญ หลักการ และวิธีการทดสอบความถูกต้องของไฟโลจีเนติกและพื้นฐานที่

เกี่ยวข้องได้ในระดับหนึ่ง

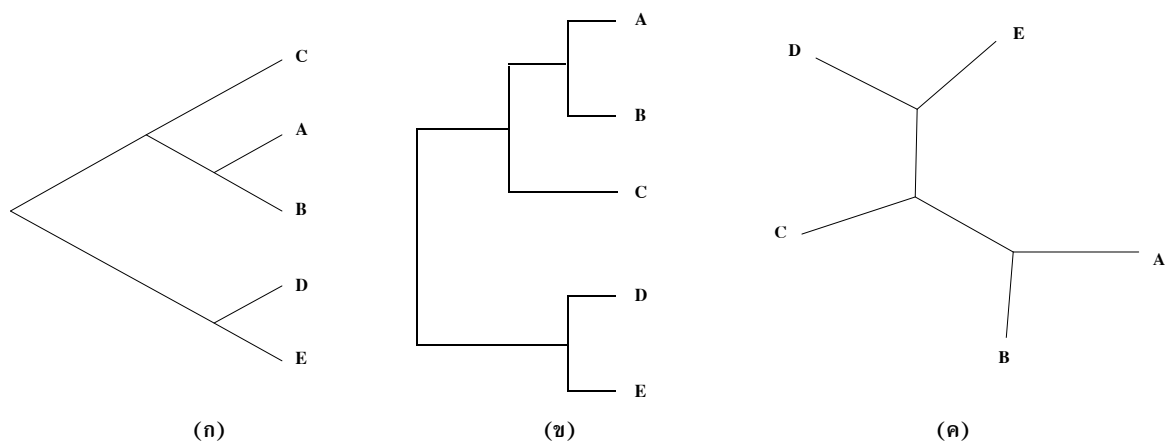
การสร้างไฟโลจีนติก

ไฟโลจีนติกสามารถสร้างโดยอาศัยข้อมูลลักษณะสัณฐานวิทยา ชากติกด้าบรรพ์หรือข้อมูลระดับโมเลกุล โดยเฉพาะลำดับนิวคลีโอไทด์และลำดับโปรตีนของยีนต่างๆ ที่เป็นออร์โธล็อก (orthologue) ในบทความนี้จะใช้ข้อมูลระดับโมเลกุล ก่อนที่จะกล่าวถึงขั้นตอนต่างๆ ควรจะต้องทำความเข้าใจองค์ประกอบที่สำคัญของไฟโลจีนติก และคำศัพท์ที่เกี่ยวข้อง

ไฟโลจีนติก มีองค์ประกอบคล้ายต้นไม้ คือ ประกอบด้วย กิ่งก้านหรือแขนง (branch) ก้านอาจแตกเป็นกิ่งย่อยแบบสองทาง (bifurcation) หรือหลายทาง (multifurcation) ซึ่งแบบหลังนี้พบน้อยมาก ตำแหน่งที่ก้านแตกเป็นกิ่งย่อยเป็นจุดต่อ เรียกว่า โหนด (node) ที่ปลายสุดของกิ่งจะเป็นลำดับเบส หรือโปรตีนสายพันธุ์ หรือสปีชีส์ของสิ่งมีชีวิต เรียกว่า ใบ (leaf) หรือแทกซอน (taxon) หรือหน่วยอนุกรมวิธาน เชิงปฏิบัติการ (OTU; operational taxonomic unit) อาจพิจารณาที่เป็นกราฟชนิดหนึ่งซึ่งมีส่วนยอด (vertex) เป็นโหนด และแขนงที่ต่อโหนดเป็นเส้นขอบ (edge) แต่ทว่าอาจมีจุดกำเนิด ร่วมหือราก (rooted tree) ซึ่งเป็นส่วนยอดที่เป็นตำแหน่งของบรรพบุรุษร่วมของแทกซอน (taxa) ทั้งหมดหรือไม่มีราก (unrooted tree) ก็ได้ ทรีที่มีรากจะทำให้ทราบทิศทาง เวลา และบรรพบุรุษ ในการวิวัฒนาการความยาวของก้านที่แตกต่างกัน แสดงระยะห่างของวิวัฒนาการ ซึ่งแต่ละก้านอาจมีอัตราเร็วของวิวัฒนาการแตกต่างกัน ระยะห่างอาจแสดงเป็นตัวเลขบนแต่ละก้านหรืออาจเขียนเป็นสเกลความยาวบอกขนาดในภาพหรืออาจเขียนในลักษณะเส้นตรงที่ยกออกจากกัน (รูปที่ 1ก.) หรืออาจเขียนด้วยเส้นแนวตั้งและแนวนอน (รูปที่ 1ข.) เพื่อให้เปรียบเทียบความยาวของก้านได้ง่าย ในกรณีเส้นในแนวตั้งเป็นเส้นเชื่อมต่อไม่มีความหมายอื่นใด หรืออาจเขียนในลักษณะที่ส่วนของใบเสมอกันและโดยความยาวของก้านไม่มีความหมาย ทรีในลักษณะนี้ เรียกว่า เคลดโดแกรม (cladogram) ถ้าพิจารณาทรีโดยสนใจเพียงโครงสร้างการแตกกิ่งก้าน ไม่สนใจความยาวของขอบ สามารถจะกล่าวว่า ทรีนั้นมีโทโพโลยี (topology) เดียวกันหรือเรียก

topological trees เมื่อสามารถบิตหรือยึดขอบต่างๆ แล้วทำให้ทรีที่มีลักษณะโครงสร้างของกิ่งก้านแบบเดียวกัน หรืออาจเรียกว่าเป็น additive tree ถ้าระยะห่างของวิวัฒนาการของสองแทกซอนที่แยกมาจากโหนดเดียวกันเท่ากับผลบวกของระยะห่างทั้งสองจากโหนดนั้น ถ้าตั้งสมมุติฐานให้ทรีมีอัตราการกลายพันธุ์ (mutation rate) คงที่สำหรับทุกสายทาง (lineages) สำหรับการวิวัฒนาการของลำดับเบสหรือโปรตีน จะเรียกสมมุติฐานนี้ว่า สมมุติฐานนาฬิกาของโมเลกุล (molecular clock assumption) เนื่องจากอัตราการกลายพันธุ์คงที่ ดังนั้น ปริมาณการกลายพันธุ์ของแต่ละเส้นขอบจะเป็นสัดส่วนกับเวลาที่ผ่านไป ดังนั้น ไม่ว่าจะใช้ความยาวของเส้นแทนปริมาณการกลายพันธุ์หรือเวลาที่ผ่านไปจะได้ตัวเลขเดียวกัน แต่ถ้าไม่ใช้สมมุติฐานนาฬิกาของโมเลกุลที่ขอบที่มีระยะเวลาเท่ากัน แต่มีอัตราเร็วของการกลายพันธุ์แตกต่างกัน ขอบที่มีอัตราเร็วการกลายพันธุ์มากกว่าจะมีปริมาณการกลายพันธุ์มากกว่าในช่วงเวลาที่ผ่านไปเท่ากันนั้น (Allman & Rhodes, 2004) กลุ่มของแทกซอนที่กำเนิดมาจากบรรพบุรุษร่วมกันเรียกว่า เคลด (clad) หรืออาจมีเพียงเคลดเดียว เรียกว่า monophyletic tree หรือ มีหลายเคลดเรียกว่า polyphyletic tree และแต่ละคู่ของสปีชีส์หรือแทกซอนที่แยกจากกันและกันหนึ่งช่วงในทรีที่อยู่ภายในทรี (internal node) หรือ กล่าวอีกอย่างก็คือ แทกซอนคู่ใด ๆ ที่เส้นขอบวิ่งไปบรรจบกัน เรียกว่า เป็นเพื่อนบ้านกัน (neighbors) สำหรับสปีชีส์หรือแทกซอนที่แยกออกจากสปีชีส์อื่นที่กำลังศึกษาในระยะแรกๆ ของการวิวัฒนาการ เรียกว่า พวกนอกกลุ่ม (outgroup) ซึ่งสปีชีส์เหล่านี้สามารถนำมาใช้เพื่อสร้างทรีแบบมีรากได้ (Deodier et al., 2005; Krane & Raymer, 2003) องค์ประกอบที่สำคัญของไฟโลจีนติก ดังแสดงในรูปที่ 1

