

การวางแบบและสถิติของ... ..

การวางแผนและสถิติวิจัยทางการเกษตร

EXPERIMENTAL DESIGNS
AND
STATISTICS FOR USE IN AGRICULTURAL RESEARCH

โดย

นายระพี สาคริก กส.บ., Cert., in Statistics and Agr.Exp. Design.

๒๕๑๒

(คัดลอกจากฉบับพิมพ์ปี พ.ศ. ๒๕๔๙)



ChangeFusion



เครือข่ายจิตอาสา
Volunteer Spirit Network

คำขอบคุณ

คณะผู้จัดทำขอกราบขอบพระคุณท่านอาจารย์ระพี สาคริก เป็นอย่างสูง
ที่ไกรฤณาอนุญาต ให้จัดพิมพ์โรเนียวตำราเล่มนี้ ขึ้นมาใหม่อีกครั้ง เพื่อเป็นวิทยาทาน
เผยแพร่แก่ผู้สนใจในวิชาสถิติ.

คณะผู้จัดทำ

น.ส.มะลิ	ศรีรุ่งโรจน์
นายไมตรี	ทองสวัสดิ์
นายสุจินต์	กีแท้

๔ มิถุนายน ๒๕๑๒

สารบัญ

	หน้า
ความมุ่งหมายของการศึกษาค้นคว้าทางการเกษตร	๒
ประชากร(Population)	๓
Mean	๔
การกระจายความถี่ (Frequency Distribution)	๑๐
การแจกแจงความถี่	๑๔
การหาค่าตัวกลาง	๒๑
Median	๒๒
Mode	๒๒
Deviation	๒๓
Standard Deviation	๒๔
Degree of Freedom	๒๕
การ Code ตัวเลข	๓๔
Replication	๔๖
Randomization	๔๘
การเปรียบเทียบความแตกต่างระหว่างสองสิ่ง	๕๔
Randomize Complete Block Design	๕๔
Paring Design	๕๗
Analysis of Variance	๖๐
Randomized Block Design	๖๓
การตรวจความแน่นอนของผลการทดลอง	๗๕
การหา Least Significant Difference (L.S.D.)	๘๒
Latin Square Design	๘๕
การหา Missing value Randomized Block Design	๙๗
การหา Missing value for Latin Square Design	๑๐๑
สหสัมพันธ์ (Correlation)	๑๐๓

Regression	๑๑๖
Analysis of Covariance	๑๓๑
✓ Split plot Design	๑๔๗
✓ Factorial Design	๑๖๐ ✓
Chi-Square	๑๖๗
✓ Check Method	๑๖๘

ปัจจุบันนี้เป็นสมัยที่วิทยาการต่าง ๆ กำลังเจริญก้าวหน้าอยู่เรื่อย ๆ โดยเหตุที่มนุษย์มีการค้นคว้าหาความรู้ ทำการทดลองค้นคว้า เพื่อหาหลักเกณฑ์ในทาง วิชาการอันจะนำมาใช้ปฏิบัติให้เป็นประโยชน์ในการสร้างความเจริญก้าวหน้าให้แก่ สังคมของมนุษย์ โดยเฉพาะอย่างยิ่งการ เกษตรอันนับว่าเป็นสิ่งจำเป็นสำหรับชีวิตแขนง จึงต้องมีการค้นคว้าทดลอง เพื่อหาความรู้ใหม่ ๆ มาใช้ปรับปรุงแก้ไขการผลิตพืชผลทาง การเกษตรให้ก้าวหน้าอยู่เสมอ ๆ ฉะนั้นการค้นคว้าทดลองจึงเปรียบเสมือนบันไดที่จะให้ไต่ขึ้นไปสู่ความเจริญรุ่งเรืองในยุคต่อ ๆ ไป แต่ถ้ามองจะพิจารณาถึงความ สำเร็จในการทดลองค้นคว้าแล้วจะพบความจริงว่า ย่อมขึ้นอยู่กับผลของงานเป็นส่วนสำคัญ ที่สุด เพราะว่าการทดลองการค้นคว้าทดลองมีความแน่นอนเป็นที่ เชื่อถือได้ที่จะนำไปปฏิบัติในอนาคต ก็ย่อมจะทำให้ปฏิบัติได้มีประสิทธิภาพยิ่งขึ้น รวมทั้งผู้ที่กระทำ การค้นคว้าทดลองในรอบหลัง ๆ ก็สามารที่จะยึดถือหลักเกณฑ์อันเป็นผลงานที่ได้ กระทำไว้แล้วดำเนินการต่อไปได้ ความเจริญก้าวหน้าก็ย่อมจะดำเนินไปอย่างไม่ หยุดยั้ง แต่ถ้ามองว่าการ ค้นคว้าทดลองนั้นได้รับผลซึ่งไม่มีหลักฐานให้เป็นที่ เชื่อถือได้แล้ว หากจะมีผู้ยึดถือผลงานนั้นเป็นหลักปฏิบัติก็ย่อมจะเป็นการเสี่ยงต่ออันตราย โดยเฉพาะอย่างยิ่งถ้านั้นต้องลงทุนลงแรงและเสียเวลามากมาย ถ้าหากนำหลักเกณฑ์ ที่ได้จากการทดลองซึ่งไม่ตรงกับความจริงตามกฎเกณฑ์ของธรรมชาติมาใช้เป็นหลัก ปฏิบัติ ก็ย่อมจะหมายถึงความเสียหายอย่างร้ายแรงโดยไม่ได้รับอะไรเป็นเครื่อง ทอแนะเลย ยิ่งกว่านั้นถ้าผู้นำเอาผลการทดลองที่ไม่ตรงกับความจริงไปใช้เป็น หลักปฏิบัติเมื่อไม่ได้รับผลตามนั้นก็จะเป็นการขาดความเชื่อถือ

ในการผลิตพืชผลทางการเกษตรใด ๆ ก็ตาม ผู้ผลิตย่อมจะมีความ มุ่งหมายกันสุดทุกอย่างเหมือนกันหมด คือ ให้ได้คุณภาพดีที่สุดในปริมาณสูงที่สุด แต่ลงทุนน้อยที่สุด ซึ่งทุนนี้รวมทั้งเงิน แรงงาน และเวลาที่ต้องเสียไป ฉะนั้นจุด มุ่งหมายอันใหญ่ยิ่งของการ ค้นคว้าทดลองจึงได้มุ่งไปสู่สิ่งเหล่านี้ และนำผลงานที่ได้ รัับออกเผยแพร่ให้ผู้ที่เกี่ยวข้องได้ยึดถือเป็นหลักปฏิบัติในอันที่จะปรับปรุงการผลิต

ของคนให้ได้รับผลดียิ่ง ๆ ขึ้นสืบไปแต่ตามที่โคกถาวมาแล้วว่าถาดลงงานที่เรา
จะยึดถือเป็นหลักปฏิบัติไม่มีสิ่งใดยืนยันเป็นหลักฐานให้เชื่อถือโคกเทากับว่าลงทุน
ลงแรงไปด้วยการเสี่ยง ยิ่งผู้ที่เขาเชื่อคำเล่าลือควยแล้ว การปฏิบัติของผู้นั้นก็ยอม
จะยึดอะไรเป็นหลักไม่ได้ เช่นถ้าใครเขาบอกว่าใส่ปุ๋ยอย่างไหนให้ผลดีก็เอาอย่าง
หรือการปลูกพบบโน้นก็แบบนี้ไม่ดีก็กระทำตาม โดยตัดสินใจเชื่อได้ง่าย ๆ ไม่มี
เหตุผลหรือหลักฐานใด ๆ ยืนยัน จึงเกิดมีปัญหาคือว่าทำอย่างไรจึงจะสามารถ
ทำให้ผลการคนควาทดลองนั้นได้รับผลเป็นที่แน่นอนเชื่อถือได้ และจะใช้อะไรเป็น
เครื่องวัดความแน่นอนของผลงานคนควาทดลอง

วิชาสถิติและการวางแผนวิจัยทางการเกษตร เป็นวิทยาการแขนงใหม่
แขนงหนึ่งที่ช่วยสร้างความเจริญก้าวหน้าให้แก่การเกษตรอย่างใหญ่หลวง จนเป็นที่
ที่รับรองยืนยันกันในหมู่ประเทศที่เจริญแล้วทั้งหลายทั่วโลก เนื่องจากเป็นวิชาการที่
ใช้เป็นเครื่องมือในการวางแผนการทดลองคนควาให้รัดกุม โดยมีเหตุผลเพื่อสร้าง
ความแน่นอนให้แก่งานเป็นครั้งแรกและครั้งสุดท้ายก็ยังใช้เป็นเครื่องมือวัดความแน่นอน
ของผลงานที่ใดอย่างถูกต้อง เพื่อจะได้ยึดถือเป็นหลักในการสรุปผลการทดลองนั้นให้
สามารถใช้เป็นกฎเกณฑ์แห่งความจริงที่เชื่อถือได้ ซึ่งจะเป็นประโยชน์แก่ผู้ที่นำไปใช้
ปฏิบัติหรือดำเนินการคนควาต่อไปโดยไม่มีการผิดพลาด ฉะนั้นผลการทดลองใด ๆ
ก็ตามอันมีเหตุที่สามารถจะเกิดความคลาดเคลื่อนไปจากความจริงได้ ถ้าไม่มีการ
ยืนยันทางสถิติแล้วก็ย่อมจะยึดถือเป็นหลักเกณฑ์ในการปฏิบัติต่อ ๆ ไปไม่ได้

ความมุ่งหมายในการคนควาทดลองทางการเกษตรทั่ว ๆ ไป อาจจำ
แนกออกเป็นแนวทางใหญ่ ๆ ได้สองอย่างด้วยกัน คือ

1. การคนควาเพื่อหาหลักเกณฑ์ใหม่ ๆ เพิ่มขึ้น (Induction)

เช่นการทดลองคนควาทดลองเปรียบเทียบพันธุ์ เปรียบเทียบส่วนผสมของปุ๋ย
เปรียบเทียบวิธีการปลูก การใส่ยาฆ่าแมลง ฯลฯ ทั้งนี้เพื่อนำหลักเกณฑ์พบใหม่
มาใช้เป็นประโยชน์ต่อไป ซึ่งนอกจากหลักเกณฑ์นั้น ๆ เราจะไม่เคยทราบกันมา

ก่อนแล้วบางทีก็อาจผู้อื่นเคยทำมาแล้วก็ได้ แต่ผลงานยังไม่เห็นชัดให้เป็นที่เชื่อถือก็
อาจจะเป็นกระทำได้ตามความมุ่งหมายเดิมเพื่อสร้างหลักฐานและความมั่นใจในผล
งานอีกทีหนึ่ง

2. การค้นคว้าเพื่อพิสูจน์สิ่งที่ทำการค้นคว้าให้เข้ากับกฎเกณฑ์ที่มีอยู่แล้ว
(Deduction) หมายความว่าใครมีหลักเกณฑ์หรือทฤษฎีเดิมอยู่แล้ว แต่เราจะต้อง
ทำการค้นคว้าเพื่อพิสูจน์ว่าผลงานของเรานั้น สามารถเข้ากับทฤษฎีหรือกฎเกณฑ์ใด
ตัวอย่างเช่นในการผสมพันธุ์ ใดมีทฤษฎีวางไว้ว่าลักษณะพอพันธุกรรมเป็นอย่างไร
อย่างนี้ ลูกที่ปรากฏออกมาจะต้องมีลักษณะผิดแผกกันไปเป็นอัตราส่วนเท่านั้น ๆ
เมื่อเราทำการค้นคว้าเกี่ยวกับการผสมพันธุ์ผลการตรวจลักษณะลูก ก็จะต้องถูกนำไป
เทียบกับอัตราส่วน ทางทฤษฎีว่าจะเข้ากฎเกณฑ์ของทฤษฎีใดหรือไม่

ประชากร (POPULATION)

ในการค้นคว้าทดลองทางการเกษตรนั้น เรามีความมุ่งหมายที่จะนำผล
การทดลองค้นคว้าไปใช้เป็นหลักเกณฑ์ที่จะปฏิบัติได้ทั่ว ๆ ไป ฉะนั้นในส่วนที่แท้จริงของ
สถิติวิจัยแล้วจำเป็นต้องมีความเกี่ยวข้องกับประชากรอยู่ด้วย

เมื่อกล่าวถึงคำว่า "ประชากร" โดยปกติเข้าใจกันว่าเป็นปวงประชาชนของ
ประเทศหรือของโลกซึ่งสุดแล้วแต่ขอบเขตที่จะนิยามไว้ แต่ในทางสถิติแล้ว คำว่า
ประชากร ย่อมหมายถึงจำนวนเต็มภายในขอบเขตที่ใดระบุไว้ของ มนุษย์ สัตว์ พืช
สิ่งของ หรือพฤติกรรมใด ๆ ที่จำเป็นต้องเกี่ยวข้องด้วย เช่นชาวพันธุ์ในแคว้นในประ
เทศไทย ก็มีจำนวนประชากรภายในขอบเขตที่ใดระบุไว้เพียงในประเทศ ถ้าเป็น
พันธุ์ชาวในแคว้นในโลก ก็จะมีขอบเขตที่ใดระบุไว้กว้างออกไปทั่วโลก เมื่อพูดถึง
นิสิตมหาวิทยาลัยเกษตรศาสตร์ ขอบเขตของประชากรก็มีเพียงในมหาวิทยาลัยเกษตรศาสตร์
ถ้ากล่าวถึงในค่านของพฤติกรรม เช่น การยิงเป้า 100 ครั้ง หรือการโยนหัวโยนก้อย
1,000 ครั้ง เป็นต้น

ถ้าหากจะกล่าวถึงขอบเขตของประชากรแล้ว เราสามารถที่จะแบ่งประชากรออกเป็นสองประเภทด้วยกัน คือ ประเภทหนึ่งเป็นประเภทประชากรที่มีขอบเขตซึ่งสามารถจะสิ้นสุดลงได้ เช่น จำนวนมนุษย์ จำนวนสัตว์ จำนวนพืช ถ้าเราสำรวจกันจริง ๆ ก็มีวันสิ้นสุดลงได้ ประชากรประเภทนี้เราเรียกว่า Finite Population ส่วนอีกประเภทหนึ่งเป็นประเภทที่ไม่มีการสิ้นสุด หรือการสิ้นสุดอยู่ไกลจนถือว่าไม่มี โดยทั่ว ๆ ไปเป็นประชากรของจำนวนพฤติกรรม เช่น การโยนหัวโยนก้อย เราจะโยนสักกี่ครั้งก็ยังไม่มีการสิ้นสุด เราเรียกประชากรประเภทนี้ว่า Infinite Population

DATA

ในการบันทึกปรากฏต่าง ๆ ไว้เป็นหลักฐานสำหรับการวิจัยนั้น โดยทั่ว ๆ ไปจำเป็นต้องบันทึกลงในลักษณะตัวเลข (data) และตัวเลขที่เราได้มาจากการสังเกตการ (Observation) โดยการวัด ซึ่ง ตรง นับจำนวนโดยตรง หรือโดยการแปรลักษณะเทียบให้เป็นตัวเลข เช่น ความเข้มข้นของสีเราสามารถแปรให้เป็นตัวเลขได้โดยเทียบกับสีมาตรฐานซึ่งมีความเข้มต่าง ๆ กัน และทุกระยะความเข้มก็มีตัวเลขกำกับไว้เป็นหน่วยของความเข้มประจำระยะนั้น ๆ

ปรากฏการณ์นี้อาจปรากฏขึ้นได้เป็นสองลักษณะด้วยกัน คือ คุณภาพ (Qualitative) ลักษณะหนึ่งกับเป็นปริมาณ (quantitative) อีกลักษณะหนึ่ง สำหรับลักษณะคุณภาพนั้นคงตัวอย่างเช่นสีของดอก ความเย็นของเมล็ด สีนัยตา เป็นต้น ซึ่งสิ่งเหล่านี้จะบันทึกได้โดยการพิจารณาและแบ่งคุณภาพ ส่วนลักษณะทางปริมาณ เช่น ความสูง ผลผลิต จำนวนดอก จำนวนเมล็ด ฯลฯ ซึ่งเราสามารถทำการวัด (Measurement) ได้ผลเป็นตัวเลขแสดงจำนวน เช่น การวัดความสูง การชั่ง การตวง การนับจำนวน เป็นต้น ในด้านของการพิจารณาลักษณะของปริมาณนี้ตัวเลขที่ปรากฏย่อมจะมีการผันแปร (variation) ไปได้ เป็นต้นว่าถ้าเราต้องการ

จะวัดความสูงของคนชาวโพทพันธุ์หนึ่งโดยวัด 10 คน จะเห็นได้ว่าความสูงของคนชาวโพทแต่ละคนใน 10 คนนั้น มีการผันแปรไปต่าง ๆ กัน เมื่อพิจารณาแล้วเราสามารถที่จะแบ่งแยกตัวเลขโดยอาศัยลักษณะการผันแปร เป็นเกณฑ์ออกได้เป็นสองลักษณะ คือ

ตัวเลขที่มีการผันแปรต่อเนื่องกัน (Continuous data) ลักษณะหนึ่ง ตัวเลขประเภทนี้อาจมีการผันแปรได้ใกล้เคียงกันมากจนกระทั่งสามารถนึกเพี้ยนกับจุดทศนิยมหลาย ๆ ตำแหน่ง ซึ่งเราไม่สามารถจะตัดสินโดยละเอียดถึงกว่าตัวเลขที่ได้นั้นอยู่ที่สุดใดเป็นคน ตัวอย่างเช่น ตัวเลขที่ได้จากการวัดความยาว ความสูง หรือตัวเลขที่ได้จากการชั่ง เราอาจได้ทศนิยมหลาย ๆ ตำแหน่ง ยิ่งการชั่งละเอียดละออมากเท่าใดก็ยิ่งได้ทศนิยมมากขึ้นเท่านั้น อีกลักษณะหนึ่งก็คือ ตัวเลขที่มีการผันแปรไม่ต่อเนื่องกัน (Discrete data) ได้แก่ตัวเลขที่ได้จากการนับจำนวนเป็นคน สมมุติว่าเรานับจำนวนคนได้ 10 คน จะนับละเอียดสักเท่าใดก็คงได้ 10 คนอยู่นั่นเอง

MEAN

แบ่งออกเป็น 3 ประเภทด้วยกันคือ Arithmetic Mean, Harmonic Mean, Geometric Mean สำหรับงานในด้านการวิจัยทางการเกษตรโดยทั่ว ๆ ไปแล้วเกี่ยวข้องกับ Arithmetic Mean ทั้งสิ้น ส่วนเรื่อง Harmonic Mean เป็นเรื่องที่เกี่ยวข้องกับการเฉลี่ยอัตราส่วน หรือความสัมพันธ์ของซึ่งกันและกัน ส่วน Geometric Mean ก็เป็นเรื่องที่เกี่ยวข้องกับการเฉลี่ยสิ่งๆ ที่เพิ่มจำนวนทบถมมากขึ้นเช่นการเพิ่มจำนวนของจุลินทรีย์ หรือการคิดดอกเบี้ยทบต้น แต่ในเรื่องของ Arithmetic Mean นับเป็นสภาพปกติ ตัวอย่างเช่นในปี ค.ศ. 1936 องค์มนตรีทางสาธารณสุขของสมาคมเภสัชกรรมอเมริกันได้ทำการวิเคราะห์ไวตามินจากน้ำมะเขือเทศกระป๋องตราต่าง ๆ ซึ่งทางสมาคม ฯ รับรองและได้ทำการหาผลเฉลี่ยแบบธรรมดา โดยใช้ Arithmetic Mean วิธีหาคือให้ผลรวมซึ่งได้จากจำนวนไวตามิน ของน้ำมะเขือเทศ 17 อย่าง = 333 เป็นตัวตั้งแล้วหารด้วยจำนวนตัวเลขที่เอามารวมกันคือ 17 ผลที่ได้รับคือ $\frac{333}{17} = 19.6$ ซึ่งเป็น Arithmetic Mean

หรือเฉลี่ยธรรมดา Arithmetic Mean นี้ยังอาจแบ่งแยกออกได้เป็น
สองลักษณะคือ แบบสามัญ (Simple) กับแบบถ่วงน้ำหนัก (Weighted)
สำหรับแบบที่เราจำเป็นต้องเกี่ยวของด้วยบ่อย ๆ ก็คือแบบสามัญ ส่วนแบบถ่วงน้ำหนัก
นั้นมักใช้กับตัวจำนวนมาก ๆ เช่นการกระจายความถี่ ซึ่งมีการจัดแบ่งเป็นพวกเป็นชั้น
ชั้นใดมีมากก็มีน้ำหนักมาก ชั้นใดมีน้อยก็มีน้ำหนักน้อย แต่แบบสามัญนั้นเป็นแบบเฉลี่ยธรรมดา
ซึ่งเรานิยมปฏิบัติกันอยู่ทั่ว ๆ ไป