จำนวนของทรีทั้งชนิดมีรากและไม่มีราก สามารถคำนวณหาได้ เมื่อทราบจำนวนของแทกซอน ซึ่งจะมีจำนวนมากมาย ดังนั้น โปรแกรมสำหรับสร้างไฟโลจีนติกจึงมุ่งไปสู่การสร้างหรือค้นหาทรีที่เหมาะสมหรือดีที่สุด ซึ่งบางวิธีจะใช้เวลาและกำลังคอมพิวเตอร์ในการวิเคราะห์มาก การสร้างไฟโลจีนติกที่มี 5 ขั้นตอน คือ การรวบรวมชุดข้อมูล การจัดเรียงลำดับ การพิจารณาเลือกวิธีวิเคราะห์ เลือกโมเดลวิวัฒนาการ และโปรแกรมการทดสอบ และการนำเสนอ



รูปที่ 1 องค์ประกอบที่สำคัญของไฟโลจีนติก ทรีในรูป ก และ ข เป็นแบบชนิดมีราก โดยรูป ข ทรีอยู่ในลักษณะเส้นนอน และเส้นตั้ง ส่วนรูป ค เป็นแบบไม่มีราก ทั้งสามทรีนี้มีโทโพโลยี เหมือนกัน A-E คือ แทกซอน โดย A, B แยกมาจากโหนดหนึ่ง และ D, E จากอีกโหนดหนึ่ง และแทกซอน A, B, C อยู่ในเคลด หนึ่งส่วนแทกซอน D, E อยู่อีกเคลดหนึ่งต่างเคลดกัน

ขั้นที่ 1 การรวบรวมชุดข้อมูล

การค้นหาและเรียกลำดับจากฐานข้อมูลนิวคลีโอไทด์ทำได้ 2 วิธี เพื่อเพิ่มโอกาสในการได้ลำดับที่คล้ายกับยีนที่สนใจ ควรใช้ทั้งสองวิธี คือ การเรียกข้อมูลโดยใช้คำสำคัญ อาจใช้โปรแกรมเรียกข้อมูล เช่น SRS (<http://srs.ebi.ac.uk>) หรือ Entrez (<http://www.ncbi.nlm.nih.gov>) จากฐานข้อมูลทั่วไป คือ วว BJ (<http://www.ddxvni.ac.th>) EMBL (<http://www.ebi.ac.uk/EMBL/>) FenBank (<http://www.ncbi.nlm.nih.gov>) หรือจากฐานข้อมูลยีนเฉพาะ เช่น TIF (<http://www.tigr.org/tdb/mdb/mdbcomplete.html>) JFI (<http://www.gi.doe.gov/JFI/microbial.html#index.html>) Sanger (<http://www.sanger.ac.uk/rosetts/Microbes/>) และ 5 CBI (<http://www.ncbi.nlm.nih.gov/Fenomes/index.html>) และการเรียกข้อมูลโดยการค้นความคล้ายกัน โดยใช้โปรแกรมบลาส (BLAST) ซึ่งมีในทุกระบบข้อมูลและเว็บไซต์ยีนส่วนใหญ่ เช่น 5 CBI BLAST (<http://www.ncbi.nlm.nih.gov/BLAST/>)

ขั้นที่ 2 การนำลำดับทั้งหมดมาจัดเรียง

เนื่องจากการจัดเรียงลำดับโดยใช้โปรแกรม เช่น Clustal หรือ ClustalW (<ftp://ftp.ebi.ac.uk/pub/software/cluster/>) ผลที่ได้อาจเกิดช่องว่าง (gaps) ในตำแหน่งที่ไม่เหมาะสม ดังนั้นสมควรที่จะนำมาปรับแต่งโดยการตรวจสอบด้วยตนเองอีกครั้ง โดยใช้โปรแกรม BioEdit (<http://www.mbio.ncsu.edu/BioEdit/bioedit.html>)

ขั้นที่ 3 การเลือกใช้วิธีวิเคราะห์โมเดลและโปรแกรมสำหรับสร้างทรี

3.1 พิจารณาเลือกที่จะสร้างทรีจากลำดับของดีเอ็นเอหรือลำดับของโปรตีน

ชนิดของกรดอะมิโนมีจำนวนหลากหลายกว่าชนิดของนิวคลีโอไทด์ (20:4) แต่ถ้าลำดับโปรตีนมีความใกล้เคียงกันมากจะมีการเปลี่ยนแปลงมากกว่าที่ระดับของดีเอ็นเอ ดังนั้น ควรใช้ลำดับของดีเอ็นเอ แต่ถ้าลำดับมีความสัมพันธ์ที่ห่างกันมาก ลำดับของกรดอะมิโนจะเก็บข้อมูลมากกว่า อย่างไรก็ตามก็ยังสามารถใช้ลำดับดีเอ็นเอสำหรับลำดับที่มีความสัมพันธ์ห่างกันได้การตัดเอาโคดอนตำแหน่งที่สามออก เพราะจะรบกวนหรือทำให้ผลคลาด (Baldauf, 2003)

3.2 พิจารณาเลือกที่จะสร้างทรีเป็นแบบมีราก หรือไม่มีราก

ในบางกรณีที่พิจารณาเพื่อตรวจสอบว่ากลุ่มยีนเป็นออร์โธลอกกันหรือไม่ อาจใช้วิธีแบบไม่มีราก แต่ส่วนมากแล้วการสร้างทรีแบบมีราก คือทรีที่แสดงบรรพบุรุษร่วมของทุกสาขาที่ศึกษาจะเป็นพื้นฐานความเข้าใจกระบวนการวิวัฒนาการ ดังนั้น จึงต้องสร้างทรีแบบมีราก วิธีการทำได้สองวิธี คือ การใส่ลำดับของพวกนอกกลุ่มเข้าไปในการสร้างทรี โดยเลือกหาพวกนอกกลุ่มจากข้อมูลหรือจากลำดับของการจัดเรียงหรือโดยใช้ยีนที่ถูกจำลองมา (duplicated gene) โดยใช้ลำดับจากเซลล์หนึ่งไปใช้เป็นพื้นฐานของอีกเซลล์หนึ่ง