โดยเหตุที่ในการวิจัยทางการเกษตรทั่ว ๆ ไปใช้ Arithmetic Mean
เป็นพื้น ฉะนั้นจึงขอถือโอกาสกล่าวสั้น ๆ ว่า " Mean " หรือตัวย่อว่า M
ก็ให้เป็นที่เข้าใจว่า Arithmetic Mean

Mean ใต้เป็นเครื่องกลวัดหรือแสดงแทนลักษณะหรือปรากฏการณ์ซึ่ง
เกิดขึ้นซ้ำ ๆ กันหลายหน โดยเราไม่สามารถจะกล่าวถึงแต่ละครั้งได้ทุกครั้งไป
นอกจากนั้นยังสามารถสรุปผลของปรากฏการณ์นั้น ๆ ให้อยู่ในความและสามารถนำไป
ใช้กล่าวอธิบายแทนปรากฏการณ์นั้น ๆ ได้ ตัวอย่างเช่นในการวัดความสูงของ
นักศึกษาในชั้นซึ่งมีความสูงจำนวน 10 คน จะได้ตัวเลขแสดงความสูงของแต่ละคน
ดังต่อไปนี้.-

ตารางที่ 1

คนที่	Variable X ความสูง (ซ.ม.)
1	170
2	169
3	164
4	171
5	165
6	159
7	167
8	166
9	175
10	174
Total	1680

$$\text{Mean } (\bar{M}) = \frac{\text{Sum of } X}{n} = \frac{1680}{10} = 168.0$$

Σ = Sum หรือผลรวมของ.....

n = จำนวนตัวเลข

จะเห็นได้ว่าความสูงของแต่ละคนไม่เหมือนกัน มีการผันแปรเปลี่ยนแปลงไปต่าง ๆ กัน เราจึงเรียกตัวเลขแต่ละตัวที่มีการเปลี่ยนแปลงได้นี้ว่า variable ส่วนการที่จะกลาวว่านักศึกษาในชั้นนี้มีความสูงเท่าใดนั้น ถ้าไม่มี Mean

ก็ยอมจะกระทำไม่ได้ ถึงหากจะนำตัวเลขทั้ง 10 ตัวมากล่าวทุก ๆ ตัว ก็จะสามารถจับความสำคัญไม่ไคว่ภายในตัวเลข 10 ตัวนี้ เราสามารถจะเชื่อถือตัวไหนได้แน่ ฉะนั้นในการหา จึงเท่ากับหาตัวแทนของเลขชุดนี้มากล่าว โดยถือเอาจุดกึ่งกลางระหว่างระยะที่ตัวเลขนี้เปลี่ยนแปลงไป ๆ มา ๆ เพื่อใช้เป็นตัวแทนของเลขชุดนี้

- การหา Mean นี้ก็โดยเอาเลขทั้งหมด 10 จำนวนมารวมกันแล้วหารด้วยจำนวนเลขทั้งหมดซึ่งในที่นี้ ผลรวมได้เท่ากับ 1680 เมื่อหารด้วยจำนวนตัวเลข 10 จะได้ผลลัพธ์เป็นผลเฉลี่ยคือ 168.0 โดยใช้ค่าความสูงเฉลี่ยของนักศึกษาทั้ง 10 คน ซึ่งเท่ากับ 168 ซม. เป็นตัวแทน เราก็สามารถกล่าวสรุปได้ว่า ความสูงเฉลี่ยของนักศึกษาในชั้นนี้เท่ากับ 168 ซม. แต่นี่เป็นเพียงตัวแทนที่ใช้สำหรับกล่าวสรุปว่า ระดับความสูงปานกลางของนักศึกษาในชั้นนี้ คือ 168 ซม. แต่เราไม่สามารถจะทราบได้ว่านักศึกษาแต่ละคนในชั้นนี้จะสูงต่ำเท่าใด ฉะนั้นจึงนับว่า Mean ยังเป็นเครื่องอธิบายที่ไม่ละเอียดเพียงพอเพราะเป็นเพียงตัวแทนของตัวเลขในชุดทั้งหมดเท่านั้น ยังไม่สามารถจะเป็นตัวแทนหรืออธิบายตัวเลขแต่ละตัวได้ เมื่อพิจารณาถึงเรื่องนี้จึงมีหลักง่าย ๆ อยู่ว่า Mean ที่เป็นตัวแทนที่ดีในเลขชุดหนึ่งถึงตัวอย่างต่อไปนี้.-

1	2
24.0	29.5
24.5	18.0
26.0	32.5
24.0	36.0
26.5	19.0
ΣX 125.0	ΣX 125.0
$M_1 = \frac{125.0}{5} = 25.0$	$M_2 = \frac{125.0}{5} = 25.0$

ผลเฉลี่ยของเลขชุดที่ 1 เท่ากับผลเฉลี่ยของเลขชุดที่ 2

จะเห็นได้ว่าตัวเลขทั้งสองชุดนี้ได้ผลเฉลี่ย 25.0 เท่า ๆ กัน แต่ในชุดที่ 1 นั้น สามารถใช้แทน X แต่ละตัวได้ดีกว่าชุดที่ 2 เพราะ Mean มีความใกล้เคียงความจริงในชุดที่ 1 มากกว่าชุดที่ 2 เพราะฉะนั้นจึงจำเป็นต้องหาวิธีว่า ทำอย่างไรจึงจะหาทางที่จะใช้ Mean เป็นเครื่องมือพยากรณ์การเปลี่ยนแปลงของตัวเลขในชุดที่ Mean เป็นตัวแทนอยู่ มิฉะนั้นเราจะไม่มีทางทราบได้ว่า Mean นั้นจะเป็นที่เชื่อถือได้หรือไม่ดังตัวอย่างที่แลวมานี้จะเห็นได้ว่า ตัวเลขทั้งสองชุดนี้ค่าของ Mean = 25.0 เท่า ๆ กัน แต่ถ้าจะพิจารณาคู ตัวเลขทีละจำนวนจะเห็นว่าตัวเลขในชุดที่ 2 ห่างไกล Mean มากกว่า ฉะนั้น Mean ในชุดที่สองจึงยังไม่น่าเชื่อถือ แต่ถ้าหากมีผู้นำ Mean มากล่าวแต่ อย่างเดียว เราก็ไม่มีทางที่จะทราบความจริงได้

การกระจายความถี่

FREQUENCY DISTRIBUTION

ในการวัดซึ่งตรงหรือตรวจสอบลักษณะปรากฏการใด ๆ ก็ตาม เราจำเป็นต้องมีการบันทึกสถิติตัวเลขเป็นรายย่อย เช่นการวัดความสูงของนักเรียนในชั้นซึ่งมีจำนวน 120 คน เราจำเป็นต้องวัดและจดบันทึกตัวเลขแสดงความสูงของนักเรียนแต่ละคนไว้ทุกคนไป แต่เราเอาความสูงของนักเรียนทั้ง 120 คนมาดู เราจะไม่ได้รับประโยชน์มาจากตัวเลขทั้ง 120 จำนวนนั้นแต่ประการใดเลย เพราะเพียงแต่เห็นตัวเลขมากมายเต็มแผ่นกระดาษ แต่ไม่สามารถจะรู้ถึงเนื้อหาสำคัญที่ตัวเลขเหล่านั้นจะแสดงถึงอะไรออกมาให้เป็นประโยชน์ จึงจำเป็นต้องหาทางจัดระเบียบตัวเลขเหล่านั้นเพื่อจะให้อ่านเข้าใจ และเห็นประโยชน์ จึงจำเป็นต้องหาทางจัดระเบียบตัวเลขเหล่านั้นเพื่อจะให้อ่านเข้าใจและเห็นประโยชน์ได้มากขึ้น

จะขอยกตัวอย่างเช่น ในการตรวจและจดบันทึกของคอกกล้วยไม้ชนิดหนึ่ง ซึ่งเป็นกรตรวจทางคุณภาพ นักคนคว้าทำการคนคว้าเรื่องนี้จะจดบันทึกการสำรวจของกลีบดอกกลีบเป็นบัญชีรายการยืดยาวเป็นรายคนไป การดำเนินการขั้นแรกก็คือตรวจความผิดเพี้ยนของสีดอกทั้งหมด แล้วจึงจัดจำแนกออกไปเป็นหมวด ๆ แต่ละหมวดก็ระบุลักษณะของสีไว้ เสร็จแล้วคววากกล้วยไม้ที่สำรวจสีดอกและจดบันทึกลงในบัญชีรายการรวมนั้น คนใดที่ดอกมีสีจัดเข้าหมวดใดก็จัดเข้าในหมวดนั้น เสร็จแล้วจึงนับจำนวนกล้วยไม้ที่ปรากฏในแต่ละหมวดหรือแต่ละชั้น และจำนวนที่ปรากฏขึ้นหรือแต่ละหมวดนี้เองเราจึงเรียกว่า "ความถี่" หรือ frequency ลักษณะหรืออาการที่ความถี่ปรากฏกระจายออกไปตามหมวดหรือตามชั้นต่าง ๆ เราเรียกว่า "การกระจายความถี่" หรือ Frequency Distribution ของลักษณะที่ทำการศึกษา ดังตัวอย่าง.

ตารางที่ 2

การกระจายของลักษณะสีกลีบดอกกล้วยไม้สกุลหวายลูกผสมชนิดหนึ่ง

	Class	Frequency
1	เหลือง	162
2	เหลืองขลิบแดง	224
3	เหลืองขลิบแดงและเส้นลายสี แดง	281
4	สีทองแดง	250
5	สีแดง	198
	รวม	1115

ในที่นี้จำนวนกล้วยไม้ลูกผสมทั้งหมดที่ทำการสำรวจมีอยู่ 1115 ต้น มีลักษณะสีดอกที่ปรากฏแบ่งออกได้เป็น 5 หมวด ซึ่งแสดงลักษณะประจำชั้นในของแรกของการวางที่ 2 ส่วนของที่ 2 เป็นการกระจายลักษณะสีดอกของกล้วยไม้ 1115 ต้น ออกไปตามชั้นต่าง ๆ ทั้ง 5 ชั้น ดังนั้น ตัวเลขในช่องที่สองจึงเป็นความถี่ปรากฏขึ้นในแต่ละชั้นนั่นเอง

ที่ไคกลาวและยกตัวอย่างมาแลวนั้นเป็นการศึกษาทางคานคุณภาพ ต่อไปนี้จะขอกล่าวถึงการศึกษาในคานปริมาณ คือลักษณะที่สามารถวัดหรือนับเป็นหน่วยมาตราได้ เช่น ความยาว ความสูง น้ำหนัก ความจุ ฯลฯ การตรวจทางปริมาณนี้โดยทั่ว ๆ ไป กระทำกันเป็นจำนวนมาก ๆ ดังนั้นนอกจากการใช้ Frequency distribution ช่วยให้ตัวเลขเหล่านั้นแสดงผลให้เห็นได้ชัดแล้ว ยังช่วยย่อให้บันทึกทางสถิติมีขนาดเล็กลง สะดวกแก่การใช้ประโยชน์และการเก็บ ดังตัวอย่างต่อไปนี้.

ตารางที่ 3

ตัวเลขได้จากการวัดความยาวของผักขาวโปก

ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต
1	10.0	101	10.4	201	9.7	301	8.7
2	10.7	102	9.7	202	11.7	302	10.2
3	9.8	103	9.7	203	8.9	303	12.0
4	10.6	104	10.4	204	11.6	304	10.4
5	11.2	105	11.0	205	11.7	305	12.0
6	9.2	106	11.2	206	9.2	306	11.9
7	9.2	107	11.5	207	10.2	307	8.5
8	11.2	108	10.9	208	10.6	308	6.4
9	8.5	109	13.6	209	10.4	309	11.4
10	9.7	110	11.0	210	11.3	310	9.8
11	12.5	111	9.5	211	10.6	311	10.2
12	8.3	112	11.1	212	10.3	312	10.9
13	9.8	113	10.5	213	10.7	313	10.1
14	10.4	114	9.8	214	11.8	314	12.1
15	10.8	115	12.2	215	11.8	315	9.5
16	10.0	116	8.5	216	9.8	316	10.4
17	9.6	117	8.4	217	9.9	317	9.9
18	8.8	118	9.3	218	9.6	318	9.2
19	11.8	119	9.1	219	9.6	319	11.3
20	9.5	120	10.6	220	9.6	320	8.8
21	9.5	121	11.5	221	13.7	321	8.4
22	9.4	122	10.5	222	9.9	322	8.7
23	10.0	123	10.4	223	5.7	323	9.9

ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต
24	10.3	124	10.4	224	9.2	324	12.1
25	7.5	125	9.4	225	9.4	325	10.5
26	8.4	126	9.0	226	10.7	326	7.0
27	11.6	127	9.4	227	8.8	327	11.5
28	11.1	128	11.3	228	8.8	328	9.7
29	9.5	129	11.8	229	11.1	329	8.2
30	11.3	130	9.6	230	10.4	330	9.6
31	8.9	131	8.4	231	11.5	331	11.6
32	9.4	132	8.7	232	11.9	332	8.4
33	10.2	133	8.5	233	10.5	333	9.5
34	11.4	134	11.4	234	7.1	334	8.6
35	9.9	135	10.8	235	10.8	335	12.9
36	10.2	136	10.6	236	9.5	336	9.2
37	10.1	137	10.5	237	8.6	337	10.2
38	10.4	138	9.9	238	12.2	338	10.3
39	9.9	139	6.7	239	11.6	339	12.4
40	9.4	140	9.5	240	10.7	340	9.4
41	8.9	141	7.0	241	9.0	341	8.9
42	12.0	142	10.6	242	8.6	342	9.9
43	13.0	143	11.2	243	10.1	343	9.8
44	9.7	144	11.0	244	11.2	344	6.9
45	8.9	145	8.3	245	9.3	345	11.8
46	7.9	146	8.4	246	11.2	346	13.1
47	8.9	147	7.0	247	7.9	347	9.4
48	9.5	148	9.6	248	9.0	348	10.5
49	11.0	149	11.7	249	9.4	349	13.6
50	9.8	150	8.9	250	7.4	350	8.7

ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต
51	10.2	151	10.9	251	10.0	351	10.2
52	8.9	152	7.8	252	8.4	352	12.1
53	9.8	153	10.0	253	10.6	353	6.2
54	8.4	154	11.9	254	10.6	354	10.8
55	8.8	155	8.7	255	9.2	355	10.9
56	9.0	156	9.4	256	10.5	356	11.0
57	9.5	157	11.7	257	9.4	357	10.5
58	13.4	158	13.1	258	10.0	358	10.0
59	10.5	159	9.6	259	9.3	359	11.5
60	7.7	160	10.2	260	11.2	360	9.2
61	10.1	161	10.4	261	8.7	361	8.4
62	8.2	162	9.5	262	9.6	362	11.8
63	11.4	163	12.0	263	9.2	363	9.7
64	7.5	164	10.6	264	9.3	364	9.0
65	9.6	165	8.7	265	9.5	365	6.3
66	6.6	166	10.9	266	9.8	366	10.3
67	9.9	167	12.2	267	10.0	367	8.6
68	12.7	168	8.0	268	11.2	368	10.2
69	10.2	169	12.8	269	11.0	369	8.0
70	9.6	170	9.7	270	11.6	370	10.0
71	10.0	171	6.5	271	10.0	371	10.4
72	11.7	172	11.5	272	13.2	372	11.4
73	8.7	173	9.8	273	7.6	373	8.9
74	8.4	174	10.0	274	10.1	374	11.5
75	10.0	175	9.8	275	11.4	375	10.8

ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต	ลำดับที่	ความยาวของ ผัก/นิ้วฟุต
76	8.3	176	9.4	276	12.0	376	10.5
77	7.9	177	6.3	277	10.0	377	11.8
78	10.0	178	10.9	278	10.3	378	10.4
79	10.9	179	7.6	279	10.5	379	10.4
80	9.7	180	7.6	280	10.9	380	9.5
81	12.5	181	10.1	281	10.8	381	11.8
82	10.5	182	8.9	282	9.0	382	9.8
83	11.3	183	10.9	283	10.4	383	8.5
84	10.8	184	8.9	284	12.3	384	10.7
85	10.7	185	9.4	285	9.7	385	10.0
86	10.5	186	5.5	286	9.1	386	10.7
87	9.5	187	6.3	287	10.0	387	9.2
88	9.1	188	10.1	288	8.8	388	10.8
89	9.1	189	8.6	289	9.8	389	10.3
90	12.1	190	10.6	290	9.7	390	8.2
91	12.5	191	8.8	291	5.4	391	9.5
92	9.0	192	9.3	292	10.1	392	6.8
93	11.2	193	13.3	293	8.3	393	10.6
94	9.7	194	11.5	294	9.0	394	10.0
95	11.0	195	11.3	295	10.7	395	12.4
96	11.1	196	10.0	296	10.0	396	8.1
97	9.7	197	10.7	297	9.3	397	7.5
98	8.9	198	9.8	298	13.5	398	10.3
99	10.1	199	6.7	299	11.0	399	8.8
100	11.4	200	8.5	300	10.2	400	10.2

จากตัวอย่างนี้ ค้นแรกเราจะทองหาระยะ (Range) ระหว่าง
 ตัวเลขน้อยที่สุดกับมากที่สุด เพื่อจะได้จัดแบ่งเป็นชั้นใดเหมาะสม เราจะเห็นได้ว่า
 ในบรรดาชาวโศก 400 ผักที่ทำการวัดความยาวนั้น ผักที่สั้นที่สุดวัดได้ 5.4 นิ้วฟุต
 และผักที่ยาวที่สุดวัดได้ 13.7 นิ้วฟุต เราจึงมีระยะความแปรปรวนทั้งหมดตั้งแต่
 5.3 ถึง 13.7 นิ้วฟุต ถ้าเราจะจัดเป็นชั้น ๆ โดยมีระยะระหว่างชั้น
 (Class Interval) เป็น 0.5 นิ้วฟุต เราก็แบ่งชาวโศกโดยอาศัยลักษณะ
 ความยาวออกได้เป็น 17 ชั้น ชั้นแรกเริ่มควย 5.3 ถึง 5.7 ชั้นที่ 2 ตั้งแต่ 5.8
 ถึง 6.3 ควรจำไว้ว่าตัวเลขติดต่อกันระหว่างชั้นต้องมีตัวเลขเดียวกัน มิฉะนั้นถ้ายัง
 เอื้อยเลขจำนวนใดมีค่าตกอยู่ระหว่างชั้นจะจัดเข้าชั้นไม่ได้ ดังนั้นค่าของตัวเลขชั้นตอน
 จึงต้องเป็นตัวเลขที่มีค่าติดกัน 1 หน่วยและต่อเนื่องกัน เช่น 5.7 เป็นค่าสูงสุดของชั้นที่ 1
 และ 5.8 ก็เป็นค่าต่ำสุดของชั้นที่ 2 รับต่อกัน เมื่อจัดค่าของชั้นแล้วก็จำเป็นต้องหา
 ค่ากึ่งกลางของแต่ละชั้น (Class value) หรือ เช่น ชั้นที่ 1 ค่ากึ่งกลางระหว่าง
 5.3 ถึง 5.7 ก็คือ 5.5 ทั้งนี้เราจะถือว่าเลขที่ตกอยู่ในขอบเขตของชั้นที่ 1
 จะมีค่าเป็น 5.5 หกค ซึ่งที่แท้จริงอาจเป็น 5.3 5.4 5.5 5.6 หรือ 5.7 ก็ได้
 ควยเหตุนี้เองเราจึงสามารถกล่าวได้ว่า ถ้าจัดขอบเขตของชั้นให้แคบลง ผลก็จะยิ่ง
 ใกล้ความจริงยิ่งขึ้น และในทำนองตรงกันข้าม ถ้าจัดขอบเขตของแต่ละชั้นให้กว้างออก
 ผลก็จะเคลื่อนไปมากยิ่งขึ้น การคำนวณในขั้นต่อไปก็คือการหาความถี่ของแต่ละชั้น
 โดยการนับดูว่าตัวเลขที่ตกอยู่ในขอบเขตของแต่ละชั้นมีจำนวนชั้นละเท่าใด แล้วจึง
 กรอกลงในรายการประจำชั้นนั้น ๆ ความถี่ที่เราเรียกว่า FREQUENCY หรือใช้
 ตัวย่อว่า f ต่อไปจึงหาผลคูณระหว่างค่ากึ่งกลางของชั้น กับความถี่ $f \times V$
 เนื่องจากเราถือเสมือนว่าตัวเลขที่ตกอยู่ในชั้นนั้นมีค่าเท่ากับค่ากึ่งกลางของชั้น
 ซึ่งใช้แทนค่าของตัวเลขทุกจำนวนในชั้นนั้น ดังนั้นผลคูณระหว่าง f กับ V
 จึงเท่ากับผลรวมของตัวเลขทุกจำนวนในชั้นนั้น ๆ นั่นเอง เช่น ในชั้นที่ 1 มีตัวเลข
 ตกอยู่ภายในขอบเขตของชั้น 3 จำนวน ก็คือ $5.5 \times 3 = 16.5$ เท่ากับเลข 3
 จำนวนรวมกันโดยรวมเท่ากับ 16.5 นั่นเอง ดังนั้นเมื่อรวมผลคูณระหว่าง f กับ V
 ทุกจำนวนเข้าด้วยกันจึงเป็นผลรวมใหญ่ (Grand Total) ของตัวเลขทั้ง
 400 จำนวน หรือ

$$\sum_{400} x = fV$$

หรือ

$$\sum x_n = fV$$

จกตารางที่ 4 เราจะสังเกตเห็นได้ว่า Frequency คานบนและคานล่าง มีน้อย และค่อย ๆ เพิ่มขึ้นในตอนกลาง ซึ่งเป็นไปตามกฎเกณฑ์ของโดลกตามธรรมชาติ หรืออาจกล่าวได้ว่า สิ่งปรากฏตามธรรมชาตินั้น ลักษณะที่อยู่ในระดับกลาง ๆ หรือระดับปกติ จะมีจำนวนมาก ดังเช่นความยาวของผักขาวโพด ความสูงของ มนุษย์ หรือสัตว์ หรือต้นไม้ น้ำหนัก ฯลฯ ผักที่ยาวมากและผักที่สั้นมากผิดปกติ เหาไคก็ยังมีจำนวนน้อยลงเท่านั้น