3.3 ตรวจสอบการจัดเรียงแล้วตัดสินใจเลือกที่จะรวมส่วนไหนและตัดส่วนไหนออกจากการวิเคราะห์

โดยพื้นฐาน คือ แต่ละคอลัมน์ของการจัดเรียงหรือแต่ละเรลิติวซ์จากแต่ละลำดับนั้นเป็นโฮโมโลกกัน นั่นคือ วิวัฒนาการจากตำแหน่งเดียวกันบนลำดับโปรตีนของบรรพบุรุษร่วมกันโดยไม่มีการแทรกหรือตัด จำไว้ว่าไฟโลจีนีในระดับโมเลกุลจะมีคุณภาพดีได้เท่าๆ กับคุณภาพการจัดเรียงลำดับที่นำมาใช้สร้าง หลักการทั่วไป คือ ให้ตัดส่วนใดก็ตามที่มีช่องว่าง (gaps) รวมทั้งคอลัมน์ที่ยังคลุมเครือออกไป เพราะไม่สามารถเชื่อมั่นได้เลยว่าบริเวณเหล่านี้จัดเรียงกันอย่างถูกต้อง และถ้าการจัดเรียงไม่ถูกต้องแล้วก็จะไม่ได้ข้อมูลไฟโลจีนีที่ต้องการ นอกจากนี้แม้ว่าจะทราบแน่นอนว่าบริเวณที่มีช่องว่างนั้นเป็นการจัดเรียงอย่างถูกต้องบริเวณนั้นก็ยังคงมีผลต่อทรีต่อไป เพราะว่าช่องว่างยังมีขนาดใหญ่ก็จะมีจำนวนตัวอักษรที่รวมเป็นกลุ่มร่วมกันมากขึ้น (เนื่องจากในความเป็นจริงช่องว่างหนึ่งๆ เป็นเพียงเหตุการณ์ทางวิวัฒนาการครั้งเดียว โดยไม่ขึ้นกับขนาด) ความสำคัญของช่องว่างช่องหนึ่งจะไม่เป็นสัดส่วนเดียวกับขนาดของช่อง

3.4 เลือกวิธีการสำหรับใช้คำนวณไฟโลจีนีในทรี

วิธีที่ใช้คำนวณแบ่งเป็นสองกลุ่ม (Ilman & Rhodes, 2004; ITC Academic Computing Health Science, 1998)

3.4.1 วิธีที่ใช้เมตริกซ์ของระยะห่าง (Distance-matrix methods)

วิธีนี้จะสรุปความแตกต่างระหว่างลำดับโดยการคำนวณค่าระยะห่างเป็นคู่ๆ ระหว่างลำดับที่มาจัดเรียงทั้งหมด อาจเรียกว่าวิธีจัดรวมกลุ่ม หรือวิธีเลขคณิต (clustering or algorithmic methods) วิธีนี้มีการสร้างตารางเมตริกซ์ ซึ่งบอกระยะห่างหรือความแตกต่างเป็นค่าตัวเลข ระหว่างลำดับเบสหรือโปรตีนต่างๆ ที่เปรียบเทียบ โดยเปรียบเทียบลำดับแต่ละคู่ และหาความแตกต่างระหว่างลำดับ จากนั้นจับกลุ่มให้ลำดับวิธีต่อไปนี้จะจัดอยู่ในกลุ่มนี้

7. FM (unweighted pair-group method with arithmetic mean) วิธีนี้จะสร้างทรีชนิดมีราก และใช้สมมติฐานนาฬิกาโมเลกุล ทรีที่ได้จะวางแทกษาทั้งหมดอยู่ในระยะห่างจากรากเท่ากัน ดังนั้นปริมาณการกลายพันธุ์จากรากจนถึงแต่ละแทกษาจะเท่ากัน

Fitch-Margoliash วิธีนี้อาศัยพื้นฐานเดียวกับวิธี 7. FM แต่จะใช้เทคนิคการรวมแทกษาที่นำมาวิเคราะห์บางแทกษาเข้าเป็นกลุ่มเดียวกันในลักษณะที่จะทำให้เหลือ 3 กลุ่ม ที่จะเปรียบเทียบ แล้วใช้สูตรคำนวณระยะห่างแบบสามจุด (three-point formula) ในแต่ละครั้ง เมื่อจับคู่แทกษาใด ๆ ได้แล้ว ก็จะแยกกระจายกลุ่มของแทกษาที่รวมไว้ในครั้งแรกออก แล้วทำการคำนวณใหม่ ทำซ้ำเช่นนี้ต่อไปจนได้ผลสมบูรณ์ วิธีนี้มักจะได้ทรีแบบเดียวกับวิธี 7. FM แต่ทรีที่ได้จะเป็นทรีที่ไม่มีราก

5 neighbor-Joining วิธีนี้มาจากการพิจารณาทรีที่ไม่มีรากที่ประกอบด้วย 4 แทกษา พบว่า ผลบวกระยะห่างระหว่างคู่ของเพื่อนบ้านทั้งสองคู่จะน้อยกว่าผลบวกระยะห่างระหว่างคู่ที่ไม่ใช่เพื่อนบ้านเสมอ เรียกว่า เป็นเงื่อนไขสี่จุด (4-point condition) และใช้เป็นพื้นฐานในการคำนวณในการทดสอบความถูกต้องของวิธีที่ใช้สร้างทรี ทั้งการจำลองการกลายพันธุ์ของดีเอ็นเอและการใช้แทกษาจริงที่ทราบทรีแน่นอน พบว่า วิธีนี้เป็นวิธีที่ให้ผลดีกว่าวิธี 7. FM และ Fitch-Margoliash แม้ว่าสองวิธีหลังจะให้ผลเชื่อถือได้ในบางสถานการณ์ แต่วิธี neighbor-Joining จะใช้ได้กับข้อมูลในช่วงกว้างมากกว่าและถ้าข้อมูลไม่อยู่บนพื้นฐานนาฬิกาโมเลกุล ซึ่งปัจจุบัน

ข้อมูลส่วนมากไม่อยู่บนสมมติฐานนี้ วิธี neighbor-joining จะดีกว่า เนื่องจากไม่ได้ตั้งบนสมมติฐานนั้น (ดังนั้น สองวิธีแรกในทางปฏิบัติจึงไม่ค่อยได้ใช้จริง)

ในชุดโปรแกรมไฟลพ์ ในขั้นตอนการสร้างทรีด้วยวิธีสร้างแมทริกซ์ระยะทางจะเริ่มต้นด้วยโปรแกรมสำหรับสร้างแมทริกซ์ (1 nadist สำหรับดีเอ็นเอ และ rodinist สำหรับ โปรตีน) แล้วนำผลไปสร้างทรีโดยใช้โปรแกรม Fitch (วิธี Fitch-Margolish) หรือ 5 eighxor (วิธี 5 eighxor-joining) หรือ o itsch อย่างใดอย่างหนึ่ง (Tuimala, 2006)