ตารางที่ 4

การกระจายความถี่ของความยาวของผักขาวโพด

Class	Class value (V)	Frequency (f)	Frequency X Class value (f.V)
5.3 - 5.7	5.5	3	16.5
5.8 - 6.2	6.0	1	6.0
6.3 - 6.7	6.5	8	52.0
6.8 - 7.2	7.0	6	42.0
7.3 - 7.7	7.5	8	60.0
7.8 - 8.2	8.0	11	88.0
8.3 - 8.7	8.5	32	272.0
8.8 - 9.2	9.0	42	378.0
9.3 - 9.7	9.5	58	551.0
9.8 - 10.2	10.0	65	650.0
10.3 - 10.7	10.5	55	577.5
10.8 - 11.2	11.0	37	407.0

Class	Class value (V)	Frequency (f)	Frequency X Class value (f.V)
11.3 - 11.7	11.5	31	356.5
11.8 - 12.2	12.0	24	288.0
12.3 - 12.7	12.5	7	87.5
12.8 - 13.2	13.0	6	78.0
13.3 - 13.7	13.5	6	81.0
Total		$f = 400$	3991.0

ในการพิจารณาความกว้างความแคบของชั้น (Class Interval)
มีหลักเกณฑ์ที่จะใช้พิจารณาอยู่หลายประการด้วยกัน คือ

1. ความกว้างของชั้นแต่ละชั้นจะต้องเป็นไปโดยสม่ำเสมอเท่า ๆ กัน
ดังตัวอย่างคือชั้นละ 0.5 ทุกชั้น และระยะความกว้างของชั้นนั้นจะต้องไม่กว้างหรือ
แคบเกินไป จะต้องพึงให้พอเหมาะแก่ระยะความแตกต่างระหว่างตัวเลขจำนวน
ต่อจำนวน ถ้ากว้างเกินไป ตัวเลขหลายจำนวนซึ่งมีค่าแตกต่างกันมาก ๆ ก็จะตกอยู่
ในชั้นเดียวกัน โดยมีค่ากึ่งกลางชั้นเป็นตัวแทน ทำให้คลาดเคลื่อนไปจากค่าที่แท้จริง
ได้มาก ผลงานก็จะหยاب แต่ก็ไม่ควรให้เล็กจนเกินไป เพราะจะทำให้มีมากขึ้นเกิน
ความจำเป็น ความถี่ของแต่ละชั้นก็จะน้อย บางชั้นอาจมีค่าเป็น 0 เลยก็มี
ถ้าขอบเขตของแต่ละชั้นแคบมาก จำนวนชั้นก็จะมีมากขึ้นและค่า frequency ของ
แต่ละชั้นก็จะลดน้อยลง แต่ถ้าขอบเขตของแต่ละชั้นกว้างออก จำนวนชั้นก็จะมีน้อยลง
และค่า frequency ของแต่ละชั้นก็จะเพิ่มขึ้น อย่างไรก็ตามผลรวมของ frequency
ทั้งหมดจะต้องเท่ากับจำนวนตัวเลขทั้งหมดที่นำมาตรวจวัดจึงจะถูกต้อง

(ตามตัวอย่างคือ 400)

2. ระยะเวลาของชั้นทั้งหมด คือ ค่าตัวเลขน้อยที่สุดในชั้นต่ำถึงค่าสูงที่สุดในชั้นสูง (ตามตัวอย่าง 5.3 ถึง 13.7) จะต้องคลุมเอาตัวเลขที่ทำการตรวจวัดใดทั้งหมด และค่าของชั้นทุกชั้นจะต้องต่อเนื่องกัน ตามลำดับเรื่อยไปโดยตลอด

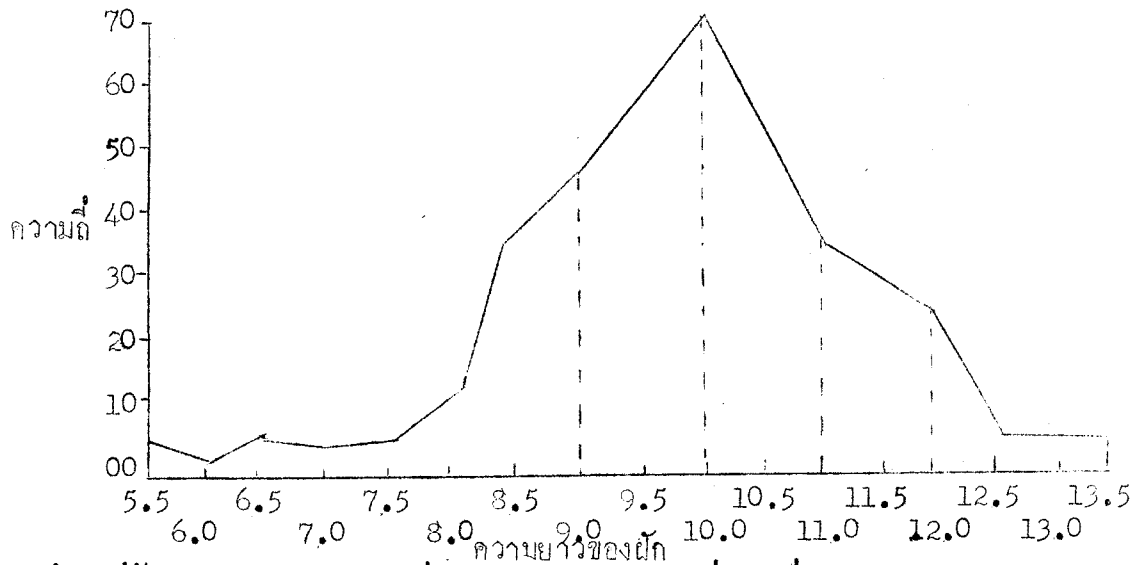
3. โดยทั่ว ๆ ไปจำนวนชั้นควรจะอยู่ในราว 10 ถึง 20 จึงนับว่าเหมาะสมไม่ควรมากเกินไป 30 ชั้น หรือน้อยกว่า 6 ชั้น

4. ควรจะวางความกว้างของชั้นให้สามารถกำหนดค่ากึ่งกลางของชั้นได้สะดวก คือให้ใดจำนวนเป็นสิบ หรือ ห้า หรือจำนวนที่แบ่งครึ่งได้สะดวก ถ้าเป็นเลขคู่ก็หาจุดกึ่งกลางยากกว่าเลขคี่ เพราะเลขคู่หากจะแบ่งครึ่งก็ไม่มีค่าตัวกลาง เนื่องจากสองข้างเท่ากันอยู่แล้ว เช่น 4 เป็นเลขคู่ ค่าตัวกลางมีค้อยู่ที่ 2 แต่อยู่ระหว่าง 2 กับ 3 ถ้าเป็นเลขคี่ เช่น 5 ค่าตัวกลางอยู่ที่ 3 คือ 1 และ 2 อยู่ข้างหนึ่ง และ 4 กับ 5 อยู่อีกข้างหนึ่ง เป็นต้น

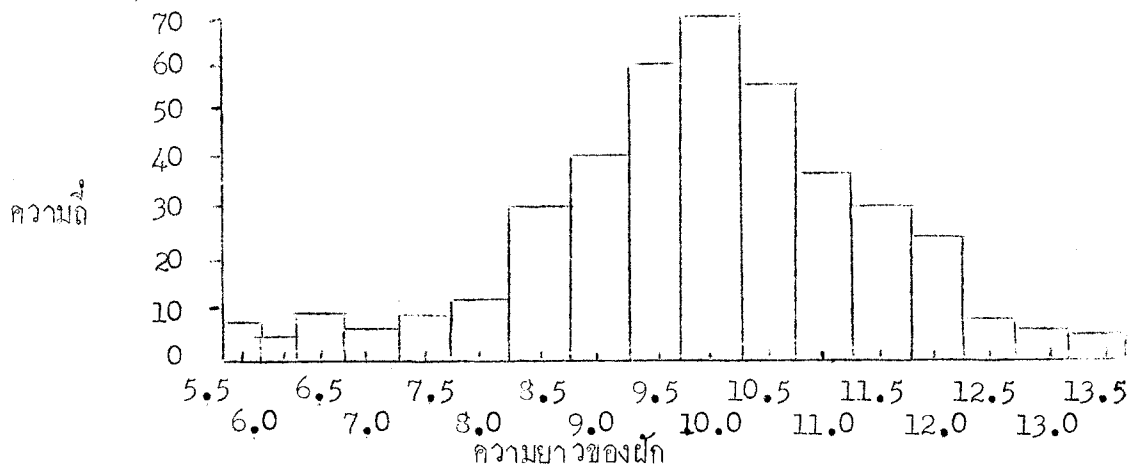
การแสดงควมกราฟ

บางทีการอ่านผลการค้นคว้าจากตัวเลขมาก ๆ ทำให้เกิดความสับสนและยุ่งยาก ผู้พบเห็นก็จะเกิดความสนใจ จึงได้คิดหาทางเปลี่ยนสภาวะทางตัวเลขให้กลายเป็นภาพเพื่อแสดงให้เห็นการเปลี่ยนแปลงได้ง่ายและสะดวกรวดเร็วยิ่งขึ้น โดยการเรียงลำดับความถี่ของแต่ละชั้นมาเปลี่ยนสภาพ การแสดง frequency ด้วยกราฟนั้น กระทำได้โดยลากเส้นแกนต้งและแกนนอนให้ไต่จากกัน สำหรับแกนต้งให้เป็นค่าของ frequency โดยตรวจดูเสียก่อนว่าค่าของ frequency ตั้งแต่ต่ำสุดถึงมากที่สุดคืออะไร ตามตัวอย่างคือ 1 ถึง 65 ดังนั้นก็ตั้งมาตราส่วนให้เป็น 0 ถึง 70 คือให้क्रमระยะของ frequency ว่างพอเหมาะ แต่ควรเป็นจำนวน 5 หรือ 10 ส่วนแกนนอนก็แบ่งตามค่ากึ่งกลางของชั้น เสร็จแล้วจึงหมายจุดโดยใช้ค่ากึ่งกลางของชั้นกับความถี่เป็นคู่ ๆ ไป เช่น ค่าของชั้นเป็น 5.5 ความถี่เป็น 3 เราก็ลากเส้นตั้งฉากกับแกนนอนตรง 5.5 และลากเส้นตั้งฉากกับแกนต้งตรง 3 ตัดกันที่ใดก็เป็นค่าความถี่หมายจุด กระทำเช่นนี้จนครบทุกคู่.

* ต่อไปจะต้องตั้งจุดหมายความว่าการทำกราฟนั้นเพื่อประสงค์อะไร ถ้าประสงค์จะทราบการเปลี่ยนแปลงความถี่จากชั้นหนึ่งไปสู่อีกชั้นหนึ่ง ก็ขีดเส้นโยงระหว่างจุดที่หมายไว้ โดยเรียงลำดับจากชั้นต้นไปชั้นสุดท้าย เราก็จะสามารถเห็นรูปร่างการเปลี่ยนแปลงของ frequency ได้ชัดเจน กราฟที่แสดงเช่นนี้ เราเรียกว่า FREQUENCY POLYGON



* แต่ถาประสงค์จะทราบอัตราส่วนหรือเปรียบเทียบระหว่างชั้นต่อชั้นเป็นสำคัญ การเขียนกราฟจำเป็นจะต้องแสดงให้เห็นสัดส่วนและความแตกต่างระหว่างชั้นต่อชั้นได้ชัดเจน ดังนั้นแทนที่จะลากเส้นโยงระหว่างจุดจึงใช้วิธีลากเส้นตั้งฉากหรือทำเป็นแท่งหนา ๆ เพื่อแสดงความสูงต่ำของ frequency ให้เห็นได้แจ่มชัดโดยอิสระ กราฟเส้นตั้งหรือ กราฟแท่งนี้ เราเรียกว่า HISTOGRAM



ในเรื่องของ Frequency Distribution นั้นการหาห้นยมกระทำกันอยู่ ๓ แบบด้วยกัน คือ การหา Mean, Median และ Mode

MEAN ก็คือผลเฉลี่ยซึ่งหาได้โดยวิธีการคำนวณทางคณิตศาสตร์ธรรมดาที่เรานิยมใช้กันทั่ว ๆ ไป คือเอาค่าของตัวเลขทุกจำนวนรวมกันหมดแล้วก็หารด้วยจำนวนตัวเลขที่เอามา รวมกัน แต่ว่า Frequency Distribution มีตัวเลขมากจำนวนด้วยกัน การหา MEAN จึงกระทำยากกว่าปกติ จึงได้อธิบายไว้แล้วถึงวิธีการหาผลรวมของตัวเลขทั้งหมด เราหาผลรวมทีละชั้นเป็นชั้น ๆ ไป โดยเอาค่าของชั้นคูณกับจำนวนตัวเลขในชั้นนั้น แล้วจึงเอาผลคูณแต่ละชั้น มารวมกันหมด ผลรวมทั้งหมดจึงอาจเคลื่อนไปได้อ่างเล็กน้อย เพราะค่าประจำชั้นมิได้ตรงกับ

ความจริงของตัวเลขในชั้นใดก็ทัวตัวเสมอไป แต่เนื่องจากตัวเลขทั้งหมดมีจำนวนมากมาย การที่คิดาค
เคลื่อนไปในทิศทางน้อยจึงถือว่าไม่สำคัญ เสร็จแล้วจึงเอาจำนวนตัวเลขทั้งหมดไปหารผลรวมของ
ค่าตัวเลขทั้งหมด

f = frequency

V =Class value

n = จำนวนตัวเลขทั้งหมด

ดึงขึ้นจากตารางที่ 3

$$M = \frac{3991.0}{400} = 9.9775$$

DEVIATION

เมื่อใดกล่าวมาแล้วในตอนท้ายเรื่องของ Mean ว่าเราจำเป็นต้องหาทางที่จะวัดการเปลี่ยนแปลงของตัวเลขแต่ละชุด เพื่อต้องการทราบว่าตัวเลขแต่ละจำนวนในชุดนั้นมีความสม่ำเสมอให้เชื่อถือได้หรือมีการแปรปรวนมากน้อยอย่างไร เราสามารถจะกระทำได้ โดยอาศัยจุดกึ่งกลาง คือ Mean เป็นหลักเสียก่อนแล้วจึงวัดความคลาดเคลื่อนของตัวเลขแต่ละจำนวนที่เคลื่อนไปจาก Mean ระยะที่ตัวเลขแต่ละจำนวนเคลื่อนไปจาก Mean นี้ เราเรียกว่า "Deviation" รวม, ๓๐๕๖

เนื่องจาก Mean เป็นจุดกึ่งกลางนั่นเอง ค่าของ Deviation จึงมีทั้งทาง "บวก" และทาง "ลบ" เพราะความคลาดเคลื่อนจะต้องเป็นไปทั้งสองข้าง และก็โดยเหตุผลที่ว่า Mean เป็นจุดกึ่งกลาง ฉะนั้นแหละทำให้ค่าของ Deviation รวมกันแล้วต้องเท่ากับ "๐" เสมอ เพราะผลรวมทั้งทางบวกและทางลบต้องได้สมดุลย์เท่ากันทั้งสองข้าง ถ้าผลรวมของ Deviation ไม่ได้ "๐" ก็ย่อมแสดงว่า Deviation ทั้งสองข้างไม่เท่ากัน ค่าของ Mean อาจผิดก็ได้เพราะไม่ได้จุดกึ่งกลาง หรือมีฉะนั้นการคำนวณตอนหนึ่งตอนใดอาจผิด ยกเว้นเสียจากในกรณีที่มีการคำนวณหา Mean จำเป็นต้องปัดเศษ หรือมีทศนิยมไม่ครบ ค่าของ Mean อาจเคลื่อนไบบ้างเล็กน้อยแล้วแต่ความละเอียดละออของงานที่ปฏิบัติพึงประสงค์ ตามที่ได้ทราบหลักเกณฑ์แล้วว่า ความคลาดเคลื่อน หรือ Deviation นั้นก็คือ ระยะที่ตัวเลขแต่ละจำนวนคลาดเคลื่อนไปจาก Mean ฉะนั้น ในการคำนวณหา Deviation เราก็จำเป็นต้องเอาค่าของ Mean ไปลบค่าของตัวเลขแต่ละจำนวน ดังตัวอย่างต่อไปนี้

ตารางที่ ๓
(ตัวเลขจากตารางที่ ๑)

ลำดับที่	variable "X"		deviation (d)		
			X - M	ผลลัพท์	
				+	-
1	X ₁	170	170 - 168	2	
2	X ₂	169	169 - 168	1	
3	X ₃	164	164 - 168		4
4	...	171	171 - 168	3	
5	...	165	165 - 168		3
6	...	159	159 - 168		9
7	...	167	167 - 168		1
8	...	166	166 - 168		2
9	...	175	175 - 168	7	
10	X _n	174	174 - 168	6	
รวม	1680			+ 19	- 19
				Total = 0	

$$M = (X_1 + X_2 + X_3 + X_4 + \dots + X_n) / n = \frac{X}{n} = \frac{1680}{10} = 168.0$$

แผนผังที่ ๑ แผนผังแสดงที่ตั้งของผลเฉลี่ยและการกระจายของเลขบริวาร

+7	+6	+5	+4	+3	+2	+1	0	-1	-2	-3	-4	-5	-6	-7	-8	-9
175	174	173	172	171	170	169	168	167	166	165	164	163	162	161	160	159

เมื่อพิจารณาจากแผนผังนี้ประกอบกับที่ได้จำกัความว่า Mean เป็นจุดกึ่งกลางนั้น จะทำให้เห็นได้ชัดเจนขึ้นว่า Mean มีใช้จุดกึ่งกลางระหว่างค่าของตัวเลขที่มากที่สุดกับตัวเลขที่น้อยที่สุด หากแต่เป็นจุดกึ่งกลางระหว่างผลรวมของระยะที่ตัวเลขแต่ละจำนวนคลาดเคลื่อนไปจาก Mean จะเห็นว่าตัวเลขที่น้อยที่สุด คือ ๑๕๕ ต่ำกว่า Mean ถึง -๕ แต่ตัวเลขที่มากที่สุด คือ ๑๗๕ สูงกว่า Mean + ๗ เท่านั้น แต่เมื่อนำค่าของระยะความคลาดเคลื่อนของทุก ๆ จำนวน ผสมกัน ทางบวกก็จะได้ ๑๕ และทางลบก็จะได้ -๑๕ ซึ่งแสดงว่าสองข้างเท่ากัน Mean จึงเป็นจุดกึ่งกลางของระยะความคลาดเคลื่อนทั้งหมดรวมกัน

STANDARD DEVIATION

ในการที่เราวัดความคลาดเคลื่อนของตัวเลขแต่ละจำนวนซึ่งเคลื่อนไปจาก Mean นั้น เมื่อใดความคลาดเคลื่อนเหล่านี้นี้มาแล้ว ก็ยังไม่สามารถที่จะสรุปลงไปได้ว่า Mean นั้นใกล้เคียงความจริงสักเท่าใด หรือเป็นตัวแทนของตัวเลขชุดนั้นได้ดีเพียงไร เนื่องจากตัวเลขแต่ละจำนวนต่างก็มีความคลาดเคลื่อนของตัวเองโดยเฉพาะ เมื่อมีตัวเลขมากจำนวนด้วยกันจึงไม่สามารถจะตัดสินอย่างไรลงไปได้ หนทางที่จะกระทำได้ดีมีอยู่หนทางเดียว คือ ใช้วิธีวางขอบเขตเพื่อให้ระบุลงไปว่าความคลาดเคลื่อนของตัวเลขในชุดนั้น สามารถเป็นไปภายในขอบเขตที่กว้างหรือแคบเพียงใด ส่วนการกำหนดขอบเขตเพื่อแสดงความแปรปรวนของตัวเลขนั้น ก็จำเป็นจะต้องวางไว้เป็นมาตรฐานสากลเพื่อให้เป็นที่ยอมรับและใช้ได้ทั้ง ๆ ไป ขอบเขตที่จะวางให้เป็นมาตรฐานนี้จำเป็นจะต้องใช้กฎเกณฑ์ในการวางอย่างเดียวกันตลอดไป และผลที่ได้รับเราเรียกว่า "มาตรฐานแห่งความคลาดเคลื่อน" หรือ "Standard Deviation (S.D.)" มาตรฐานแห่งความคลาดเคลื่อนนี้เป็นผลที่ได้จากการเฉลี่ยความแปรปรวน (Variation) ซึ่งนักสถิติสากลได้ประชุมตกลงรับรองแล้ว การคำนวณกระทำโดย เมื่อหา Deviation (d) แล้วจากตารางที่ ๓ ให้เอา Deviation แต่ละตัวยกกำลังสองเป็น d^2 แล้วรวมกันทั้งหมดเป็น Sum of Square ($\sum d^2$) แล้วหารเฉลี่ยด้วยจำนวนตัวเลขทั้งหมดคั่นด้วย ๑ หรือใช้สูตรว่า $n - 1$ ซึ่งเราเรียกว่า degree of freedom (D.F.) ผลที่ได้รับคือ $\sum d^2 / (n-1)$ เรียกว่า Mean square หรือ Variance เมื่อเราถอดกำลังสองออก $\sqrt{\frac{\sum d^2}{n-1}}$ ผลที่ได้รับก็เป็น Standard deviation หรืออาจกล่าวได้ว่า Variance หรือ Mean square (M.S.) นั้น ก็คือ Standard deviation ยกกำลังสองนั่นเอง

$$M.S. = S.D.^2 \quad \text{หรือ} \quad S.D. = \sqrt{M.S.}$$

$$\text{และ Variance หรือ } M.S. = \frac{\text{Sum of Square (S.S.)}}{\text{Degree of freedom (D.F.)}}$$

ดังตัวอย่างต่อไปนี้

ตารางที่ ๔
(ตัวเลขจากตารางที่ ๓)

ลำดับ	X	deviation (d) (X - M)	d ²
1	170	2	4
2	169	1	1
3	164	-4	16
4	171	3	9
5	165	-3	9
6	159	-9	81
7	167	-1	1
8	166	-2	4
9	175	7	49
10	174	6	36
	$\Sigma = 1680$	$\Sigma d = 0$	$\Sigma d^2 = 210$

$$\begin{aligned}
 M &= \frac{1680}{10} \\
 &= 168.0 \\
 S.D. &= \sqrt{\frac{\Sigma d^2}{n-1}} \\
 &= \sqrt{\frac{210}{10-1}} \\
 &= \sqrt{23.333} \\
 &= 4.83
 \end{aligned}$$

การที่ค่าของ S.D. ต้องมีเครื่องหมายบวกลบก็ขึ้นอยู่กับว่าค่า S.D. เป็นขอบเขตของความคลาดเคลื่อนเฉลี่ยที่ตัวเลข X แต่ละตัวเคลื่อนไปจาก Mean ก็ย่อมจะต้องมีการเป็นไปได้ทั้งสองทาง คือทางบวกและทางลบ และค่าของ S.D. นี้ใช้เขียนตามหลัง Mean คือ $M \pm S.D. = 168.0 \pm 4.83$ ซึ่งหมายความว่ามีความวุ่นวายตั้งแต่ $168.0 + 4.83$ ถึง $168.0 - 4.83$ หรือจาก ๑๗๒.๘๓ ถึง ๑๖๓.๑๗ ซึ่งขอบเขตของ S.D. นี้จะคลุมเอาค่าของ X ไว้ประมาณสองในสามส่วน อีกหนึ่งในสามส่วนจะตกอยู่นอกขอบเขตเสมอ

ในการที่เราต้องยกกำลังสองของ Deviation นั้น ถ้าหากจะอธิบายกันอย่าง
ธรรมดา ๆ โดยไม่อาศัยหลักทางคณิตศาสตร์เข้ามาเกี่ยวข้องแล้ว ก็นับว่าไม่มีเหตุผลเพียงพอที่
จะตอบปัญหานี้ได้ แต่ถ้าวัดกันโดยคุณของการคำนวณที่ได้รับก็นับว่าได้ผลแน่นอน และถา
ยั้งได้มีการศึกษาวิชาสถิติมากขึ้นก็ดูเหมือนจะทำให้เข้าใจเรื่องนี้ยิ่งขึ้นทุกที ลองสังเกตดูโดย
ละเอียดจะเห็นว่า ค่าของ deviation เอง ที่เรากำหนดหามาได้ทุกตัวนั้นไม่ได้ช่วยให้ Mean
มีความหมายชัดเจนขึ้นแต่อย่างใดเลย แต่กลับพบกับปัญหาสำคัญอันหนึ่ง คือ เมื่อเอาค่าของ
deviation รวมกันเป็นยอดใหญ่ของ deviation จะได้ค่าเท่ากับ ๐ ทุกครั้ง เนื่องจาก
Mean เป็นจุดกึ่งกลาง เมื่อเอาค่าของ deviation ซึ่งเป็นระยะที่ตัวเลขเคลื่อนไปจาก
mean ทั้งทางบวกและทางลบรวมกัน ก็จะหักลบกันหมด ทำให้หมดหนทางที่จะคำนวณต่อไปได้อีก
แต่ถ้าวางยกกำลังสอง deviation แต่ละตัว เครื่องหมายลบก็จะหมดไป เสร็จแล้วนำไปรวม
กันทีหลัง ก็จะได้อายุของ deviation ตามต้องการ หากแต่เป็นค่าของผลรวมของ deviation
ยกกำลังสอง เราจึงจำเป็นต้องหารด้วย $(n-1)$ เพื่อเฉลี่ยผลรวมของกำลังสองนี้ให้เหลือเพียง
หน่วยเดียว เพราะมาตรฐานความคลาดเคลื่อนที่เราได้มานั้น เรานำไปใช้แสดงความคลาด
เคลื่อนของตัวเลขแต่ละตัวจากผลเฉลี่ย ฉะนั้น ในการคำนวณเราจึงจำเป็นต้องเฉลี่ยด้วย $n-1$
เมื่อเฉลี่ยเสร็จแล้วจึงถอดกรณฑ์สอง เพราะเราได้อยกกำลังสองเอาไว้ ผลที่ได้จึงเป็นมาตรฐาน
ความคลาดเคลื่อน (Standard deviation) หรือ Standard error S.D. หรือ S.E.
ใช้แสดงขอบเขตแห่งความคลาดเคลื่อนที่ variable $X(X_i)$ จะเคลื่อนไปจาก mean เพื่อให้
mean มีความหมายชัดเจนยิ่งขึ้นอีก

ในการหารเฉลี่ย deviation นี้เกี่ยวกับการหา mean เพราะการหา mean
เราหารด้วย n ซึ่งเป็นจำนวนของ Sample และตัวของ mean เองก็เป็นจุดกึ่งกลาง
ของ Sample ที่เรานำมาหาเท่านั้น เราไม่สามารถจะทำการวัดจากประชากรทั้งหมดได้
เป็นต้นว่าเราจะทำการวัดความสูงของชาวพันธุ์บางพระ เราไม่สามารถจะวัดได้ทั้งหมดทั่วโลก
หรือแม้แต่ทั่วประเทศไทย เราจึงจำเป็นต้องเก็บตัวอย่างมาจำนวนหนึ่ง ฉะนั้น mean ที่หา
ได้จึงเป็นตัวแทนของตัวอย่าง แต่ในการกล่าวถึงความสูงของชาวบางพระเราหมายถึงชาว
บางพระทั่ว ๆ ไป เพราะคนอื่น ๆ ก็อาจมีตัวอย่างชาวบางพระของคนและอาจต้องการความสูง

ไปใช้ประโยชน์ ถ้าเรากล่าวถึงความสูงของข้าวบางพระภายในวงตัวอย่างของเราก็นับว่าแคบเกินไป ไม่ทำให้ผู้อื่นยึดถือเป็นหลักเกณฑ์ทั่ว ๆ ไปได้ เมื่อเรากอยู่ในฐานะที่สามารถทำการวัดได้แต่เพียงจำนวนน้อย แต่ต้องการทางรายงานผลที่รัดกุมให้เป็นที่ยอมรับ สำหรับการหา

mean นั้นเราไม่มีทางที่จะกระทำได้จากจำนวนประชากรทั้งหมด เราจึงไม่มีทางจะทราบ

mean ที่แท้จริงของประชากร เช่น ข้าวบางพระทั่วทั้งประเทศไทยได้ แต่ก็มีทางที่จะตั้ง

ขอบเขตเพื่อให้มาตรฐานความคลาดเคลื่อนนี้สามารถคลุมไปถึงประชากรทั้งหมดได้โดยการคำนวณจากความคลาดเคลื่อนของตัวเลขภายในตัวอย่างที่รวบรวมมาโดยวิธีการสุ่ม

(Randomization) เราจะเห็นได้ว่าถ้าเราหารเฉลี่ย Sum of square ด้วย n

เช่นเดียวกันกับการหารเฉลี่ยเพื่อหา Mean แล้ว มาตรฐานความคลาดเคลื่อนที่ได้ก็เป็นเพียง

ขอบเขตของ Sample เท่านั้น เช่น จากตารางที่ ๔

$$\begin{aligned}\text{ถ้า } S.D. &= \pm \sqrt{\frac{\sum d^2}{n}} \\ &= \pm \sqrt{\frac{210}{10}} \\ &= \pm 4.58\end{aligned}$$

$$\text{คือ } M \pm S.D. = 168.00 \pm 4.58$$

$$\text{หรือ จาก } 168.00 \pm 4.58 \quad \text{ถึง } 168.00 - 4.58$$

$$\text{จาก } 163.42 \quad \text{ถึง } 172.58$$

ซึ่งเป็นขอบเขตที่แสดงความคลาดเคลื่อนของตัวเลขภายใน Sample ทั้ง ๑๐ จำนวนนั้น จะเห็นได้ว่าประมาณ ๒ ใน ๓ ส่วนของ Sample ตกอยู่ภายในขอบเขต และ

๑ ใน ๓ ส่วนตกอยู่นอกขอบเขตภายใน X_{90} จำนวนนี้ X_6 ซึ่งเท่ากับ ๑๕๕, $X_9 = ๑๗๕$ และ $X_{10} = ๑๗๕$ ตกอยู่นอกขอบเขตสำหรับ X_6 นั้นตกอยู่นอกขอบเขตทางด้านต่ำ แต่

X_9, X_{10} ตกอยู่ทางด้านสูง ถ้าหากจำนวนของ X ยิ่งมากขึ้น และวิธีการสุ่มตัวอย่างรัดกุมดี อัตราส่วนก็จะยิ่งใกล้ ๒ ใน ๓ มากขึ้น แต่จุดหมายอันใหญ่ยิ่งของเรานั้นอยู่ที่ต้อง

การทราบขอบเขตที่แสดงความคลาดเคลื่อนของตัวเลขจากประชากรทั้งหมด ทั้ง ๆ ที่เราไม่สามารถจะวัดจากประชากรทั้งหมดได้ วิธีการหาขอบเขตนั้นก็หาเฉลี่ย Sum of

square ด้วย Degree of freedom $(n-1)$ แทนที่จะหารด้วย n ซึ่งเป็นขนาดของ Sample

ผลที่ได้จะขยายวงกว้างขวางออกไปจนถึงขอบเขตความคลาดเคลื่อนของประชากร ดังได้แสดงไว้แล้วในตารางที่ ๕ แต่ถ้าเป็นกรณีจำนวนตัวอย่างมีมาก ๆ เป็นร้อยพันเรือนั้น การหาค่า n หรือ $n-1$ ก็ไม่มีความแตกต่างเกิดขึ้น เพราะค่าของ n สูงมาก ถ้าจะลบด้วย ๑ ก็คงไม่มีความหมายใด ๆ แต่ในด้านการวิจัยทางการเกษตรแขนงนี้เราเกี่ยวข้องกับเลขจำนวนน้อย จึงจำเป็นต้องเกี่ยวข้องกับ Degree of freedom หรือ $n-1$ อยู่เสมอ *

* "DEGREE OF FREEDOM" นี้ เป็นปัญหาหนึ่งซึ่งการอธิบายปัญหานี้จะต้องเกี่ยวข้องกับทฤษฎีทางคณิตศาสตร์ ในการวิจัยปัญหาทางการเกษตรนั้น ขั้นแรกเราจะต้องวัดตรวจสอบตัวอย่างสักจำนวนหนึ่งซึ่งเราสุ่มเอามา และถือว่ามันมีลักษณะเป็นตัวแทนของพืชหรือสัตว์พันธุ์ใดจำนวนตัวอย่างที่นำมานั้น สมมุติว่าเท่ากับ n เมื่อเราคำนวณหา Mean ของตัวอย่างนั้นแล้วเรามีความมุ่งหมายที่จะใช้ Mean ที่ได้นั้นแทนประชากรของสิ่งนั้นทั้งหมด โดยการคำนวณหา deviation เพื่อนำไปตั้งขอบเขต แต่ตัวเลขที่นำมาคำนวณมาจากตัวเลขเพียงส่วนน้อย ย่อมจะไม่มีคามอิสระที่จะปรวนแปรได้ทุกตัว เนื่องจากผลรวมของ deviation จะต้องเป็น ๐ ตายตัวเสมอ ฉะนั้น ภายหลังจากที่เราเลือกตัวเลขแล้วจนกระทั่งเหลือตัวสุดท้าย เราจะไม่มีโอกาสเลือกตัวสุดท้ายเลย เพราะถ้าเลขตัวสุดท้ายเปลี่ยนไปได้ ผลรวมก็จะไม่เป็น ๐ เสมอไป เช่น สมมุติว่ามีเลขอยู่ ๕ จำนวน คือ ๐, -๑, +๒, -๔, +๔ = ๐ การเลือกครั้งที่ ๑ อาจเป็นจำนวนหนึ่งจำนวนใดใน ๕ จำนวนนั้น เมื่อเลือกครั้งที่ ๑ แล้ว การเลือกครั้งที่ ๒ ก็ยังมีโอกาสเลือกได้ ๔ จำนวน จนกระทั่งเลือกครั้งที่ ๓ แล้ว สมมุติว่าเหลืออีก ๒ จำนวน คือ ๐ กับ -๕ เลขสองจำนวนที่เหลือนี้สำหรับการเลือกครั้งที่ ๔ ซึ่งยังคงได้เลขที่เปลี่ยนแปลงได้ อาจเป็น ๐ ก็ได้ หรือเป็น -๕ ก็ได้ สมมุติว่าในการเลือกครั้งที่ ๔ เราได้ ๐ เลขตัวสุดท้ายต้องเป็น -๕ แน่ ๆ ในครั้งที่ ๕ จะไม่มีโอกาสเลือกได้ เพราะตัวเลขที่เหลือจำนวนสุดท้ายนั้นค่าตายตัว เรียกว่าเสียคามอิสระในการที่จะ variate ไป หรือเสีย degree of freedom ฉะนั้น degree of freedom จึงเท่ากับจำนวน $n-1$ นี้คือคำอธิบายเหตุผลที่ degree of freedom ต้องเสียไป ๑ หรือมีตัวเลขที่มีความอิสระจริง ๆ เพียง $n-1$ เท่านั้นที่เกี่ยวข้องกับการตั้งขอบเขตความคลาดเคลื่อน ส่วนจำนวนสุดท้าย

๑ จำนวนที่เสียไปนั้นก็เท่ากับมอบให้กับผลเฉลี่ยซึ่งนำมาวัดความคลาดเคลื่อนนั่นเอง เพราะฉะนั้นในการหา Mean square หรือ Variance จึงใช้สูตร $\frac{\sum d^2}{n-1}$

ถ้าหากได้ศึกษาคณิตศาสตร์เกี่ยวกับการสถิติแล้วอาจจะได้ยินหรือได้เกี่ยวข้องกับหลักเกณฑ์ของ Least squares ด้วย Least square นี่คือหลักเกณฑ์ว่าด้วยโอกาสที่ผลรวมของความคลาดเคลื่อนยกกำลังสอง ($\sum d^2$) จะมีค่าน้อยที่สุด เราอาจกล่าวได้ว่าถ้าเราวัดความคลาดเคลื่อนโดยอาศัย Mean เป็นหลักแล้ววัดความคลาดเคลื่อนของ X แต่ละตัวที่เคลื่อนไปจาก Mean แล้ว $\sum d^2$ จะมีค่าน้อยที่สุดในเลขชุดนั้นเราจะไม่สามารถหาค่าของ $\sum d^2$ ที่น้อยกว่านั้นได้อีก เช่น จากตัวอย่างในตารางที่ ๔ เราจะเห็นว่า ๒๐๐ เป็นค่าที่น้อยที่สุด ถ้าหากเราใช้เลขจำนวนอื่นที่ไม่ใช่ Mean เป็นเครื่องหมาย deviation แล้ว เราจะได้อะไร $\sum d^2$ สูงกว่า ๒๐๐ เสมอ เพราะว่า Mean เป็นตัวกลางของ deviation นำหนักตัวเลขที่เป็นค่าของ deviation ทั้งสองข้าง คือทางบวกและทางลบ ย่อมเท่ากัน ถ้าไม่ใช่ Mean นำหนักทางบวกหรือทางลบจะฉุดทำให้ค่าของ deviation ทางด้านนั้นสูงขึ้น ยิ่งยกกำลังสองก็ยิ่งสูงมากขึ้นไปอีก

ทฤษฎีว่าด้วย Least Squares อันเกี่ยวข้องกับ Mean นี้ช่วยให้เราเห็นความสำคัญของ degree of freedom ได้กระจ่างแจ้งยิ่งขึ้น เพราะว่า Mean ของตัวอย่างนั้นย่อมจะต้องไม่ตรงกับ Mean ของประชากรที่แท้จริง และโดยเหตุที่เราต้องคำนวณจากตัวอย่างเสมอ ดังนั้น $\sum d^2$ ที่ได้จาก Mean ของตัวอย่างก็ย่อมจะต้องเล็กกว่า $\sum d^2$ ที่ได้จาก Mean ของประชากร ถ้าหากเราเอา Mean ของประชากรมาคำนวณกับตัวอย่างได้ แต่ก็ไม่สามารถทราบ Mean ของประชากรได้ ฉะนั้นการหาร $\sum d^2$ ด้วย $n-1$ จึงช่วยแก้ปัญหาเรื่องนี้ได้

การที่เราสามารถคำนวณหา Mean และ Standard Deviation ซึ่งเป็นเครื่องแสดงขอบเขตความคลาดเคลื่อนของประชากรที่เคลื่อนไปจาก Mean ดังได้กล่าวมาแล้วทั้งหมดนั้น ยืนยันว่าไม่สามารถทำให้ Mean มีความหมายชัดเจนเพียงพอ ยิ่งกว่านั้นสิ่งที่นิยมใช้กันโดยทั่วไปก็คือ Mean จึงทำให้มีผู้นิยมหา Mean กันอยู่เรื่อย ๆ แม้แต่ในการ