3.4.2 วิธีที่ใช้ข้อมูลที่ไม่เกี่ยวข้องกัน เช่น ตัวอักษร (Discrete data methods)

วิธีนี้อาจเรียกว่าวิธีการค้นหาทรี (tree searching methods) โดยใช้ข้อมูลในลักษณะตัวอักษรตลอดทั้งลำดับ วิธีที่จัดอยู่ในกลุ่มนี้คือ พาร์สิโมนี (parsimony) วิธีนี้ยึดถือหลักทรีที่เหมาะสมคือ ทรีที่มีการกลายพันธุ์หรือการเปลี่ยนแปลงน้อยที่สุด ดังนั้น จะนับจำนวนการกลายพันธุ์ที่น้อยที่สุดในแต่ละทรีเรียกเป็นคะแนนพาร์สิโมนี (parsimony score) แล้วเลือกทรีที่คะแนนน้อยสุดเป็นคำตอบ (O'Rand & Szymura, 2003)

แมกซิมัมไลค์ลิฮูด (maximum likelihood) วิธีนี้เป็นการหาค่าไลค์ลิฮูดของพารามิเตอร์โมเดลที่เป็นพื้นฐานของข้อมูล คือ การนำเอาข้อมูลที่มีอยู่มาใช้ประมาณค่าพารามิเตอร์ โดยใช้วิธีการแทนค่าพารามิเตอร์ต่าง ๆ ลงในโมเดลของความน่าจะเป็น แล้วพิจารณาคุณค่าพารามิเตอร์ที่มีค่าไลค์ลิฮูดสูงสุดหรือทำให้ข้อมูลนั้นน่าจะเป็นไปได้ที่สุด (Urbell, 2007) ในโปรแกรมไฟลพ์จะใช้ตารางสถิติในการคำนวณค่าไลค์ลิฮูดของการเปลี่ยนแปลงที่จะเกิดขึ้นที่ตำแหน่งหนึ่ง ๆ วิธีนี้จะคำนวณหาความน่าจะเป็นที่ชุดข้อมูลหนึ่ง ๆ จะเหมาะสมกับทรีที่สร้างมาจากชุดข้อมูลนั้น โดยใช้โมเดลของวิวัฒนาการของลำดับที่กำหนดให้ ดังนั้น ในขั้นแรกจะเปรียบเทียบข้อมูลกับชุดของโมเดลของการวิวัฒนาการของลำดับ จากนั้นจะเลือกโมเดลที่สามารถอธิบายรูปแบบของการเปลี่ยนแปลงลำดับได้ดีที่สุด (แต่ผู้ใช้โปรแกรมอาจจะเลือกโมเดลใดก็ได้ตามต้องการ) จากนั้นจึงนำโมเดลการวิวัฒนาการของลำดับที่ได้นั้นไปใช้ในการวิเคราะห์หาค่าไลค์ลิฮูด การวิเคราะห์จะเริ่มต้นจากการสร้างทรีจากชุดข้อมูลนั้น จากนั้นย้ายส่วนของแขนงไปส่วนต่างๆ จนทรีที่ได้ค่าคะแนน likelihood (ค่าความน่าจะเป็นของความเหมาะสมกับข้อมูลสูงสุด) ซึ่งค่าคะแนนแปรตามโทโพโลยี และความยาวของแขนงของทรี (Tuimala, 2006)

วิธีของเบย์ (Bayesian method) เป็นวิธีหาค่าไลค์ลิฮูดอีกชนิดหนึ่งที่ได้รับความนิยม มีสมมติฐานมากขึ้นและใช้อัลกอริทึมที่แตกต่างกัน (Mau et al., 1999)

ในชุดโปรแกรมไฟลพ์ จะมีโปรแกรมหลายชนิดให้เลือกใช้ตามวิธีคำนวณที่ต้องการ (แสดงในตารางที่ 1) เมื่อเปรียบเทียบทั้งสองวิธีวิธีหาระยะทางมักจะใช้ในการสร้างทรีในขั้นต้น ถ้านำไปใช้เป็นคำตอบทันทีต้องระมัดระวัง ส่วนวิธีพาร์สิโมนีและไลค์ลิฮูดเป็นวิธีที่ดีกว่า เนื่องจากเป็นวิธีที่สำรวจความสัมพันธ์ระหว่างทรีและองค์ประกอบอื่นๆ อย่างละเอียด แต่โดยทั่วไปแล้วส่วนมากให้ผลเหมือนกัน ความแตกต่างของทั้งสองกลุ่มที่สำคัญคือ ความยากง่าย

ความรวดเร็วและข้อมูลที่ได้รับบางครั้ง นอกจากนี้ไม่ว่าจะด้วยเหตุผลใดก็ตาม การใช้วิธีใดวิธีหนึ่งอาจให้ผลไม่ถูกต้อง ดังนั้น โดยทั่วไปทรีที่ได้ควรจะทดสอบโดยใช้มากกว่า 1 วิธี

3.5 พิจารณาเลือกโมเดลของการวิวัฒนาการ

วิธีคำนวณระยะทางของลำดับอาจใช้โมเดลการวิวัฒนาการที่แตกต่างกันได้ โดยโมเดลนี้คือ สูตรทางคณิตศาสตร์ที่จะชดเชยปัญหาการแทนที่ที่เกิดขึ้นซ้ำๆ กันหลายครั้ง (multiple substitutions) และการแทนที่ที่มีน้ำหนักไม่เท่ากันระหว่างการแทนที่แบบทรานสลิชัน คือ ระหว่างเบสภายในกลุ่มพหุมีดินหรือภายในกลุ่มพิวรีน (transition; Ts) และการแทนที่แบบทรานสเวอร์ชัน คือ ข้ามกลุ่มพหุมีดินหรือพิวรีน (transversion; Tv) ชนิดของโมเดลที่สำคัญ โดยเฉพาะที่สามารถเลือกได้จากโปรแกรมไฟลพ์และลักษณะของโมเดลโดยสรุปดังแสดงในตารางที่ 2

โมเดลหลายชนิดพัฒนาขึ้นเพื่อใช้ประเมินความแตกต่างที่แท้จริงระหว่างลำดับ โดยอาศัยการเปรียบเทียบชนิดของเบสในลำดับ ในปัจจุบันโมเดลหลายชนิดจะให้คะแนนแตกต่างกันระหว่างลำดับ โมเดลจะแตกต่างกันตามน้ำหนักที่ใช้ในการแทนที่ที่ต่างกัน เช่น จะให้น้ำหนักการเกิดทรานสเวอร์ชันมากกว่าการเกิดทรานสลิชัน (อัตราส่วน Ts/Tv มักใช้ที่ค่า 1-2) หรือการปรับค่าแบบแกมมา (gamma corrections) เนื่องจากการนับความแตกต่างธรรมชาติระหว่างสองลำดับ อาจจะประเมินการวิวัฒนาการที่เกิดขึ้นจริงต่ำไป โดยเฉพาะที่ตำแหน่งที่มีวิวัฒนาการรวดเร็วมากและเกิดการกลายพันธุ์ซ้ำๆ (multiple mutations) คือมีการเปลี่ยนแปลงไปจากสิ่งที่ย่อยเปลี่ยนแปลงไปแล้ว วิธีสร้างไฟลพ์จีโนมดิกทรีมักจะสันนิษฐานว่าทุกตำแหน่งมีวิวัฒนาการในอัตราเดียวกัน ซึ่งในความเป็นจริงไม่ถูกต้อง เช่น โคดอนตำแหน่งที่สาม จะมีวิวัฒนาการเร็วกว่าตำแหน่งที่สอง เนื่องจากตำแหน่งที่สามไม่มีผลเปลี่ยนชนิดกรดอะมิโนเหมือนตำแหน่งที่สอง อัตราที่ต่างกันนี้สามารถแก้ได้โดยการใช้การกระจายแกมมา (gamma distribution) โดยรูปร่างการกระจายแบบแกมมาถูกกำหนดโดยพารามิเตอร์ที่เรียกว่า อัลฟา (α) ซึ่งจะกำหนดสัมประสิทธิ์ของการผันแปร (coefficient of variation; CV) ของอัตราการแทนที่ระหว่างตำแหน่ง ซึ่งค่า $CV = 0.1/\sqrt{\alpha}$ และค่าโดยประมาณของอัลฟา สำหรับยีนที่เก็บรหัสโปรตีนส่วนมากคือ 0.5 ซึ่งสามารถหาจาก โปรแกรม Tree-uzzle (Schmidt, et al., 2002) สำหรับลำดับโปรตีนมีโมเดลวิวัฒนาการให้เลือก 5 ชนิด คือ JTT (Jones-Taylor-Thornton), M (Dayhoff), MB (Blosum), oimura, และ Categories (Tuimala, 2006)

3.6 เลือกโปรแกรมที่จะใช้

โปรแกรมที่มีความสมบูรณ์มากที่สุดและใช้กันอย่างกว้างขวาง เช่น Hf LI, Mega, และ M ซึ่งต่างรวมโมเดลและวิธีการคำนวณหลากหลาย โดยโปรแกรม Mega3.1 ใช้ง่ายที่สุด เนื่องจากเป็นลักษณะของวินโดวส์ และให้เลือกรายการจากแถบรายการ ส่วน M มีความซับซ้อนมากที่สุดที่อยู่เว็บไซต์ของโปรแกรมต่างๆ ดังแสดงในตารางที่ 3

ตารางที่ 1 วิธีคำนวณของโปรแกรมต่างๆในชุดโปรแกรมไฟลิป เวอร์ชัน 3.66

โปรแกรม	พื้นฐานวิธีคำนวณ
Dnadist	Distance matrix (สำหรับดีเอ็นเอ)
Protdist	Distance matrix (สำหรับโปรตีน)
Dnapars	Parsimony (สำหรับดีเอ็นเอ) ใช้อัลกอริทึมชนิดฮิวริสติก (heuristic)
Dnapenny	Parsimony (สำหรับดีเอ็นเอ) (branch-and-bound algorithm, วิธีนี้รับรองว่าจะได้รหัสสั้นที่สุด)
Protpars	Parsimony (สำหรับโปรตีน)
Fitch	Fitch-Margoliash (สำหรับสร้าง ตรี จากระยะห่าง)
Neighbor	Neighbor-joining (สำหรับสร้าง ตรี จากระยะห่าง)
Kitsch	สำหรับสร้าง ตรี จากระยะห่าง
Modeltest	Maximum likelihood
MrModeltest	Maximum likelihood
MrBayes	Bayesian

ตารางที่ 2 โมเดลวิวัฒนาการชนิดต่างๆ (Allman & Rhodes, 2004)

โมเดล	คำอธิบายลักษณะโมเดล
Jukes-Cantor distance	เป็นโมเดลที่ใช้พารามิเตอร์ค่าเดียว (α) ในการกำหนดอัตราการกลายพันธุ์ คือ ใช้สมมติฐานว่า การแทนที่ทั้งหมดสามารถเกิดได้เท่า ๆ กัน (และที่สมดุลเบสทุกชนิดมีความถี่เท่ากับ 0.25)
Kimura distance	เป็นโมเดลที่กำหนดอัตราการเปลี่ยนแปลงแตกต่างกัน 2 ชนิด คือ อัตราหนึ่งสำหรับการแทนที่แบบทรานสชัน (อัตราเบตา) และ อีกอัตราสำหรับการแทนที่แบบทรานสเวอร์ชัน (อัตราแกมมา) (ที่สมดุลเบสทุกชนิดมีความถี่ 0.25) โมเดลนี้เหมาะกับข้อมูลที่เห็นว่าการแทนที่ทั้งสองแบบเกิดในอัตราที่ไม่เท่ากัน
F84 distance	โมเดลนี้กำหนดพารามิเตอร์ของอัตราเป็นสองอัตรา คือ อัตราสำหรับการแทนที่แบบทรานสชัน และ อัตราสำหรับการแทนที่แบบทรานสเวอร์ชัน (แต่ความถี่ที่สมดุลของแต่ละเบสจะมีค่าแตกต่างกันได้)
LotDet distance	โมเดลนี้ใช้การแปลงค่าด้วยการหาดีเทอร์มิแนนท์ของแมทริกซ์ ตามด้วยการหาค่าลอการิทึมฐานธรรมชาติ (Lockhart et al., 1994) โมเดลนี้ใช้เมื่อมีความแตกต่างความถี่ของเบสสูงมากระหว่างลำดับในตรี แต่โมเดลนี้ไม่สามารถใช้กับอักษรที่ยังคลุมเครือ เช่น เบส P, R, หรือ N เป็นต้น ดังนั้น ลำดับจะต้องทราบแน่นอนทั้งหมด (Tuimala, 2006)

ตารางที่ 3 โปรแกรมสำหรับสร้างไฟโลจีเนติกทรี

โปรแกรม	ที่อยู่เว็บไซต์
PHYMLIP	http://evolution.genetics.washington.edu/phymlip.html
PAML	http://abacus.gene.ucl.ac.uk/software/paml.html
TREE-PUZZLE	http://www.tree_puzzle.de
Mega3.1 (latest version: 4.0)	http://www.megasoftware.net/
PAUP*	http://paup.csit.fsu.edu/index.html
Phylowin	http://pbil.univ-lyon1.fr/software/phylowin.html
Treeview	http://taxonomy.zoology.gla.ac.uk/rod/treeview.html
Seaview	http://pbil.univ-lyon1.fr/software/seaview.html
ClustalW	http://www.bioinformatics-toolkit.org/Help/Topics/clustalw.html