วิจัยสิ่งเดียวกันต่างคนต่างก็มีการรวบรวมตัวอย่างของตนเอง และต่างคนต่างก็หา Mean จากตัวอย่างที่ตนรวบรวมได้ สมมุติว่าในการวัดความสูงของชาวพม่าในแคว ไค้มอบในหนึ่งคนควา 10 คน ไปทำการรวบรวมตัวอย่างและวัดความสูงพร้อมทั้งคำนวณหา Mean โดยแต่ละคนก็มุ่งไปวัดความสูงของชาวพม่าในแควคนละ 100 คน โดยวิธีสุ่ม ซึ่งถือว่า ความสูงของชาวพม่าในแคว 100 คน นั้นจะใช้แทนความสูงของประชากรชาวพม่าในแควทั่วทั้งประเทศ ไทย ดังนั้นใน 10 คนนี้ต่างคนต่างก็มีตัวเลขอยู่คนละ 100 จำนวน พร้อมทั้งหา Mean เสร็จเรียบร้อยแล้ว เมื่อนักคนควาทั้ง 10 คนนี้มาประชุมกัน เราจะพบว่าชาวพม่าในแควพันธุ์เดียวกันมี Mean ถึง 10 ค่า ไม่เหมือนกัน จึงเกิดมีปัญหาคือเราจะเชื่อถือค่าของ Mean ค่าใดเพื่อใช้แทนความสูงของชาวพม่าในแควได้ โดยเหตุที่เราไม่สามารถจะทราบได้ว่า Mean ที่แท้จริงของประชากรชาวพม่าในแควนั้นอยู่ที่ไหน และเราก็ไม่ทราบอีกนั่นแหละว่า Mean ของชาวพม่าพันธุ์เดียวกัน ซึ่งต่างคนต่างหามาได้ใน 10 คน หรือถ้าหากมีคนหัวอีก 100 คน หรือ 1,000 คน ฯลฯ ก็คงได้ไม่ตรงกัน นั่น Mean ใดที่จะเชื่อถือได้ เพราะต่างคนไปวัดมาจริง ๆ และกระทำอย่างถูกต้องตามหลักเกณฑ์ทุกอย่าง เมื่อเป็นเช่นนั้นถ้าไม่มีสิ่งใดเป็นมาตรฐานตัดสินแล้ว นักคนควาก็คงจะเถียงกันตลอดไป คอยเหตุผลดั่งกล่าวแล้ว นักสถิติจึงลงความเห็นว่ ในบางโอกาสแทนที่จะใช้มาตรฐานความคลาดเคลื่อนที่แสดงขอบเขตของความคลาดเคลื่อนที่ X แต่ละตัวเคลื่อนไปจาก Mean ดังกล่าวมาแล้วข้างต้น เราควรจะใช้มาตรฐานความคลาดเคลื่อนที่แสดงขอบเขตของ Mean ซึ่งหาได้จากครั้งอื่นอันคลาดเคลื่อนไปจาก Mean ที่แสดงไว้นั้น มาตรฐานนี้เราเรียกว่ามาตรฐานความคลาดเคลื่อน Mean (Standard Error of Mean, S.E.M) ซึ่งมีขอบเขตคลุมกว้างออกไปไม่ว่าเราจะหา Mean เท่าแล้วเท่าอีกอยู่เรื่อย ๆ ไป โดยไม่มีวันเสร็จสิ้น Mean ที่ใดมาทุก ๆ ครั้งนั้น เราไม่มีโอกาสทราบได้ว่า Mean ตัวใดที่ใกล้เคียงกับ Mean จริงของประชากร เพราะตัว Mean ที่แท้จริงของประชากรเราก็ไม่มีโอกาสจะหาได้ แต่ตามทฤษฎีที่ว่าด้วย Probability curve นั้น เราสามารถกล่าวได้ว่าจำนวนของ Mean ที่ใกล้เคียงความจริงที่สุดนั้นจะมีจำนวนมากกว่า Mean ที่มีค่าไกลกว่า ฉะนั้นจึงสามารถกล่าวได้ว่า Mean ที่เราหาได้นั้นมีค่าของ S.E.M ต่ำหรือมีขอบเขต

แถบก็แสดงว่า ๒ ใน ๓ ของ Mean ทั่ว ๆ ไปอยู่ใกล้เคียงกับ Mean ของตัวอย่างของ
เรานั้น ซึ่งเท่ากับว่า Mean ของตัวอย่างนั้นมีความแน่นอนสูง หรือใกล้เคียงความเป็นจริงมาก

ตามธรรมชาติความคลาดเคลื่อนในระหว่าง Mean ต่อ Mean ควบกันย่อมจะมีขอบ
เขตแคบกว่าความคลาดเคลื่อนระหว่าง variable X แต่ละตัวด้วยกัน และ S.E. นั้น
จะเป็นปฏิภาคส่วนกลับกับ S.D. และกรณีของจำนวนรายการในตัวอย่างนั้น ๆ ดังนั้นสูตร
ของ S.E. จึงเป็นดังนี้

$$S.D. = S.E._M \sqrt{n}$$

$$S.E._M = \frac{S.D.}{\sqrt{n}}$$

$$\text{but } S.D. = \sqrt{\frac{\sum d^2}{n-1}}$$

$$S.E._M = \pm \sqrt{\frac{\sum d^2}{n-1}} \cdot \frac{1}{n} \quad \text{or} = \pm \sqrt{\frac{\sum d^2}{n(n-1)}} \quad +++$$

จะขอยกตัวอย่างวิธีคำนวณหา S.E._M ดังต่อไปนี้ จากตัวเลขในตารางที่ ๔
เราได้อาคของผลรวมของกำลังสองของ deviation หรือที่เรียกว่า Sum of Square (S.S.)
ซึ่งเท่ากับ ๒๑๐

$$S.E._M = \pm \sqrt{\frac{\sum d^2}{n(n-1)}}$$

$$= \pm \sqrt{\frac{210}{10(10-1)}}$$

$$= \pm \sqrt{\frac{210}{90}}$$

$$= \pm \sqrt{2.333}$$

$$= \pm 1.527$$

$$M \pm S.E._M = 168.00 \pm 1.527$$

$$\text{or } 168.00 + 1.527 \text{ to } 168.00 - 1.527$$

$$\text{ขอบเขต} = 166.473 \text{ to } 169.527$$

ซึ่งแสดงว่าในเรื่องเดียวกันนี้ ถ้าหากเราจะไปหา Mean จากตัวอย่างอื่น ๆ อีก สักตัวอย่าง เราจะได้ Mean บามากมายสักเพียงใดก็ตาม ๒ ใน ๓ ของจำนวน ทั้งหมดจะตกอยู่ภายในขอบเขตนี้และอีก ๑ ใน ๓ จะตกอยู่นอกขอบเขต

สูตร $\sqrt{\frac{\sum d^2}{n(n-1)}}$ ที่เรากำนวนหา S.E._M นั้น เป็นสูตรที่ได้มาโดยตรงจาก ทฤษฎี ดังที่ได้อธิบายมาแล้วเป็นตอน ๆ จากสูตรนี้เอง เราสามารถจะดัดแปลงไปได้อีกตาม ความเหมาะสมและลักษณะของการใช้งาน ดังต่อไปนี้

สูตรเดิมที่กล่าวมาแล้วนั้นเราจะคำนวณได้ ก็ต่อเมื่อเราได้คำนวณหา Mean ของ ตัวอย่างเสียก่อน แล้วจึงเอา Mean ไปลบ X แต่ละตัวเพื่อหา "d"

ต่อไปนี้สมมติว่าเราจะหา Mean ขึ้นสักจำนวนหนึ่ง ซึ่งเราเรียกว่า guess mean (M_g) เป็น Mean ที่เราสมมุติขึ้นโดยประมาณค่าปานกลางระหว่างค่าของเลข ตัวอย่างนั้น ดังต่อไปนี้

ตารางที่ ๕
เลขจากตารางที่ ๔

No.	X	d	d ²
๑	๑๓๐	๕	๒๕
๒	๑๒๔	๔	๑๖
๓	๑๒๔	๑	๑
๔	๑๓๑	๖	๓๖
๕	๑๒๕	๐	๐
๖	๑๕๔	๖	๓๖
๗	๑๒๓	๒	๔
๘	๑๒๖	๑	๑
๙	๑๓๕	๑๐	๑๐๐
๑๐	๑๓๔	๙	๘๑

$$\sum d = ๓๐ \quad \sum d^2 = ๓๐๐$$

Mean ค่า = ๑๖๕.๐๐

ค่าของ d มีความจริงรวมทั้งหมดสิบตัว = ๓๐ (ความจริง d ต้องเท่า ๐)

เฉลี่ยค่าของ d มีความจริงแต่ละ $\frac{30}{10} = 3$

แสดงว่า Mean ค่าไป ๓ ต้องแก้ไขโดยเอา ๑๖๕.๐๐ + ๓.๐๐ =

๑๖๘.๐๐ ค่า d นิดไป ค่า d^2 ก็ต้องนิดไปด้วย และค่า d^2 ก็นิดไปด้วย วิธีแก้ เอาค่า

$\sum d = 30$ ซึ่งแสดงความจริงทั้งหมด ยกกำลังสอง แล้วหารเฉลี่ยด้วยจำนวน เราเรียก
กันว่า

$$\text{Correction Factor} = \frac{(\sum d)^2}{n}$$

(C.F.)

$$\text{C.F.} = \frac{30^2}{10} = 90$$

แล้วเอา C.F. ซึ่งเท่ากับ ๙๐ นี้ไปแก้ความจริงของ $\sum d^2$ โดยหักออก

$$= 300 - 90$$

$$= 210 \text{ ซึ่งตรงกับค่า } \sum d^2 \text{ ที่ถูกต้องในตารางที่ ๕}$$

คำอธิบายตารางที่ ๕

สมมติว่าเราเอา Mean โดยประมาณค่า มีค่าเป็น ๑๖๕.๐๐ เราก็เริ่มการ
คำนวณหา d อย่างเดียวกันกับที่ทำมาแล้ว โดยเอา Mean ๑๖๕ ซึ่งเอาขึ้นนั้นไปลบ X
แต่ละตัว เสร็จแล้วหา d^2 เราจะสังเกตเห็นว่า Mean ที่เราเอาขึ้นนั้นผิดจาก Mean
จริง เพราะผลรวมของ d = ๓๐ แสดงว่า Mean ที่เอานั้นไม่ได้อยู่ที่จุดกึ่งกลาง ถ้า
Mean อยู่จุดกึ่งกลางจริงผลรวมของ d ต้องเท่ากับ ๐ เสมอ เพราะค่าทางบวกและ
ทางลบจะเท่ากันพอดี นี่แสดงว่า Mean ผิด และทำให้ d แต่ละตัวผิดไปด้วย เมื่อปรากฏ
ว่าผลรวมของ d = ๓๐ ก็แสดงว่า d ทั้ง ๑๐ ตัวนั้นมีความจริงไปถึง ๓๐

$$\text{เฉลี่ยความจริงของ } d = \frac{30}{10}$$

$$\text{เฉลี่ย } d \text{ มีความจริงแต่ละ } = 3$$

แสดงว่า Mean ที่เดานั้นต่ำกว่าความจริงไป ๓ จึงได้ทำให้ d แต่ละตัวเกินไปข้างบวก ๓
เราจะแก้ Mean เค้าให้ถูกต้องได้โดยเอา ๓ นี้ไปบวก $๑๖๕.๐๐ + ๓.๐๐ = ๑๖๘.๐๐$

เพราะฉะนั้น สูตรในการคำนวณ Mean จริงจาก Mean เค้า (M_g) $M = M_g + \frac{\sum d}{n}$
(d ในที่นี้หมายถึง d ที่ได้จาก Mean เค้า) ถ้าค่าของ $\sum d$ ที่ได้จาก Mean เค้า
เป็นลบ ก็เป็นไปตามเครื่องหมาย เพราะแสดงว่า Mean ที่เราเดานั้นสูงเกินความจริง
ของแฉกควมการลบออก

สูตรที่ใช้ในการคำนวณ S.E.M จาก M_g

$$\begin{aligned} S.E.M &= \pm \sqrt{\frac{\sum d^2 - (\sum d)^2/n}{n(n-1)}} \quad \text{or} \quad \sqrt{\frac{\sum d - C.F.}{n(n-1)}} \\ &= \pm \sqrt{\frac{300 - 30^2/10}{10(10-1)}} \\ &= \pm \sqrt{\frac{210}{90}} \\ &= \pm 1.527 \end{aligned}$$

ได้ผลถูกต้องตามความจริงทุกประการ

โดยเหตุนี้เราอาจไม่ต้องเสียเวลาบวกเลขค่าของ X แต่ละตัว เพื่อรวมแล้ว
หารหา Mean เราอาจจะประมาณ Mean เอาตามใจชอบ เสร็จแล้วแก้ความผิดเอาที่หลัง
ก็ได้ผลเท่ากัน

ยังมีวิธีการคำนวณ S.E.M อีกวิธีหนึ่งที่ยากกว่าวิธีที่กล่าวมาแล้ว แต่ก็เป็นวิธี
การที่ได้มานมาจากวิธีแรก และวิธีที่สอง เป็นชั้น ๆ

จากวิธีการเดา Mean ที่ได้กล่าวมาแล้วนั่นเอง สมมติว่าเราจะเดาให้เท่ากับ a
เราก็กระทำได้ ดังต่อไปนี้

ตารางที่ ๖
ตัวเลขเดิมจากตารางที่ ๔ และที่ ๕

No.	X	$(X - M_g)$	d^2
๑	๑๓๖	๑๓๖	๒๔๔๐๐
๒	๑๖๔	๑๖๔	๒๔๔๖๑
๓	๑๖๔	๑๖๔	๒๖๘๔๖
๔	๑๓๑	๑๓๑	๒๔๒๔๑
๕	๑๖๕	๑๖๕	๒๗๒๒๕
๖	๑๕๕	๑๕๕	๒๔๒๔๑
๗	๑๖๓	๑๖๓	๒๖๗๔๙
๘	๑๖๖	๑๖๖	๒๗๕๕๖
๙	๑๓๕	๑๓๕	๓๔๒๒๕
๑๐	๑๓๕	๑๓๕	๓๐๒๒๖
๑๖๘๐		๑๖๘๐	๒๔๒๔๕๐

สมมติว่าเราเอา Mean (M_g) = 0 *

เมื่อคำนวณหา d โดยเอา M_g ไปลบ X แต่ละตัวก็เท่ากับเอา ๐ ไปลบ

ค่าของ d = ค่าของ X

เมื่อใช้สูตรคำนวณ S.E._M จาก Mean เดิม = $\pm \sqrt{\frac{\sum d^2 - (\sum d)^2/n}{n(n-1)}}$

เนื่องจากค่าของ d กับของ X เท่ากัน เราอาจเขียนสูตรใหม่จากสูตรเดิม

คือ ในเมื่อ Mean เดิมเป็น ๐

$$\begin{aligned}
 S.E._M &= \pm \sqrt{\frac{\sum X^2 - (\sum X)^2/n}{n(n-1)}} \\
 &= \pm \sqrt{\frac{282450 - (1680)^2/10}{10(10-1)}}
 \end{aligned}$$

$$= \pm \sqrt{2.33} = \pm 1.527$$

$$\begin{aligned} \text{แทนค่าตามสูตร} \quad S.E._M &= \pm \sqrt{\frac{282450 - (1680)^2 / 10}{n(n-1)}} \\ &= \pm \sqrt{\frac{21}{9}} = \pm \sqrt{2.33} = \pm 1.527 \end{aligned}$$

วิธีคำนวณโดยใช้หา Mean จริงและหา deviation จริง ($\sqrt{\frac{\sum d^2}{n(n-1)}}$)

กับวิธีคำนวณโดยหา Mean ให้มีค่าเป็น 0 ($\sqrt{\frac{\sum X^2 - (\sum X)^2 / n}{n(n-1)}}$) แม้ทั้งสองวิธีนี้

จะได้ผลลัพธ์เท่ากัน แต่โอกาสที่ควรจะใช้วิธีการคำนวณแบบใดนั้น จำเป็นจะต้องพิจารณาเลือกใช้วิธีการที่สั้น ไม่เสียเวลาและสะดวกสำหรับวิธีที่คำนวณจากค่าของ d จริง ๆ นั้น แม้จะเป็นวิธีสั้น เข้าใจง่าย แต่ทำให้เครื่องคิดเลขก็หนักว่าไม่สะดวก เพราะต้องมีวิธีการลบเพื่อจะหา d แต่ละตัว ทำให้เสียเวลาและอาจผิดได้ง่าย ส่วนวิธีคำนวณจากค่าของ X นั้น มีวิธีการบวกกับคูณเป็นส่วนใหญ่ นับว่าสะดวกในการใช้เครื่องคิดเลข แต่ถ้าหากไม่มีเครื่องคิดเลข การคำนวณวิธีนี้จะได้รับความลำบากเกี่ยวกับการที่ตัวเลขมากขึ้นทุกที ผลสุดท้ายเวลาแทนค่าในสูตรจะได้ตัวเลขมาก เพราะมีแต่การบวกและยกกำลังสองเรื่อย ๆ ดังจะเห็นได้จากยอดของรายการตัวเลขในตารางที่ ๖ ซึ่งแสดงตัวเลขจำนวนมาก เมื่อเปรียบเทียบกับตารางที่ ๔ ซึ่งมีตัวเลขน้อย แต่เมื่อคำนวณเสร็จแล้วจะได้ผลลัพธ์เท่ากัน

ในการใช้สูตร X แม้จะทำให้ตัวเลขมากขึ้นทำให้ไม่สะดวกในการคำนวณในเมื่อไม่มีเครื่องคิดเลข แต่เราก็สามารถแก้ปัญหานี้ได้โดยทำการ code ตัวเลขให้น้อยลงเสียก่อนจึงจะทำการคำนวณได้

การ code ตัวเลขนี้มีอยู่หลายวิธีด้วยกัน แต่ที่เราใช้โดยทั่ว ๆ ไปนั้น ใช้ code ด้วยวิธีลบ ความมุ่งหมายของการ code ด้วยวิธีลบ ก็เพื่อจะตัดทอนตัวเลขที่เรา นำมาเปรียบเทียบกับนั้นให้มันน้อยลง แต่เพื่อเป็นการยุติธรรมจึงจำเป็นต้องลบออกตัวเลข เท่า ๆ กันด้วย ยกตัวอย่างเช่น เราต้องการวัดความเหลื่อมล้ำค่าสูงระหว่างนาย ก. กับ นาย ข. สมมุติว่านาย ก. สูง ๑๖๐ ซม. นาย ข. สูง ๑๖๕ ซม. แสดงว่านาย ข. สูงกว่านาย ก. $๑๖๕ - ๑๖๐ = ๕$ ซม. แต่ถ้าวัดความสูงของนาย ก. และนาย ข. ออกคนละเท่า ๆ กันคือ -๑๕๐ ผลที่ได้รับมีความสูง code ของนาย ก. = ๑๐ ซม. ส่วน ความสูง code ของนาย ข. = $๑๖๕ - ๑๕๐ = ๑๕$ ซม. เมื่อนำความสูง code ของ ทั้งสองคนมาเปรียบเทียบกับกัน คือ $๑๕ - ๑๐ = ๕$ ซม. เท่ากัน แทนที่เราจะเอา $๑๖๕ - ๑๖๐$ แต่เอา $๑๕ - ๑๐$ จะสะดวกกว่า

ตารางที่ ๗

จากตารางที่ ๖ ถ้าเราใช้คำนวณโดยการ Code ด้วยวิธีลบ

(Coded by - 165)

No.	X	X.coded - 165	X^2
1	170	$170 - 165 = 5$	25
2	169	$169 - 165 = 4$	16
3	164	$164 - 165 = -1$	1
4	171	$171 - 165 = 6$	36
5	165	$165 - 165 = 0$	0
6	159	$159 - 165 = -6$	36
7	167	$167 - 165 = 2$	4
8	166	$166 - 165 = 1$	1
9	175	$175 - 165 = 10$	100
10	174	$174 - 165 = 9$	81
		Sum of X = 30	Sum of X^2 = 300

$$\begin{aligned}
 S.E._M &= \pm \sqrt{\frac{\sum X^2 - (\sum X)^2 / n}{n(n-1)}} \\
 &= \pm \sqrt{\frac{300 - (30)^2 / 10}{10(10-1)}} = \pm 1.527
 \end{aligned}$$

ดังที่ได้ยกตัวอย่างมาแล้วว่า ถ้าเราต้องการหาความคลาดเคลื่อน ความแตกต่าง หรือความเหลื่อมล้ำค่าสูง ผลของการคำนวณด้วยวิธี code จะแสดงผลที่แท้จริงให้ปรากฏ เช่น การเปรียบเทียบความสูงระหว่างนาย ก. กับนาย ข. เป็นต้น ฉะนั้นในการหา $S.E._M$ ซึ่งก็เป็นการวัดความคลาดเคลื่อนไปมาของตัวเลข เราก็จะไต่ค่าของ $S.E._M$ ที่แท้จริงเช่นกัน ดังจะเห็นได้ว่า ผลที่จากการคำนวณหา $S.E._M$ ด้วยวิธี code และไป code ให้ผลเหมือนกัน แต่วิธี code ทำให้สับสนขึ้นมาก ดังจะเปรียบเทียบได้จากตารางที่ ๖ และที่ ๗ ในตารางที่ ๗ นั้น ตัวเลขน้อยลงจนกระทั่งสามารถคำนวณหา X^2 ได้โดยการคิดในใจ

แต่ในการคำนวณหา Mean ที่แท้จริงจากวิธี code นั้น เราก็กระทำได้ แต่เมื่อไต่ค่าของ mean ให้ code แล้วจะต้อง decode กลับเป็น Mean ที่แท้จริงด้วย ดังตัวอย่างในตารางที่ ๗ ผลรวมของ $X \text{ coded} = 30$ ในการหา Mean

$$\begin{aligned}
 M &= \frac{\sum X}{n} = \frac{30}{10} \\
 &= 3.0
 \end{aligned}$$

๓.๐ เป็น Mean ที่ได้จาก code

เนื่องจากตัวเลขที่แท้จริงนั้น เราลบไว้แต่แรกทุก ๆ ตัว ๆ ละ ๑๖๕ เมื่อเราทราบผลเฉลี่ยของเลขหนึ่งตัว คือ ๓ เมื่อต้องการทราบผลเฉลี่ยที่แท้จริงจึงต้องเอา ๑๖๕ ไปบวกกลับคืน ซึ่งเท่ากับ ๑๖๘ ตรงกับ Mean ที่แท้จริงในตารางที่ ๔ ผลของการ code จะทำให้เลขลดน้อยลงเพียงใดนั้น ย่อมขึ้นอยู่กับค่าของตัว code ที่เราวางลงไป เนื่อง จากตัว code เป็นเลขที่เราสมมุติขึ้นไว้ ส่วนการใช้ได้ผลดีหรือไม่นั้น ย่อมขึ้นอยู่กับ การสมมุติตัว code ที่เหมาะสม โดยอาศัยหลักเกณฑ์ดังต่อไปนี้

๑. การสมมติตัว code ควรจะประมาณเอาตัวที่ใกล้ตัวกลางที่สุด โดยสังเกตดูค่าของตัวเลขแต่ละจำนวน แล้วหาค่าตัวใดมากที่สุด เสร็จแล้วจึงวางค่าของตัว code ให้ประมาณกลาง ๆ ค่าของตัว code ยิ่งใกล้ตัวกลางเข้าไปเท่าใด ค่าของ $\sum X^2$ ก็จะยิ่งลดน้อยลง และถ้าเมื่อใดค่าตัว code ไปตรงกับค่าตัวกลาง หรือ Mean พอดี $\sum X^2$ ก็จะมีค่าน้อยที่สุด Mean code ก็จะมีค่าเป็น ๐ และ correction factor ($\sum X^2/n$) ในสูตร X ก็จะเป็น ๐ ด้วย

๒. ถ้าค่าของตัว code ต่ำกว่า Mean ผลรวมของ X code จะเป็นบวกเสมอ และ Mean code ก็จะเป็นบวกด้วย เมื่อ decode กลับ ก็จำเป็นจะต้องเอาค่าของตัว code ไปบวกกับ Mean code ให้เป็น Mean จริง

๓. ถ้าค่าของตัว code สูงกว่า Mean ผลรวมของ X code จะเป็นลบเสมอ และ Mean code ก็จะเป็นลบด้วย เมื่อ decode กลับ ก็จำเป็นจะต้องเอาค่าของ Mean code ไปลบตัว code เพื่อให้ได้ Mean จริง

ความจริงวิธีการ code นี้เป็นเพียงวิธีการคำนวณธรรมดาเพื่อจะช่วยให้การคิดเลขสะดวก และสั้นเข้าเท่านั้น ไม่ได้เข้าไปเกี่ยวข้องกับหลักเกณฑ์ในทางสถิติแต่อย่างใดเลย แต่ว่าจากหลักเกณฑ์ในการวางตัว code ข้อ ๑ นั้น ดู ๆ ก็คล้ายกับว่าวิธีการ code จะเข้ามาเกี่ยวข้องกับหลักเกณฑ์ทางสถิติด้วย ทั้งนี้เพราะเหตุว่าวิธีการ code ด้วยวิธีแบบนี้มีวิธีการอย่างเดียวกันกับการคำนวณหา deviation ในทางสถิติ หากแต่ว่าถ้าค่าของตัว code ไม่ตรงกับค่าของ Mean ตัว code นั้นก็เปรียบเสมือน Mean เถานั้นเอง แต่หาหาค่าของตัว code บังเอิญไปตรงกับ Mean เข้า ค่าของตัว code ก็คือค่าของ Mean และการคำนวณหา X code ก็คือหา deviation ที่แท้จริง $\sum X^2$ ก็จะเท่ากับ d^2 ดังนั้น correction factor จึงเป็น ๐ ดังตารางเปรียบเทียบดังต่อไปนี้

X	X coded by - 168	X^2
170	2	4
169	1	1
164	-4	16
171	3	9
165	-3	9
159	-9	81
167	-1	1
166	-2	4
175	7	49
174	6	36

Sum of coded X = 0

Sum of X^2 = 210
==

X	$d = (X - M)$	d^2
170	2	4
169	1	1
164	-4	16
171	3	9
165	-3	9
159	-9	81
167	-1	1
166	-2	4
175	7	49
174	6	36
$\sum X = 1680$	$\sum d = 0$	$\sum d^2 = 210$

$$\begin{aligned}\text{Mean} &= \frac{1680}{10} \\ &= 168.0\end{aligned}$$

$$\text{S.E.}_M = \pm \sqrt{\frac{\sum x^2 - (\sum x)^2/n}{n(n-1)}}$$

$$= \pm \sqrt{\frac{210 - 0}{10(10-1)}}$$

$$= \pm 1.527$$

$$\text{or S.E.}_M = \pm \sqrt{\frac{\sum d^2}{n(n-1)}}$$

$$= \pm \sqrt{\frac{210}{10(10-1)}}$$

$$= \pm 1.527$$

ฉะนั้น ในการที่จะเลือกใช้สูตร $\sqrt{\frac{\sum d^2}{n(n-1)}}$ หรือ $\sqrt{\frac{\sum x^2 - (\sum x)^2/n}{n(n-1)}}$ รวมทั้ง

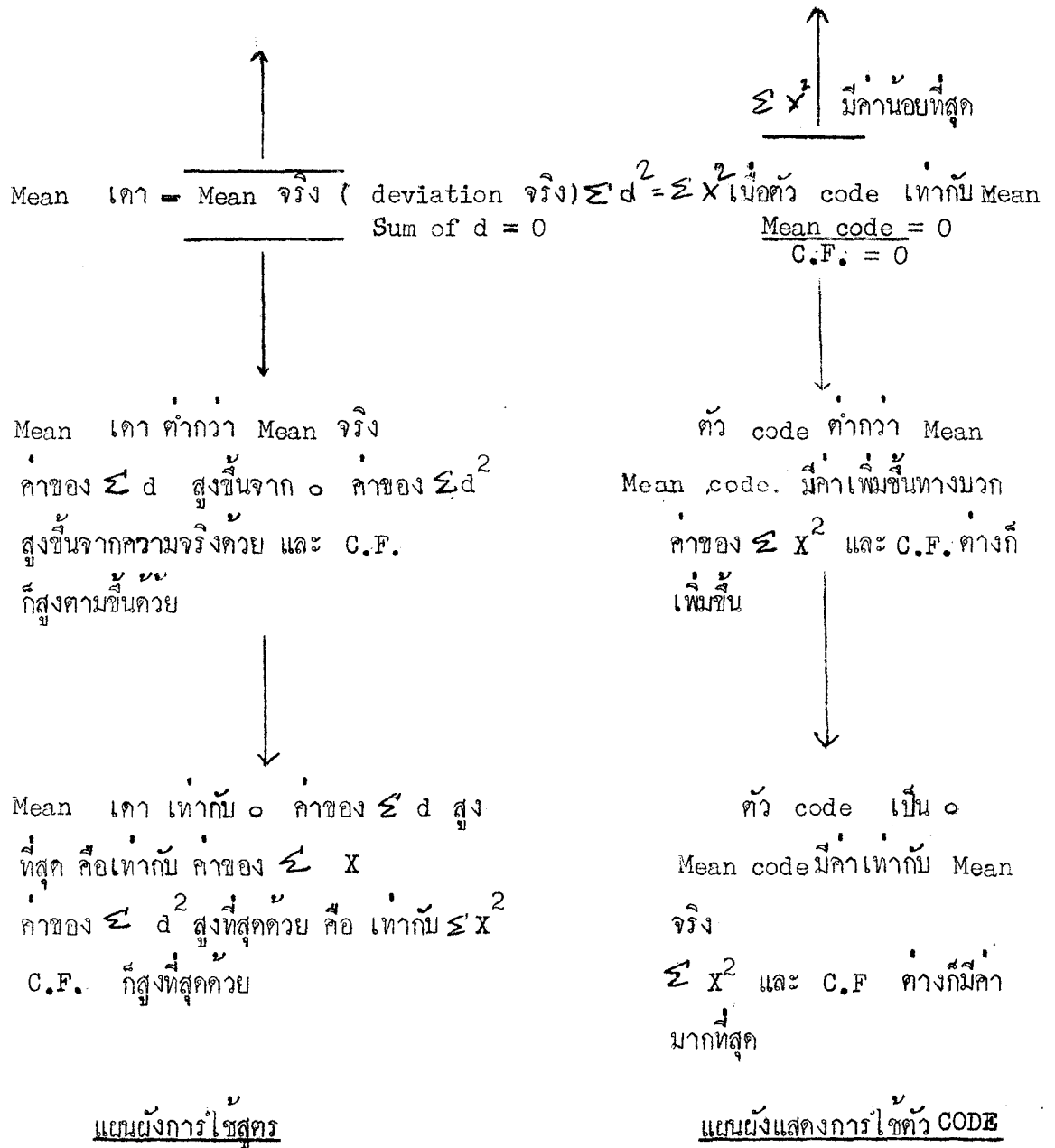
ถ้าจะใช้วิธี code เข้ามาเกี่ยวข้องของควยแล้ว ก็ควรจะทำความเข้าใจกับแผนผังความสัมพันธ์ของการเปลี่ยนแปลงในสิ่งเหล่านี้ไว้มาก เพื่อจะได้ยึดถือเป็นหลักปฏิบัติในการเลือกใช้สูตร หรือวางตัว code ได้เหมาะสม

C.F. ก็สูงขึ้นจากความจริงควยค่าของ $\sum d^2$ จะสูงขึ้นจากความจริงค่าของ $\sum d^2$ จะต่ำกว่า ๐ คือเพิ่มขึ้นทางลบ Mean เดากสูงกว่า Mean จริง

$\sum x^2$ และ C.F. ต่างก็เพิ่มขึ้น

Mean code มีค่าลดลงเป็นลบ

ตัว code สูงกว่า Mean



จะเห็นว่า แม้ทั้งสองวิธีการนี้มีความมุ่งหมายไปคนละทาง แต่วิธีการคำนวณก็ยังมีส่วนสัมพันธ์กัน เช่น สมมุติว่าในการ code เลขใดค่าของตัวเลขน้อยลงเพื่อสะดวกในการคำนวณ บังเอิญตัว code ที่ตั้งขึ้นนั้นไปตรงกับค่าของ Mean เข้าพอดี การหา $\sum X$ code ก็คือวิธีหา $\sum d$ นั่นเอง เมื่อเรารวมเป็น $\sum X$ ก็จะได้ 0 เหมือนกับ $\sum d$ ในทางองเดียวกัน

ถ้าเราหา $\bar{x}^2$ $\bar{x}^2$ ก็มีค่าเท่ากับ d^2 ฉะนั้น C.F. จึงจำเป็นต้องเท่ากับ ๐ ด้วย นอกจากนั้น Mean code ก็คือ $\frac{\sum d}{n} = 0$ และเนื่องจากตัว code เท่ากับ Mean อยู่แล้ว เมื่อ decode กลับ ก็เท่ากับเอา Mean จริงมาบวกกับ ๐ ก็ยังมีค่าตามเดิม

ความมุ่งหมายในการวิจัยทางการเกษตรโดยทั่ว ๆ ไป นั้นมีจุดหมายที่สำคัญอยู่

๓ ประการ คือ

๑. เพื่อทำการเปรียบเทียบ (Comparison) เป็นต้นว่าพืชพันธุ์ไหนจะให้ผลผลิตสูงกว่ากัน ไร่ปุ๋ยชนิดไหนหรือส่วนผสมอย่างไรจึงจะทำให้พืชชนิดนั้นเจริญเติบโตเร็วที่สุด วิธีการปลูกอย่างไรจึงจะทำให้พืชเจริญงอกงามที่สุด หรืออาหารพืชชนิดไหนที่ช่วยสร้างความเจริญเติบโตให้แก่สัตว์ที่สุด เป็นต้น

๒. เพื่อหาความสัมพันธ์ซึ่งกันและกัน (Relation) เป็นต้นว่า ถ้าปุ๋ยมากขึ้น พืชจะโตเร็วจริงหรือไม่ พืชที่มีลำต้นแข็งแรงอวบอ้วนจะให้ผลผลิตสูงกว่าหรือเปล่า โถที่มีกระดุกของหอยทากวางจะไขดกจริงหรือไม่ เหล่านี้เป็นการค้นคว้า โดยมีความมุ่งหมายที่จะหาความสัมพันธ์ของลักษณะทางธรรมชาติทั้งสิ้น

๓. เพื่อศึกษา (Study) ลักษณะประจำของพืชและสัตว์พันธุ์นั้น ๆ ไว้ เพื่อเป็นเครื่องประกอบการพิจารณาในการปฏิบัติให้ถูกต้อง หรือเพื่อใช้ประกอบการแก้ไขปรับปรุงคุณภาพให้ดียิ่ง ๆ ขึ้นไป สำหรับความมุ่งหมายข้อนี้เป็นเพียงส่วนปลีกย่อย เพราะผลสุดท้ายกับก็จะไปสิ้นสุดลงที่ข้อ ๒ หรือข้อ ๑ หรือทั้งสองข้อรวมกัน

ในการทดลองเปรียบเทียบสิ่งใด ๆ ก็ตาม ก่อนที่เราจะตัดสินลงไปได้ว่าสิ่งใด ดีแล้ว เราจำเป็นต้องพิจารณาโดยรวมรอบ และตัดสินด้วยความยุติธรรมเที่ยงตรง ผลการตัดสินของเราจึงจะเป็นที่แน่นอนเชื่อถือได้ เป็นต้นว่าเราจะเอา นาย ก. กับ นาย ข. มาขึ้นเปรียบเทียบความสูงกัน ก่อนที่เราจะตัดสินลงไปได้ว่าใครสูงกว่าใครนั้น จำเป็นจะต้องสำรวจดูเสียก่อนว่า เราให้ความเสมอภาคแก่ทั้งสองคนนั้นเพียงพอแล้วหรือยัง เช่น ยืนอยู่บนพื้นระดับเดียวกันหรือเปล่า ใส่รองเท้าสูงเท่า ๆ กันหรือไม่เช่นนั้น เป็นต้น

ถ้าหากเราต้องการจะเปรียบเทียบผลผลิตของพืชสองพันธุ์ เช่นกะหล่ำปลีพันธุ์จีน กับกะหล่ำปลีพันธุ์ไคเป็นอาทิ ก่อนอื่นเราจำเป็นจะต้องศึกษาเสียก่อนว่าสาเหตุอันใดที่อาจทำให้เกิดความแตกต่างขึ้นได้นอกจากพันธุ์ เป็นต้นว่า ความอุดมสมบูรณ์ของดิน การใส่ปุ๋ย การให้น้ำ การปฏิบัติรักษาอื่น ๆ ฯลฯ เมื่อเราทราบคิแล้ว ก็พยายามทำสิ่งเหล่านี้ให้เสมอภาคกันคือให้กะหล่ำปลีทั้งสองพันธุ์ที่จะเปรียบเทียบกันนั้นได้รับสิ่งแวดล้อมเหล่านี้เหมือนกัน เป็นต้นว่า ใส่ปุ๋ยอย่างเดียวกัน จำนวนเท่ากัน ระยะเวลาเดียวกัน รดน้ำเท่า ๆ กัน พรวนดินเหมือนกัน ฯลฯ เราก็สามารถที่จะเปรียบเทียบผลผลิตกันได้ ฉะนั้นนักค้นคว้าทางการเกษตรจึงจำเป็นต้องมีความรู้ทางการเกษตร เป็นอย่างดี รวมทั้งได้ศึกษาถึงความต้องการของพืชหรือสัตว์นั้น ๆ และสาเหตุต่าง ๆ ที่เข้ามากระทบทำให้ผลการค้นคว้าคลาดเคลื่อนไปจากความเป็นจริง เพื่อจะได้หาทางป้องกันแก้ไขสิ่งเหล่านี้ได้ โดยการวางแผนการทดลองที่รัดกุมและรอบคอบไว้ทุก ๆ ด้าน ถ้าหากนักค้นคว้ามีพื้นฐานความรู้ไม่เพียงพอ หรือวางแผนการค้นคว้าไม่รัดกุม ผลการค้นคว้าก็ย่อมคลาดเคลื่อนไปจากความเป็นจริง ซึ่งย่อมจะหมายความว่าเสียเวลา เสียเงินลงทุน เสียแรงงานไปโดยไม่ได้รับผลตอบแทนใด ๆ โดยเฉพาะถ้ารายงานผลการค้นคว้านั้นมีความจริง หากมีผู้เชื่อถือนำไปปฏิบัติแล้วก็ย่อมจะทำให้ได้รับความเสียหายสืบทอดไป ซึ่งจะทำให้ขาดความเชื่อถือในผลงานครั้งต่อ ๆ ไปด้วย สาเหตุที่ทำให้เกิดความคลาดเคลื่อนขึ้นในการทดลองนั้นเราสามารถจะแบ่งออกได้เป็นสามข้อใหญ่ ๆ คือ

๑. สาเหตุที่เรารู้และควบคุมได้
๒. สาเหตุที่เรารู้แต่ควบคุมไม่ได้
๓. สาเหตุที่เราไม่รู้

สาเหตุที่เรารู้และควบคุมได้นี้ โดยทั่ว ๆ ไปเป็นสาเหตุเกี่ยวแก่การปฏิบัติทั่วไป เพื่อทะนุบำรุงหรือควบคุมแก้ไขเจริญเติบโตตามกฎเกณฑ์ของธรรมชาติ เช่นการใส่ปุ๋ย การพรวนดิน การรดน้ำ การให้อาหาร ส่วนผสมของอาหาร ฯลฯ เราสามารถจะควบคุมได้โดยการใส่ปุ๋ยเท่า ๆ กัน พรวนดินเหมือนกัน รดน้ำเท่า ๆ กัน ถ้าเป็นสัตว์ก็ให้อาหารที่ส่วนผสมอย่างเดียวกัน จำนวนเท่า ๆ กัน ฯลฯ เป็นต้น

ไม่สามารถจะนำรายละเอียดมากกล่าวได้ทั้งหมด ยิ่งกว่านั้นสิ่งที่ปรากฏขึ้นแต่ละครั้งนั้นต่างก็ไม่เหมือนกัน เราจึงไม่สามารถจะนำเอาปรากฏการณ์ครั้งหนึ่งครั้งใดมากกล่าวแทน ปรากฏการณ์นั้นทั้งหมดได้ และเราก็ไม่สามารถจะเอาทั้ง ๑๐๐ ครั้ง ๑,๐๐๐ ครั้งมากกล่าวทั้งหมดเพื่อจะอธิบายปรากฏการณ์นั้นได้ นักคนกว่าทดลองจึงได้อธิบายปรากฏการณ์นั้นโดยวิธีการเฉลี่ย และการเฉลี่ยนี้ก็ได้อาจจากการวัดปรากฏการณ์ที่ปรากฏขึ้นหลาย ๆ ครั้ง เพื่อจะหาค่าตัวกลางมาใช้อธิบายแทนปรากฏการณ์นั้น ๆ การที่จะทำการวัดหลาย ๆ ครั้งนี้ การวัดจะต้องกระจายออกไปโดยให้เป็นไปตามโชค (Chance) แต่ไม่ซ้ำของเดิม และไม่มีการลำเอียงหรือเลือกที่รักมักที่ชัง การกระจายการวัดวิธีนี้เราเรียกว่า Replication

ฉะนั้น การกระทำ Replication ก็มีจุดมุ่งหมายที่จะหากลุ่มตัวอย่างซึ่งใช้แทนประชากรและเราก็จะได้รายละเอียดซึ่งเป็นค่าตัวกลางของกลุ่มตัวอย่าง เพื่อจะใช้อธิบายหรือแสดงปรากฏการณ์นั้น ๆ ในการเปรียบเทียบปรากฏการณ์ใด ๆ ก็ตาม เราก็นำตัวแทนนี้ไปเปรียบเทียบกัน นอกจากนี้ที่ได้อีกว่ามาแล้ว Replication ยังได้ช่วยให้สิ่งที่จะนำมาแสดงหรือเปรียบเทียบกัน ได้รับความกระชับกระเทือนจากสาเหตุแวดล้อมตามธรรมชาติสม่ำเสมอและคล้ายคลึงกัน เพื่อจะได้แสดงผลที่แน่นอนและใกล้เคียงความเป็นจริง ซึ่งจะทำให้การตัดสินใจของการเปรียบเทียบนั้นๆ เป็นไปโดยยุติธรรม และได้ผลเป็นที่เชื่อถือได้