การสร้างหรืออย่างคร่าวๆ อาจใช้โปรแกรมคลัสตัล (Clustal) ซึ่งคำนวณหาระยะทางวิวัฒนาการ (evolutionary distances) โดยใช้ค่าสำคัญในเมนู ทรี มี 4 ขั้นตอน คือ 1. ตัดส่วนของ การจัดเรียงลำดับที่มีช่องว่างออกโดยเลือกคำสั่ง delete positions with gaps 2. ถ้าลำดับเหมือนกันน้อยกว่า 95 เปอร์เซ็นต์ ควรเลือกที่จะปรับการวัดระยะทางสำหรับการเกิดการแทนที่หลายครั้ง (multiple substitutions) โดยเลือกคำสั่ง correct for multiple substitutions แล้ว 3. ตรวจสอบว่าได้เลือกใช้ฟอร์แมตที่ถูกต้องสำหรับผล โดยเลือกที่ output format แล้วเลือก link จากนั้นเปลี่ยนที่คำสั่ง bootstrap label options จาก branches ไปเป็น nodes แล้วปิด 4. คำนวณหาระยะทาง โดยเลือกคำสั่ง bootstrap 5J tree ขั้นตอนสุดท้ายนี้จริง ๆ แล้ว ประกอบด้วยหลายขั้นตอนย่อย คือ มีโปรแกรมคำนวณระยะทาง สร้างทรีด้วยวิธี neighbour-joining จากนั้น ประเมินทรีที่ได้จากการทดสอบสถิติที่เรียกว่า บุทสแตรพ์ฟิง (รายละเอียดในขั้นที่ 4) จากนั้น จึงดึงข้อมูลทั้งหมดมารวมกันลงในไฟล์ผลลัพธ์อันเดียว ซึ่งจะประกอบด้วยทรีที่มีค่าบุทสแตรพ์ฟิงที่เหมาะสม เมื่อต้องการดูผลจะต้องใช้โปรแกรมสำหรับใช้ดูทรี เช่น ทรีวิว (Treeview) วิธีการสร้างทรีจากโปรแกรมคลัสตัลนี้มักเป็นวิธีที่ดีสำหรับการเริ่มต้น เช่น การจัดเรียงลำดับจำนวนมาก ๆ อย่างรวดเร็ว เพื่อเลือกเอาชุดตัวแทนมาใช้หรือเพื่อตัดสินใจต่อไปว่า ข้อมูลที่ได้มานั้นมีความน่าสนใจเพียงพอกี่จะศึกษาต่อในรายละเอียดหรือไม่

ขั้นที่ 4 การทดสอบ

4.1 บุทสแตรพ์ฟิง (Bootstrapping)

เป็นวิธีทดสอบที่ง่ายที่สุด สำหรับตรวจสอบความถูกต้องของไฟโลจีนิติกทรี หลักการสำคัญของการวิเคราะห์อยู่บนพื้นฐานง่าย ๆ คือ ถ้าความเกี่ยวข้องระหว่างลำดับสองลำดับมีมาก ๆ การเปลี่ยนแปลงเบสเพียงบางเบสแบบไม่จำเพาะเจาะจงไม่น่าจะมีผลต่อความเกี่ยวข้องนี้ ตรงกันข้ามถ้าทั้งสองลำดับเกี่ยวข้องกันน้อยมาก การเปลี่ยนแปลงบางเบสจะมีผลอย่างมาก คือ เป็นการทดสอบว่าชุดของข้อมูลทั้งหมดสอดคล้องกับทรีหรือไม่ หรือว่าทรีที่ได้เป็นเพียงทรีหนึ่งในบรรดาทรีอื่น ๆ ทั้งหมดที่ตีเท่าเทียมกันจำนวนมาก โปรแกรมที่ใช้ คือ เซคบูท (Seqboot)

ขั้นตอนการทดสอบนี้ประกอบด้วย 3 ขั้นตอน คือ การสุ่มตัวอย่างย่อย ๆ ในแต่ละคอลัมน์มาจากคอลัมน์ใด ๆ ของชุดข้อมูล (ตัวอย่างย่อยนี้มีขนาดความยาวเท่าลำดับข้อมูลเดิม) จากนั้นนำมาสร้างทรีต่าง ๆ สำหรับแต่ละตัวอย่างย่อย เช่น จำนวน 1,000 ทรี แล้วคำนวณความถี่ที่ส่วนต่าง ๆ ของทรีนั้นถูกสร้างขึ้นในแต่ละตัวอย่างย่อยที่สุ่มขึ้นมาเหล่านั้น เช่น ถ้าพบการจับกลุ่มของลำดับแบบหนึ่งในทุกทรีของทุกตัวอย่างย่อย ก็หมายความว่า ค่าสนับสนุนบุทสแตรพ์ฟิง (bootstrap support) เป็น 100 เปอร์เซ็นต์ แต่ถ้าพบเพียง 2 ใน 3 ของทรีของตัวอย่างย่อย ก็แสดงว่าค่าสนับสนุนบุทสแตรพ์ฟิงเท่ากับ 67 เปอร์เซ็นต์เท่านั้น ค่าจะเขียนไว้ที่โนดของทรีเรียก ความน่าจะเป็นโพสทีเรียรี่ (posterior probability)

4.2 แจ็คไKnife (Jack-knife)

เป็นเปอร์เซ็นต์จำนวนครั้งที่เคลตปรากฏ เมื่อจำนวนเปอร์เซ็นต์ของตัวอักษรถูกเอาออกอย่างสุ่มจากชุดข้อมูลแล้วทำการวิเคราะห์อีกครั้ง

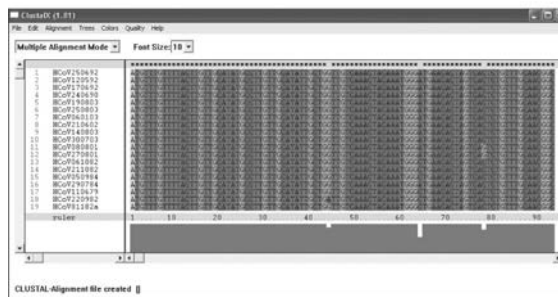
4.3 แขนงยาว (Long branches)

ปัญหาแขนงยาว คือ การที่มีการโน้มเอียงของลำดับที่แยกออกมาก ๆ จะจับกลุ่มกันในทรี โดยไม่เป็นความสัมพันธ์แท้จริง อย่างน้อยส่วนหนึ่งเกิดเนื่องจากลำดับวิวัฒนาการอย่างรวดเร็ว หรือ ลำดับที่ไม่มีสายสัมพันธ์ใกล้ชิดจะมีการกลายพันธุ์จำเพาะมากมาย แต่เนื่องจากความจำกัดที่จำนวนชนิดกรดอะมิโนมี 20 ชนิด หรือนิวคลีโอไทด์มี 4 ชนิด ทำให้การเปลี่ยนแปลงจำนวนมากเริ่มที่จะเกิดความคล้ายกัน ถ้าแขนงมีความยาวมาก ความคล้ายกันเหล่านี้จะกลบกลืนของไฟโลจีนีที่แท้จริง แล้วลำดับเหล่านั้นจะดึงดูดกันเอง ซึ่งจะเป็นปัญหาหลายอย่าง ปัญหาหนึ่ง คือ ทำให้ค่าบุทสแตรพ์ฟิงของทรีไม่ตี วิธีการทดสอบว่ามีลักษณะของแขนงยาวเกิดขึ้นในชุดข้อมูลหรือไม่ คือ ตัดลำดับเหล่านี้ออกไปจากชุดข้อมูล แล้วหาค่าบุทสแตรพ์ฟิงใหม่ดูว่าค่าจะเพิ่มขึ้นหรือไม่ การใช้วิธีอื่น เช่น แมกซิมัมไลค์ลิชูดจะถูกผลกระทบจากปัญหานี้น้อยกว่า วิธีที่ดีที่สุดในการแยกส่วนของแขนงยาว ทำโดยการหาลำดับที่อยู่ระหว่างกลางมาใส่รวมในทรี แล้วตรวจสอบว่าลำดับนั้นเข้ากันดีกับลำดับที่เหลือหรือไม่

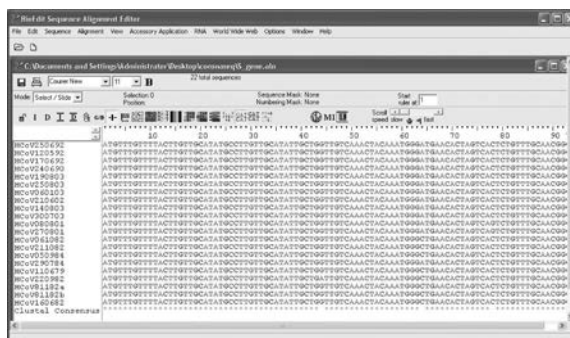
ขั้นที่ 5 การนำเสนอข้อมูล

การนำเสนอไฟโลจีนิติกทรีมีหลักปฏิบัติที่ยอมรับ คือ 1. ความยาวของแขนงจะเขียนตามสเกล (คือ เป็นสัดส่วนกับปริมาณวิวัฒนาการ) ทุกครั้ง แม้ว่าความสัมพันธ์ระหว่างความยาวของแขนงและช่วงเวลาจริงจะห่างไกลกัน ถ้าเขียนตรงไปตรงมาอาจไม่น่าเชื่อถือสำหรับยื่น ๆ เดียว แต่ความยาวจะยังคงให้ภาพที่ดีสำหรับแสดงอัตราการเปลี่ยนแปลงสัมพัทธ์ตลอดทรี 2. ให้ออกค่าบุทสแตรพ์ฟิงเป็นเปอร์เซ็นต์ ซึ่งจะทำให้ทรีอ่านและเปรียบเทียบกันได้ง่ายและรายงานผลเพียงถ้าค่าบุทสแตรพ์ฟิงสูงกว่าหรือเท่ากับ 50 เปอร์เซ็นต์ โดยการแก้ไขที่ไฟล์ของทรี 3. เลือกใช้ชื่อ 677 ที่มีความหมาย

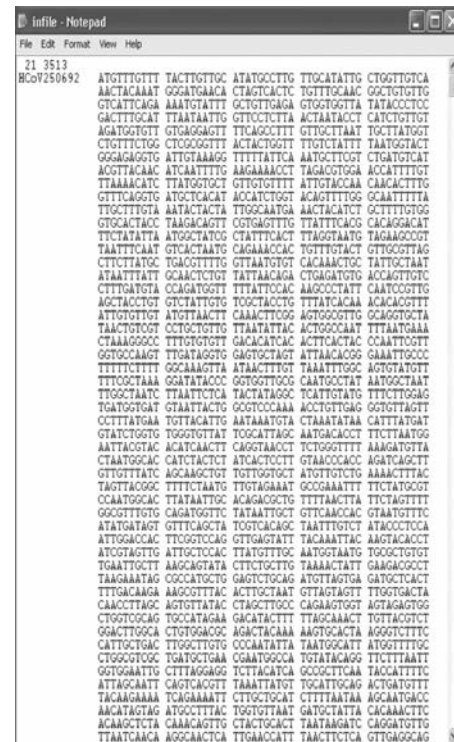
ตัวอย่างการสร้างไฟโลจีนิติกทรีของไวรัสโคโรนา สายพันธุ์ 229E จากลำดับดีเอ็นเอของยีน spike (spike) ซึ่งมีความยาว 3513 นิวคลีโอไทด์ จำนวน 21 ลำดับ แสดงในรูปที่ 2 ลำดับนิวคลีโอไทด์ถูกเรียกมาจากฐานข้อมูล GenBank แล้วนำมาจัดเรียงในฟอร์แมตของ fasta จากนั้นใช้โปรแกรม Clustal ในการจัดเรียงลำดับทั้งหมดดังแสดงในรูปที่ 2ก จากนั้นใช้โปรแกรม Bioedit ในการตรวจสอบความเหมาะสมของการจัดเรียง ในกรณีนี้ไม่มีช่องว่างและการจัดเรียงจากโปรแกรมดีแล้วจึงไม่มีการแก้ไข (รูปที่ 2ข) จึงบันทึกไฟล์ด้วยนามสกุล phy เพื่อให้ไฟล์อยู่ในฟอร์แมตไฟล์ แสดงในรูปที่ 2ค เปลี่ยนชื่อเป็น infile แล้วนำไปวิเคราะห์ด้วยโปรแกรมเซคบูท นำผลลัพธ์ที่ได้มาวิเคราะห์ด้วยวิธีพาร์สิโมนี โดยใช้โปรแกรม w napars และวิเคราะห์หาทรีที่เป็นตัวแทนทรีทั้งหมดที่ได้ด้วยโปรแกรม Consense ตามลำดับ จากนั้นจึงเปิดดูทรีด้วยโปรแกรมสำหรับใช้ดูภาพและแก้ไขทรี ซึ่งมีมากมาย ในที่นี้ใช้ทรีวิว (age, 1W6) และ เลือกแสดงทรีแบบไม่มีราก บันทึกไฟล์แบบภาพกราฟิก แล้วใช้โปรแกรมแต่งภาพต่าง ๆ ตกแต่งเพื่อนำเสนอ ในกรณีนี้มีการใช้สีพื้นหลังเพื่อชี้ให้เห็นว่าไวรัสทั้งหมดนี้สามารถจัดจำแนกเป็นสี่กลุ่ม (รูปที่ 2ง) ทรีที่ได้จากการวิเคราะห์วิธีนี้มีความคล้ายคลึงกับทรีที่มีการตีพิมพ์ก่อนหน้านี้ (Chixo & Birch, 2006) ซึ่งวิเคราะห์โดยใช้โปรแกรม w nadist ด้วยโมเดล kimura ค่า transition/transversion ratio เท่ากับ 1.62 ค่า CV เท่ากับ 2.5 (โดยแทนค่าอัลฟาเท่ากับ 0.16)