ถ้าหากเราจะกล่าวกันตามหลักเกณฑ์ธรรมศาสตร์แล้ว สมมุติ ว่าเราจะบอกกับใครสักคนหนึ่งหรือหลายคนว่าใส่ปุ๋ยชนิดนี้คนไม่จะเจริญเติบโตดีกว่าปุ๋ยชนิดนั้น ก่อนอื่นเขาจะขออนุญาตว่า ทดลองดูแล้วหรือ ซึ่งแสดงว่าการคนกว่าทดลองเป็นสิ่งสำคัญ ถ้าเราตอบเขาว่า ได้ทดลองดูแล้วได้ผลจริงดังที่กล่าวมานั้น เขาจะขออนุญาตอีกว่า ทดลองใส่ปุ๋ยกี่ต้น ทดลองกี่ครั้งกี่หน ถ้าเราตอบเขาว่าครั้งเดียว ต้นเดียว เขาอาจไม่เชื่อผลการทดลองอันนั้น เพราะบางทีคนไม่อาจงามเพราะสาเหตุอื่นก็ได้ เช่นดินหรือเครื่องปลูกทรงนั้นดีเป็นพิเศษ แต่ถาหากตอบว่าได้ทดลองดูทั้ง ๑๐ ครั้ง ๆ ละ ๒๐ ต้น ได้ผลอย่างนั้นจริง ๆ เขาก็จะสนใจแสดงว่าการทำ Replication มีความสำคัญในการช่วยเพิ่มความแน่นอนในผลงานนั้นได้เป็นที่เชื่อถือได้ แต่ถาผู้ฟังคนนั้นมีความรู้ทางการวางแผนคนกว่าทดลองบาง เขาอาจสงสัยอีก

หน้อยหนึ่งและย้อนถามมาอีกว่า คุณไม่ที่ใช้ในการทดลองปูนั้นได้อย่างไร ไปเลือกเอาต้นงาม ๆ มาทดลองกับปูอย่างหนึ่ง และเอาต้นด้อย ๆ ไปทดลองกับปูอีกอย่างหนึ่งหรือเปล่า ถ้าเราตอบว่าเปล่า คุณไม่ที่ใช้ในการทดลองนี้ไม่ได้มาด้วยวิธีการสุ่ม คือหยิบเมล็ด ๆ กันมาไม่มีการเลือก แล้วก็มาแบ่งออกเป็นสองส่วนเท่า ๆ กัน แบบตาม บูดตามกรรมโดยไม่มีการเลือก พวกหนึ่งใส่ปูอย่างหนึ่ง อีกพวกหนึ่งใส่ปูอีกอย่างหนึ่ง นี่ก็แสดงว่าวิธีการ Random มีความจำเป็นในการวางแผนทดลองด้วย ถ้ายังอธิบายรายละเอียดต่อไปอีกว่าการร่อนน้ำก็ร่อนเท่า ๆ กัน การพรวนดินก็เหมือนกัน ฉีดยาฆ่าแมลงก็ฉีดด้วยกัน ฯลฯ ผลการทดลองที่เราแสดงออกมานั้นก็เป็นที่เลื่อมใสแก่ผู้ได้พบเห็นอย่างแน่นอน

วิธีการสุ่ม (RANDOMIZATION) เป็นวิธีการที่นักสถิติทั่วโลกยอมรับกันว่ามีความจำเป็นสำหรับใช้ในการค้นคว้าทดลอง หรืออาจกล่าวได้ว่า ในการค้นคว้าทดลองเพื่อเปรียบเทียบสิ่งใด ๆ ก็ตาม เราจำเป็นต้องเก็บตัวอย่างของสิ่งนั้น ๆ มาทำการวัด โดยเหตุที่เราไม่สามารถจะวัดสิ่งนั้นได้ทุก ๆ ชิ้น ทุก ๆ อัน ทุก ๆ ต้น แต่ผลงานเราจำเป็นต้องแสดงคลุมไปถึงทุก ๆ ชิ้น ทุก ๆ อัน ทุก ๆ ต้น จึงจำเป็นต้องใช้วิธีการสุ่มเพื่อจะนำตัวอย่างที่ได้ซึ่งเป็นที่ยอมรับกันว่าใกล้เคียงความจริง เพื่อมาใช้ในการทดลอง สรุปความว่าวิธีการสุ่มมีความประสงค์ที่จะเฉลี่ยหรือกระจายลักษณะของการปรากฏการณ์ซึ่งปรากฏขึ้นโดยทั่วไป ตามธรรมชาติของประชากรของสิ่งนั้น ๆ ให้มีสัดส่วนเหมือนกับที่ปรากฏขึ้นกับตัวอย่างที่รวบรวมกันได้ เป็นต้นว่าในการเปรียบเทียบพันธุ์พืช ๒ พันธุ์ เราจำเป็นต้องสุ่มดินในบริเวณแปลงทดลอง โดยมีความประสงค์ที่จะให้ดินที่ใช้ปลูกพืชนั้นมีความอุดมสมบูรณ์คล้ายคลึงกับดินในบริเวณทั้งหมด และในเมื่อกระทำทั้ง ๒ พันธุ์ ต่างพันธุ์ต่างก็สุ่มด้วยกัน ทั้งสองพันธุ์ก็就会有ความเสมอภาคในเรื่องความอุดมสมบูรณ์ของดินด้วย จึงสามารถได้หลักเกณฑ์อีกข้อหนึ่งว่า วิธีการสุ่มช่วยให้สิ่งที่จะนำมาเปรียบเทียบกันนั้นได้รับความเสมอภาคโดยได้รับความกระทบกระเทือนจากสิ่งแวดล้อมเท่า ๆ กัน

เราจะจัดวางแผนทดลองอย่างไรจึงจะทำให้เปรียบเทียบได้รับความเสมอภาค ปัญหาที่กล่าวมานี้ ถ้าเราจะใช้หลักเกณฑ์เกี่ยวกับการสุ่มเป็นหลักสำคัญแล้ว เราสามารถที่จะแบ่งออกได้เป็นสองวิธี คือ

๑. วิธีการสุ่มเต็มที่ (Full Randomization System) วิธีนี้ใช้แบบสุ่มลงไป
ในบริเวณพื้นที่ทั้งหมด เปรียบเสมือนหว่านเมล็ดพืชลงในบริเวณที่มีเนื้อที่กว้างมาก ๆ อาจทำ
ให้ส่วนผสมบางอย่างไม่เท่ากัน หรือส่วนผสมแต่ละจุดแต่ละตำแหน่งติดกันไปไต่มา เป็นต้นว่าเรา
ต้องการเปรียบเทียบผลผลิตของพืชพันธุ์ A และพันธุ์ B ในการนี้กระทำพันธุ์ละ ๑๐ ซ้ำ
หรือ ๑๐ แปลง ภายในบริเวณแปลงทดลองจึงจำเป็นต้องแบ่งออกเป็น ๒๐ แปลง ถ้าจะใช้สุ่ม
โดยการจับสลาก และใช้วิธีสุ่มเต็มที่ เราก็เขียน ๑๐ ใบ และ ๑๐ ใบ เมื่อทำการจับ
สลากได้ว่าได้อะไร

B. A. A. A. B. A. B. B. B. B. | A. B. A. A. A. A. A. B. B. B.

เราจะเห็นได้ว่าอาจมีโอกาสเป็นไปได้ที่ A หรือ B อาจจะไปรวมกันเป็นกลุ่ม ลวง
นำนักความอุดมสมบูรณ์ของดินอยู่ทางหนึ่งทางใด ไม่กระจายออกไปที่แปลงเหมือน ๆ กัน
วิธีการสุ่มเต็มที่จึงเป็นวิธีที่ไม่เหมาะในการที่จะใช้กับตัวอย่างซึ่งมีจำนวนน้อย ๆ เพราะถ้า
หาก A มีโอกาสที่จะออกมาก่อนหลาย ๆ คน และบังเอิญไปตกครั้งที่คนที่เข้าแล้ว ซึ่งจะ
ต้องออกทีหลังก็จะถูกแต่คนแล้ว ซึ่งเป็นการเสียเปรียบกันได้

๒. วิธีการสุ่มเป็นแปลงหมู่ (Block System) วิธีการนี้ช่วยกระจายให้
รายละเอียดของการทดลองทดลองได้มีโอกาสได้รับการเปลี่ยนแปลงของสิ่งแวดล้อมเฉลี่ย
เท่า ๆ กัน โดยอาศัยหลักตัดแบ่งเนื้อที่ต้นใหญ่ให้เป็นหน่วยย่อย ๆ แล้วใส่ส่วนย่อยแต่ละส่วน
ก็สิ่งเปรียบเทียบครบถ้วน ฉะนั้น สิ่งเปรียบเทียบ เช่น พันธุ์พืชแต่ละพันธุ์ก็มีโอกาสที่จะ
กระจายไปอยู่ทุก ๆ หน่วยในบริเวณการทดลองนั้น หน่วยนี้เราเรียกแปลงหมู่ (Block)
แต่เราก็จะทั้งวิธีการสุ่มเสียมิได้ หากแต่วิธีการสุ่มจะกระทำภายในแต่ละแปลงหมู่ทุก ๆ
แปลงหมู่ เราจึงสามารถให้คำจำกัดความของแปลงหมู่ หรือ Block ได้ว่า คือหน่วยที่
รวมเอาสิ่งเปรียบเทียบ (Treatment หรือ Variety) ที่มีในการกันคว่ำไว้ โดย
ครบถ้วน แต่มีเพียงสิ่งละหนึ่งแปลงเดี่ยว (Plot) เท่านั้น เพราะฉะนั้นจำนวนของ

Plot ที่อยู่ภายใน Block แต่ละ Block จะเท่ากับจำนวนสิ่งเปรียบเทียบทั้งหมดในการ
กันคว่ำเรื่องนั้นพอดี เมื่อ ๑ Block มีสิ่งเปรียบเทียบอยู่สิ่งละ ๑ Plot ถ้าในการกันคว่ำ

จำเป็นจะต้องทำ ๕ ซ้ำ ก็เท่ากับต้องทำ ๕ Block นั้นเอง ฉะนั้นจำนวนซ้ำหรือ Replication จึงเท่ากับจำนวน Block เสมอไป ที่ได้อธิบายมาแล้วนี้เป็นหลักเกณฑ์สำคัญอันหนึ่งในการวางแผนการกั้นข้าว ซึ่งเป็นที่ยอมรับและนิยมใช้กันทั่วไปและเป็นหลักเกณฑ์ที่รัดกุมที่สุด ระบบแปลงหมู่ที่ได้อธิบายมาแล้วนี้จะเรียกให้เข้าใจง่ายซึ่งยิ่งขึ้น เราอาจเรียกได้ว่าเป็นการวางแผนแปลงหมู่ที่สมบูรณ์ (Complete Block Design)

ตัวอย่างที่ ๑ สมมติว่าถ้าเราจะทดลองกั้นข้าวเปรียบเทียบผลผลิตของพืช ๒ พันธุ์ คือ พันธุ์ A และ B ซึ่งจะต้องกระทำ ๕ ซ้ำ ก่อนอื่นเราต้องแบ่งเป็น ๕ Block, Block ละ ๒ plot เพราะมี ๒ พันธุ์ เสร็จแล้วเราก็ใช้วิธีสุ่มภายในแต่ละ Block ดังแผนผังต่อไปนี้

I		II		III		IV		V	
B	A	B	A	A	B	B	A	A	B

จะเห็นได้ว่าจากผลการสุ่มแต่ละ Block นั้น B หรือ A อาจออกก่อนก็ได้ แต่ภายใน Block, Block หนึ่งต้องมีทั้ง A และ B ครบถ้วน ส่วนเลขประจำ Block นั้นตามหลักสากลนิยมใช้เลขโรมัน

ตัวอย่างที่ ๒ ในการวางแผนการกั้นข้าวเปรียบเทียบ ๓ อัตรา (๓ Treatments) คือ อัตราที่ ๑, ๒ และที่ ๓ ซึ่งจะต้องกระทำ ๕ ซ้ำ ก่อนอื่นเราจะต้องแบ่งเป็น ๕ Block, Block ละ ๓ Plots เพราะมีสิ่งเปรียบเทียบทั้งหมด ๓ Treatments เสร็จแล้วเราก็ใช้วิธีสุ่มภายใน Block ว่า Treatment ไหนจะอยู่ใน Plot ไหน ดังแผนผังต่อไปนี้

I			II			III			IV			V		
2	1	3	1	3	2	1	2	3	2	3	1	3	1	2

ส่วนการวาง Block นั้น ไม่จำเป็นจะต้องวางเรียงกันเป็นแถวเป็นแนวให้ตรงเปียงตรงตามทิศทางที่บริเวณที่ใช้ทดสอบมีรูปร่างไม่เป็นสี่เหลี่ยมหรือไม่สามารถจะวาง Block ให้เรียงเป็นแถวเป็นแนวได้ เราอาจแยกแต่ละ Block หรือหลาย ๆ Block ออกไปตามสภาพหรือรูปร่างของพื้นที่นั้น ๆ แต่จะไปแยก Plot ที่อยู่เรียงใน Block เดียวกันออกไม่ได้ เมื่อจะเอา Block ใดไปไว้ที่ใดก็ต้องยกไปทั้ง Block ส่วนเนื้อที่ของ Block และเนื้อที่ของแต่ละ Plot ก็ต้องคงเดิม ในกรณีที่จำเป็นต้องกระจาย Block นั้นก็จะ Random ด้วยว่า Block ที่ใดควรจะอยู่ตรงจุดไหน ดังตัวอย่างเช่น

ตัวอย่างวางแผนที่ ๓ ในการทดลองเปรียบเทียบพันธุ์พืช ๓ พันธุ์ โดยการกระทำ ๑๐ ซ้ำ (๑๐ Block) แต่ละ Block ใช้วิธีการ Random เพื่อบรรจุพันธุ์พืชทั้ง ๓ พันธุ์เข้าไว้พันธุ์ละ ๑ Plot ขนาดของ plot แต่ละ plot กว้าง ๒.๕๐ เมตร ยาว ๖.๐๐ เมตร

I			II			III			IV			V		
3	2	1	2	3	1	1	3	2	3	1	2	2	1	3
1	3	2	2	1	3	2	3	1	1	2	3	1	3	2
VI			VII			VIII			IX			X		

ในกรณีที่ไม่สามารถจะเรียง Block ดังกล่าวแล้วได้โดยเหตุที่พื้นที่มีรูปร่างไม่เป็นสี่เหลี่ยมหรือกว้างยาวพอตามขนาดของ Block ที่กำหนดไว้ได้ สมมติว่าบังเอิญมีบ่อน้ำขวางอยู่ เราอาจเอาแผนผังดังกล่าวแล้วเป็นหลัก เสร็จแล้วจึง Random เอา Block วางลงไป เมื่อ Random ตรงไหนถูก Block ใดก็นำ Block นั้นวางลงไป ส่วนการ Random พันธุ์ภายในแต่ละ Block ก็คงเดิม

IX			VII			III			X			VIII			VI		
1	2	3	2	1	3	1	3	2	1	3	2	2	3	1	1	3	2
3	1	2	3	2	1	2	1	3	2	3	1						
IV			I			V			II			อาคาร					

วิธีการสุ่มหรือ Randomization นั้น สามารถปฏิบัติได้ ๓ วิธี คือ

(๑) วิธีจับสลาก เป็นวิธีการที่ง่ายและธรรมดาที่สุด สมมุติว่าในหนึ่ง Block มี ๔ พันธุ์ เราก็เขียนเบอร์ในกระดาษ ๔ แผ่น ตั้งแต่เบอร์ ๑ ถึง ๔ เบอร์ ๆ ละ ๑ แผ่น ในขณะที่เดียวกันเราก็ให้เบอร์ประจำพันธุ์ทั้ง ๔ พันธุ์ให้ตรงกับเบอร์ในกระดาษ เสร็จแล้ว บวนกระดาษทั้ง ๔ เบอร์นั้นกลุ่กเกลากันแล้วหยิบขึ้นมาดูทีละใบ หยิบครั้งแรกตรงกับเบอร์ใด ก็รอกเบอร์นั้นลงในแปลงแรกของ Block นั้น กระทำจนครบ ๔ แปลง Block ที่สอง ก็กระทำเช่นเดียวกันจนครบทุก ๆ Block

(๒) วิธีใช้ตารางสุ่ม วิธีนี้นิยมใช้ในการเก็บตัวอย่างจำนวนหนึ่งโดยแบ่งส่วนออกจากสิ่งของจำนวนมาก เป็นต้นว่ามีข้าวโพดอยู่ ๕๐๐ ผัก ซึ่งเราต้องการวัดความยาวของ ผัก แต่เราต้องการจะได้ตัวอย่างมาวัดจริง ๆ เพียง ๑๐ ผัก ส่วนผลของการวัดนั้นเราใช้ แทนความยาวของข้าวโพด ๕๐๐ ผักได้ การที่จะดึงเอาตัวอย่างข้าวโพด ๑๐ ผักออกมาให้เป็นตัวแทนของข้าวโพด ๕๐๐ ผักนั้น เราสามารถใช้ตารางสุ่มได้ ดังตัวอย่างต่อไปนี้

เราจะได้เห็นว่า ตารางสุ่มแบ่งออกเป็น ๓ ภาค ภาคแรกมีเลขจำนวนละตัวเดียว ภาคที่สองมีเลขจำนวนละ ๓ ตัว ก่อนอื่นเราจะต้องดูว่า จำนวนของทั้งหมดของเราเป็นเท่าใด ในที่นี้เท่ากับ ๕๐๐ จึงใช้ตัวเลขในภาคที่ ๓ ต้องการดึงตัวอย่างออกมา ๑๐ ผัก ดูเลขใน Column ที่ (๑๐) ตรงลงมา มี ๒๑๒, ๖๑๔, ๓๘๒, ๔๕๔, ๐๘๑, ๘๐๗, ๐๔๕, ๒๑๗, ๒๓๓, ๖๗๕, ๓๓๑ ๔๑๑, ๑๑๗, ๘๕๗, ๓๐๕, ๓๖๒ ดูเลขที่มีค่าไม่เกิน ๕๐๐ ก็ออกมาตาม ลำดับจนครบ ๑๐ จำนวน ดังที่ได้ขีดเส้นใต้ไว้

(๓) วิธีใช้เครื่องสุ่ม มีผู้ประดิษฐ์เครื่องมือขนาดเล็กแบบกระเป๋าคือใช้มือหมุน เพื่อสุ่มในการเก็บตัวอย่างเช่นเดียวกับตารางสุ่ม เครื่องมือนี้ยังไม่มีแพร่หลายนัก

วิธีการสุ่มตามวิธีที่สองและที่สามนั้นใช้สำหรับเก็บตัวอย่างจำนวนน้อยจาก ประชากรจำนวนมาก ๆ แต่ในการวางแผนผังวิจัยเกี่ยวกับการเปรียบเทียบโดยทั่ว ๆ ไป ที่จำเป็นจะต้องแบ่ง Block เป็น Plot นั้น เราใช้วิธีที่หนึ่งโดยเฉพาะ

ในการจับฉลากนี้ ถ้าหากเราต้องการเปรียบเทียบพันธุ์พืช ๔ พันธุ์ โดยกระทำ ๔ Block ชั้นแรกตัดกระดาษสองชุด ๆ ละสี่แผ่น ชุดหนึ่งเขียนชื่อพันธุ์ A, B, C, D, อีกชุดหนึ่งเขียนเลข ๑, ๒, ๓, ๔ ในแผ่นนี้ เราใส่เลขประจำ plot ไว้ทุก plot เวลาจับก็จับพร้อม ๆ กันข้างละ ๑ ใบ สมมติว่าในคานพันธุ์จับได้ D และในคาน plot จับได้เลข ๒ ก็เอาพันธุ์ D ไว้ใน plot ที่ ๒ กระทำเช่นนั้นจนครบสี่ครั้ง หมกฉลากพอได้เป็นเสร็จ Block ที่ ๑ ส่วน Block ๒, ๓, ๔, ก็กระทำเช่นเดียวกัน

ใน Block หนึ่งๆ นั้นความจริงเราจับเพียง ๓ ครั้งก็พอแล้ว เพราะครั้งสุดท้ายคือเบอร์ที่เหลือ เป็นเบอร์ตายตัวซึ่งเรายอมทราบดี เพราะเหลือเพียงชุดละ ๑ ใบเท่านั้นยอมจะเปลี่ยนแปลงเป็นเบอร์อื่นไม่ได้ ซึ่งก็เข้าอยู่ในหลักเกณฑ์ของ Degree of Freedom

หลักที่ควรจำในการปฏิบัติเรื่องนี้ก็คือ

๑. จะจัดแปลงใน Block ที่ ๒ โดยลอกแบบจาก Block ที่ ๑ ไม่ได้
จำเป็นต้องสุ่มใหม่ทุกครั้งไป

๒. การเอาแผนผังการทดลองที่ใช้ในแห่งหนึ่งแล้วนำไปใช้ในที่อีกแห่งหนึ่งเป็นการไม่ดี หากจะกระทำการทดลองซ้ำในที่อื่น แม้ว่าจะมีจำนวนพันธุ์และจำนวน Block เท่าเดิมก็ควรที่จะสุ่มหรือจับฉลากเอาใหม่