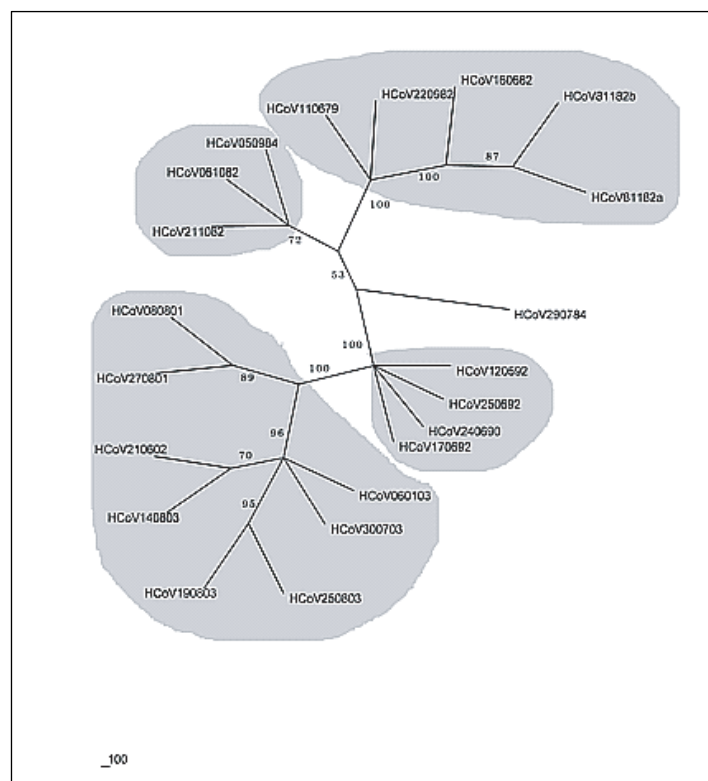
(ဂ)



(๒)



(ค)



(१)

รูปที่ 2 ตัวอย่างการสร้างไฟโลเจเนติกทรีของไวรัสโคโรนาสายพันธุ์ 229E ไฟโลเจเนติกทรีสร้างจากลำดับดีเอ็นเอของยีน spike จำนวน 21 ไวรัสที่แยกได้ในเวลาที่แตกต่างกัน เลขเรียกข้อมูล (accession number) ตั้งแต่ DQ243964-DQ243987 ก.) ผลจากการจัดเรียงด้วยโปรแกรม clustalX ข.) การตรวจการจัดเรียงด้วยโปรแกรม Bioedit ค.) การตรวจดูไฟโลพอร์เมทไฟลิก่อนใช้วิเคราะห์ด้วย PHYLIP3.66 ง.) ไฟโลจีนีติกทรีแบบไม่มีราก ตัวเลขค่าบูทสแตรฟีแสดงที่ตำแหน่งโนด สเกลระยะห่างวิวัฒนาการ (evolutionary distance) แสดงด้วยเส้นขีดได้ภาพ แทกซอนแสดงด้วยชื่อย่อไวรัสพร้อมวันที่ไวรัสถูกแยกได้ เพื่อให้การนำเสนอชัดเจนได้ใช้สีพื้นหลังเพื่อแสดงกลุ่มของไวรัส

บทสรุป

การสร้างไฟโลจีนติกที่สามารนำไปใช้เชิงประยุกต์โดยนำมาวิเคราะห์เพื่อหาข้อมูลที่เป็นประโยชน์ในด้านต่างๆ เช่น การควบคุมโรคติดต่อ การออกแบบยารักษาและการค้นหาสารมีฤทธิ์ทางชีวภาพ ขั้นตอนในการสร้างที่สำคัญมี 5 ขั้นตอน คือ การรวบรวมลำดับการจัดเรียงลำดับการเลือกใช้อัลกอริทึมวิเคราะห์โมเดลและโปรแกรมสำหรับสร้าง การทดสอบ และการนำเสนอข้อมูล

เอกสารอ้างอิง

- Allman, E. S., & Rhodes, J. A. (2004). *Mathematical models in Biology: An introduction*. New York: Cambridge University Press.
- Baldauf, S. L. (2003). Phylogeny for the faint of heart: A tutorial. *Trends in Genetics*, 19, 345-351.
- Chibo, D., & Birch, C. (2006). Analysis of human coronavirus 229E spike and nucleoprotein genes demonstrates genetic drift between chronologically distinct strains. *Journal of General Virology*, 87, 1203-1208.
- Deonier, R. C., Tavaré, S., & Waterman, M. S. (2005). *Computational genome analysis: An introduction*. New York: Springer Science Bussiness Media, Inc.
- Fry, B. G., Vidal, N., Norman, J. A., Vonk, F. J., Scheib, H., Ramjan, S. F., et al. (2006). Early evolution of the venom system in lizards and snakes. *Nature*, 439, 584-588.
- Harrison, C. J., & Langdale, J. A. (2006). A step by step guide to phylogeny reconstruction. *The Plant Journal*, 45, 561-572.
- ITC-Academic Computing Health Science. (1998). *Phylip: Phylogenetic analysis workshop*. Virginia: University of Virginia.
- Jiang, P., Faase, J. A. J., Toyoda, H., Paul, A., Wimmer, E., & Gorbalenya, A. E. (2007). Evidence for emergence of diverse polioviruses from C-cluster coxsackie A viruses and implications for global poliovirus eradication. *Proceedings of the National Academy of Sciences of the United States of America*, 104, 9457-9462.
- Joshi, B., Lee, K., Maeder, D. L., & Jagus, R. (2005). Phylogenetic analysis of eIF4E-family members. *BMC Evolutionary Biology*, 5, 48.
- Krane, D. E., & Raymer, M. L. (2003). *Fundamental concepts of bioinformatics*. San Francisco: Benjamin Commings.
- Lockhart, P. J., Steel, M. A., Hendy, M. D., & Penny, D. (1994). Recovering evolutionary trees under a more realistic model of sequence evolution. *Molecular Biology and Evolution*, 11, 605-612.
- Madisch, I., Hofmayer, S., Moritz, C., Grintzalis, A., Hainmueller, J., Pring-Akerblom, P., et al. (2007). Phylogenetic analysis and structural predictions of human adenovirus penton proteins as a basis for tissue specific adenoviral vector design. *Journal of Virology*, 81, 8270-8281.
- Mau, B., Newton, M. A., & Larget, B. (1999). Bayesian phylogenetic inference via Markov chain Monte Carlo methods. *Biometrics*, 55, 1-12.
- Page, R. D. M. (1996). TREEVIEW: An application to display phylogenetic trees on personal computers. *Computer Applications in the Biosciences*, 12, 357-358.
- Purcell, S. (2007). BGIM: Maximum likelihood estimation primer. Retrieved July 16, 2007, from http://statgen.iop.kcl.ac.uk/bgim/mle/sslike_1.html
- Schmidt, H. A., Strimmer, K., Vongron, M., & von Haeseler, A. (2002). TREE_PUZZLE: Maximum likelihood phylogenetic analysis using quartets and parallel computing. *Bioinformatics Applications Note*, 18, 502-504.
- Smith, W. L., & Wheeler, W. C. (2006). Venom evolution widespread in fishes: A phylogenetic road map for the bioprospecting of piscine venoms. *Journal of Heredity*, 97, 206-217.
- Tuimala, J. (2006). *A primer to phylogenetic analysis using the PHYLIP package*. Espoo: CSC-Scientific Computing Ltd.
- Turner, S. M., Chaudhuri, R. R., Jiang, Z. D., DuPont, H., Gyles, C., Penn, C. W., et al. (2006). Phylogenetic comparisons reveal multiple acquisitions of the toxin genes by enterotoxigenic *Escherichia coli* strains of different evolutionary lineages. *Journal of Clinical Microbiology*, 44, 4528-4536.