สมุดตารางที่เขียนโดยศาสตราจารย์ทางวิชาสถิติสองคนคือ Fisher และ Yates มีตารางคำนวณทางสถิติหลายอย่าง และยังมีตารางเลขสำหรับสุ่มเพื่อใช้ในการวางแผนการทดลองด้วย ในตารางนั้นมีอยู่ ๒ หน้า ๆ ละ ๕ หน้ ๆ ละ ๕ Column ชั้นแรกจับฉลากเลือกหน้า ต่อไปก็จับฉลากเลือกหมู่ แล้วก็เลือก Column สมมติว่าได้นหน้า ๒ หมู่ ๓ Column ๔ มีเลข ๒๔, ๓๖, ๓๕, ๕๔, ๒๔ ตัวแรกเอา ๕ ตัวที่สองเอา ๔ ตัวที่สามเอา ๓ ตัว

$๒๔ \div ๕ = \text{เศษ } ๔ \text{ เลขประจำ } A \text{ ก็เท่ากับ } ๔$

$๓๖ \div ๔ = \text{เศษ } ๐ \text{ ซึ่งหมายความว่าไม่มี จึงนับให้เป็นจำนวนเต็มคือเศษ } ๔ \text{ ซึ่งไปตรงกับ } A \text{ จึงเพิ่มอีก } ๑ \text{ ให้ } B = ๕$

$$๗๕ \div ๓ = \text{เศษ } ๑ = C$$

$$๕๔ \div ๒ = \text{เศษ } ๐ \text{ อีกหรือจำนวนเต็มคือเศษ } ๒ = ๒ \text{ D แต่ C เป็น } ๑ \text{ อยู่แล้ว}$$

D จะต้องนับ ๒ เป็น ๑ เพราะฉะนั้น D ก็ตกอยู่ที่ ๓ ส่วน E เหลือสุดท้ายก็ต้องเป็น ๒

$$C = 1, E = 2, D = 3, A = 4, B = 5$$

ที่กล่าวมานี้เป็นวิธีการสุ่มอีกแบบหนึ่ง

การเปรียบเทียบความแตกต่างระหว่างสองสิ่ง

โดยเหตุที่แต่ละสิ่งต่างก็มีความคลาดเคลื่อนประจำลักษณะ และในการเปรียบเทียบเราก็ใช้ Mean ซึ่งเป็นตัวแทนของลักษณะนั้น ๆ มาเปรียบเทียบ แต่โดยเหตุที่เราต้องการผลการเปรียบเทียบที่แน่นอน จึงจำเป็นจะต้องนำเอามาตรฐานความคลาดเคลื่อนประจำ Mean เขามาเกี่ยวข้องกับ เพื่อต้องการจะทราบว่า Mean นั้น ๆ มีความแน่นอนเป็นที่เชื่อถือได้มากน้อยเพียงใด และความคลาดเคลื่อนของแต่ละ Mean นั้น จะทำให้ผลการเปรียบเทียบขอบเขตแน่นอนเป็นที่เชื่อถือได้แค่ไหน

การเปรียบเทียบความแตกต่างระหว่างสองสิ่งนี้สามารถกระทำได้สองแบบด้วยกันคือ

๑. โดยการวางแผนแบบสุ่มในแปลงหนู (Randomized Complete Block Design) ดังที่ได้อธิบายวิธีการวางแผนผังไว้แล้ว ถ้าประสงค์จะทำการทดลองที่ซ้ำก็ทำจำนวน Block ให้เท่ากับจำนวนใน ๑ Block ประกอบขึ้นด้วย ๒ แปลงเดี่ยว (Plot) เสร็จแล้วในแต่ละ Block ก็สุ่มวางภายใน ๒ Block นั้นถ้าเปรียบเทียบพันธุ์ พันธุ์ไหนจะอยู่ใน plot ใด

ตัวอย่าง มีพันธุ์พืชอยู่ ๒ พันธุ์ คือพันธุ์ A และพันธุ์ B ต้องการจะเปรียบเทียบผลผลิตทั้งสองพันธุ์ว่าพันธุ์ไหนจะให้ผลผลิตสูงกว่ากัน การวางแผนกระทำแบบสุ่มในแปลงหนู โดยการกระทำพันธุ์ละ ๑๐ ซ้ำ เมื่อเก็บผลผลิตแต่ละแปลงมาชั่งได้หนักก็คิดเป็นโลกรัม ดังนี้

พันธุ์ A

X_A	d_A	d_A^2
40	7	49
35	2	4
36	3	9
33	0	0
31	-2	4
32	-1	1
27	-5	25
24	-9	81
39	6	36
32	-1	1
330	0	210

พันธุ์ B

X_B	d_B	d_B^2
31	2	4
29	0	0
30	1	1
28	-1	1
27	-2	4
25	-4	16
27	-2	4
22	-7	49
37	8	64
34	5	25
290	0	168

$$M_A = \frac{330}{10} = 33.0$$

$$M_B = \frac{290}{10} = 29.0$$

เมื่อเราหา Mean ของแต่ละพันธุ์ได้แล้ว เราก็เอา Mean ของผลผลิตไปลบ ผลผลิตแต่ละแปลงในพันธุ์เดียวกันเพื่อจะหา d ซึ่งแสดงความคลาดเคลื่อนที่ผลผลิตแต่ละแปลง นั้นคลาดเคลื่อนไปจาก Mean ซึ่งใช้เป็นตัวแทนของผลผลิตพันธุ์นั้น ในช่อง d นี้เราจะพบ เครื่องหมายบวกลบเสมอไป เพราะ Mean จะตกอยู่ตรงกลางเนื่องจากเป็นผลเฉลี่ย ถ้าคิด เลขถูกต้องแล้วผลรวมทางบวกและลบจะเท่ากัน เมื่อรวมเป็นยอดใหญ่ของ d จะเท่ากับ 0 เสร็จแล้วจึงเอา d แต่ละกานายกกำลังสอง

สรุปแล้วให้กระทำการหา S.E._M ของแต่ละพันธุ์ทั้งพันธุ์ A และพันธุ์ B ตาม หลักเกณฑ์ที่ได้อธิบายและยกตัวอย่างไว้แล้วตามตัวอย่างดังต่อไปนี้

$$S.E._{M_A} = \pm \sqrt{\frac{\sum d_A^2}{n(n-1)}} = \pm \sqrt{\frac{210}{10(10-1)}} = \pm 1.528$$

$$S.E._{M_B} = \pm \sqrt{\frac{\sum d_B^2}{n(n-1)}} = \pm \sqrt{\frac{168}{10(10-1)}} = \pm 1.366$$

$$\text{หรือถ้าจะใช้สูตร } S.E._M = \pm \sqrt{\frac{\sum X^2 - (\sum X)^2/n}{n(n-1)}} \text{ ก็จะได้ผลลัพธ์เท่ากัน}$$

$$M_A \pm S.E.M_A = 33.0 \pm 1.528$$

$$M_B \pm S.E.M_B = 29.0 \pm 1.366$$

เราได้ผลเฉลี่ยของพันธุ์ A และพันธุ์ B แล้ว แต่ผลเฉลี่ยนี้มีมาตรฐานความคลาดเคลื่อน คืออยู่ควยทั้งสองพันธุ์ ในการเปรียบเทียบเพื่อต้องการหาความแตกต่างระหว่างผลผลิตจึงกระทำ โดยเทียบผลเฉลี่ยเฉย ๆ ไม่ได้ เพราะผลเฉลี่ยแต่ละอันก็แสดงว่าสามารถจะเคลื่อนไปได้ไม่ คงที่ ฉะนั้นผลต่างที่ได้จากการเปรียบเทียบธรรมดา คือ

$$D = M_A - M_B$$

= 33.0 - 29.0 = 4.0 จึงเป็นความแตกต่างที่ยังไม่สามารถจะลงความเห็น แน่ชัดได้ว่าพันธุ์ A ให้ผลผลิตสูงกว่าพันธุ์ B 4.0 ก.ก. นั้น A จะดีกว่า B จริง หรือไม่จึงจำเป็นต้องหาความคลาดเคลื่อนของทั้งสองพันธุ์นี้มารวมกันเพื่อจะวัดดูว่าความคลาดเคลื่อนนั้นจะใหญ่เกินความแตกต่างหรือไม่ โดยการคำนวณดังนี้

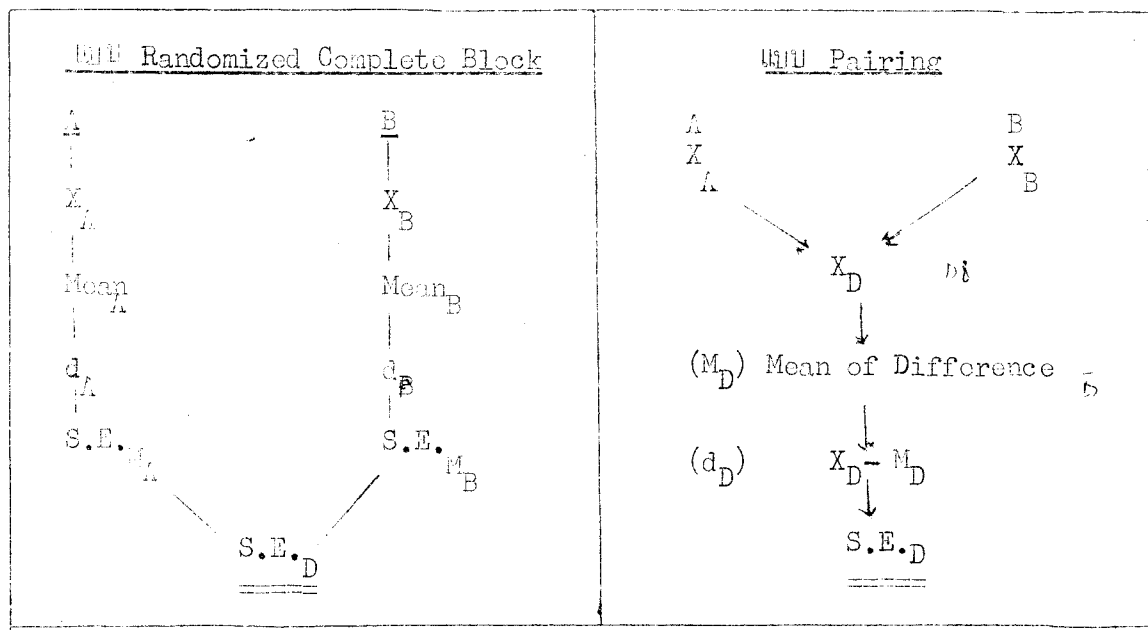
$$\begin{aligned} S.E._D \text{ (มาตรฐานความคลาดเคลื่อนของผลต่าง)} &= \sqrt{S.E._M_A^2 + S.E._M_B^2} \\ &= \sqrt{1.528^2 + 1.366^2} \\ &= 2.050 \end{aligned}$$

เราก็ได้ผลต่างที่สมบูรณ์คือ $D \pm S.E._D = 4.00 \pm 2.050$

ต่อไปเราคำนวณ "t" ratio ของผลการทดลองนี้ซึ่งได้จากการเทียบอัตราส่วนระหว่าง ความแตกต่างระหว่างสองสิ่งนี้กับความคลาดเคลื่อนของความคลาดเคลื่อนแตกต่าง ทั้งนี้เพื่อ ต้องการทราบว่า ความแตกต่างทั้งหมดนั้นเป็นกี่เท่าของความคลาดเคลื่อนอันเกิดจากสาเหตุ อื่น ๆ ที่บารบกวนผลการทดลอง ถ้าหากสิ่งที่เกิดจากการรบกวนของสาเหตุภายนอกใหญ่กว่า ความแตกต่างทั้งหมด ก็แสดงว่าความแตกต่างที่เกิดจากตัวเองซึ่งเราต้องการจะเปรียบเทียบ นั้นไม่เป็นที่เชื่อถือได้ ดังนั้น

กระทั้งการเปรียบเทียบชั้นสุดท้ายก็คิดเพี้ยนไปจากวิธีการที่กล่าวมาข้างบนแล้ว เป็นกรณีการเปรียบเทียบความเจริญระหว่างพืชมัน้อง ซึ่งตามเหตุผลแล้วพืชน้องก็ควรจะเปรียบเทียบกันเองเป็นคู่ ๆ แม้จะทำหลาย ๆ ซ้ำก็เท่ากับทำหลาย ๆ คู่ ในการเปรียบเทียบขั้น ๒ พืชในบริเวณที่ดินซึ่งมีการเปลี่ยนแปลงความอุดมสมบูรณ์ไม่มากนัก บางที่เราให้เหตุผลว่า การวางแผนจัดกันเป็นคู่ ๆ จะช่วยให้พืชทั้งสองพันธุ์ได้รับความอุดมสมบูรณ์ของดินเหมือนกัน แต่ในโอกาสเช่นนี้ก็ไม่ค่อยนิยมใช้กันในปัจจุบัน เพราะถือกันว่าวิธีการแรกเหมาะกว่า โดยให้เหตุผลว่าพืชทั้งสองพันธุ์จะได้มีอิสระในตัวเองได้เต็มที่ เมื่อได้มีการวางแผนแบบจับคู่แล้ววิธีการคำนวณก็จำเป็นต้องกระทำแบบจับคู่ด้วย กล่าวคือ แทนที่จะคำนวณหา Mean ของ A และ B โดยแยกกันหา $S.E.M_A$ กับ $S.E.M_B$ โดยอิสระ เสร็จแล้วจึงหา D และ $S.E.D$ ที่หลัง แต่การคำนวณจับคู่เราคำนวณหาผลต่างในระยะเริ่มต้น เสร็จแล้วเมื่อกำหนดหา Mean ก็เป็น D เดียวทีเดียว ส่วนการหา deviation "d" ก็เป็นของผลต่างตลอดจนการคำนวณหา $S.E.M$ ที่ได้ก็เป็นของ $S.E.D$ เดียว ฉะนั้นการคำนวณจึงสั้นกว่าวิธีก่อน

แผนผังเปรียบเทียบวิธีการคำนวณเพื่อเปรียบเทียบสองสิ่ง



ตัวอย่าง Pairing Method

ใช้ตัวเลขจากตัวอย่างที่แล้ว

พันธุ์ A	พันธุ์ B	หรือ X_D	$X_D - M_D$ (d_D)	d_D^2
40	31	9	5	25
35	29	6	2	4
36	30	6	2	4
33	28	5	1	1
31	27	4	0	0
32	25	7	3	9
28	27	1	-3	9
24	22	2	-2	4
39	37	2	-2	4
32	34	-2	-6	36
		40	0	96

$$D \text{ or } M_D = \frac{40}{10} = 4.00$$

$$S.E._D = \pm \sqrt{\frac{\sum d_D^2}{n(n-1)}} \text{ or } \sqrt{\frac{\sum X_D^2 - (\sum X_D)^2/n}{n(n-1)}}$$

(Note : n = number of pair)

$$= \pm \sqrt{\frac{96}{10(10-1)}} = 1.033$$

$$\text{Calculated "t" ratio} = \frac{4.00}{1.033} = 3.872 \text{ ++ Highly Significant}$$

ในการเปิดตารางมาตรฐาน t ผิดกับวิธีอื่น เนื่องจากวิธีตัวเลข
จำเป็นต้องลบกันเป็นคู่ ๆ ตามแผนผังการทดลอง จะตัดกันไม่ได้ ฉะนั้นจำนวนตัวเลข
จึงเสียความอิสระไปบาง การดู D.F. ของตาราง t จึงใช้จำนวนคู่ตัวหนึ่ง
หรือ $n = 1$

$$(n = \text{จำนวนคู่}) \quad D.F. = 10 - 1 = 9$$

$$\text{เปิดตาราง "t" ระดับ 1:95 หรือ } t_{5\%} = 2.262$$

$$\text{" 1:99 " } t_{1\%} = 3.250$$

ในการทดลองโดยวางแผนแบบจับคู่นี้ปรากฏว่าอัตราส่วน "๗" ที่คำนวณได้จากการทดลองมีค่าเท่ากับ ๓.๘๗๒ $\pm$ ๐.๒๕๐ ซึ่งป็นระดับ t_{17} จึงสามารถสรุปผลการทดลองได้ว่าพันธุ์ A ให้ผลผลิตสูงกว่าพันธุ์ B จริง

ในกรณีที่ค่าของ "๗" ที่คำนวณได้สูงกว่าระดับ ๑ % นิยมใช้เครื่องหมายดอกจัน ๒ ดอกใส่ไว้ที่ค่าของ "๗" นี้ถ้าสูงกว่าระดับ $t_{5\%}$ แต่ยังต่ำกว่าระดับ $t_{1\%}$ ก็ใช้ดอกจันหนึ่งเพียงดอกเดียว แต่หากต่ำกว่าระดับ $t_{5\%}$ คือสรุปได้ว่าไม่มีความแตกต่างที่จะถือเป็นสำคัญได้ ก็ไม่ใส่เครื่องหมายดอกจันเลย

ANALYSIS OF VARIANCE

วิธีการเปรียบเทียบเพียงสองชนิดที่ได้อธิบายมาแล้วนั้นยังนับว่าเป็นวิธีการที่ให้ผลน้อยเกินไป ไม่พอแก่ความต้องการของนักก้นหาทดลอง เพราะในการค้นหาทดลองแต่ละครั้งแต่ละคราวจำเป็นต้องลงทุนลงแรง เสียเวลา ยิ่งเกี่ยวข้องกับพืชหรือสัตว์ก็ยิ่งจำเป็นต้องใช้เวลามาก ฉะนั้นในการค้นหาทดลองครั้งหนึ่ง ๆ ก็ควรจะกระทำให้ครบถ้วน เป็นต้นว่าในการเปรียบเทียบพันธุ์พืช ก็ควรจะเปรียบเทียบครั้งละหลาย ๆ พันธุ์ หรือในการเปรียบเทียบส่วนผสมของอาหารกับการเจริญเติบโตของสัตว์ ก็ควรมีอาหารหลาย ๆ อัตราส่วน มิใช่กระทำเพียงสองพันธุ์หรือสองส่วน แต่ก็ไม่ได้หมายความว่าต้องกระทำหลาย ๆ สิ่งหลาย ๆ อย่างจนกระทั่งมากเกินไป ทำให้ยุ่งเหยิงไม่ทั่วถึง ซึ่งอาจเป็นเหตุให้ผลงานเสียไปก็ได้ จึงควรจะประมาณปริมาณของงานให้พอเหมาะแก่ความสามารถที่จะรับงานนั้น ๆ ได้ผลดี ในการเปรียบเทียบพร้อม ๆ กันมากกว่าสองสิ่งขึ้นไปนั้น เราอาจจะทำตามวิธีเดิมที่ได้กล่าวมาแล้วได้โดยการหา Mean ของแต่ละสิ่งแต่ละพันธุ์ เสร็จแล้วเราก็จะต้องหา S.E._M ประจำ Mean ของแต่ละพันธุ์นั้นเป็นราย ๆ ไปทุกพันธุ์ เมื่อต้องการจะเปรียบเทียบคู่ใดก็หาผลต่าง (D) ระหว่าง Mean คู่หนึ่งและหา S.E._D เป็นรายคู่ไปทุก ๆ คู่ แต่วิธีนี้เป็นวิธีที่ยาวและไม่สะดวก ถ้าหากเราจำเป็นต้องเปรียบเทียบกันถึง ๑๐ พันธุ์ ๒๐ พันธุ์ การจับคู่ให้ครบทุกคู่ก็จะยิ่งยืดยาวและเสียเวลามาก ทั้งผลงานก็ไม่มีความละเอียดละออ เพราะ S.E. นั้นเป็นเครื่องแสดงความ

คลาดเคลื่อนอันเกิดมาจากมูลกรณี่หลายประการ เช่นความคลาดเคลื่อนอันเกิดจากระดับความ
อุดมสมบูรณ์ของดินหรือเครื่องปลูก การรบกวนของศัตรู ความผันแปรของฝนฟ้าอากาศตาม
ธรรมชาติ.....ซึ่งมูลกรณี่เหล่านี้มีทั้งที่เราควบคุมได้บ้างไม่ได้บ้าง ที่ควบคุมไม่ได้นั้นก็ทั้งที่
เรารู้สาเหตุและไม่รู้สาเหตุ แต่ S.E. ที่เราหามาแล้วในคอนแรกเป็น S.E. ที่รวมเอาความ
คลาดเคลื่อนอันเกิดจากกรณี่หลายประการไว้ทั้งหมดโดยมิได้แบ่งแยกแสดงรายละเอียด เราจึง

เรียกว่า Experimental Error หรือเรียกว่า S.E. ของ Total (S.E._T)
ที่จะบอกความรายละเอียดอยู่ภายในคือ ความคลาดเคลื่อนที่เกี่ยวกับพันธุ์หรือชนิดของปุ๋ยที่เรา
นำมาเปรียบเทียบกับ หรือ S.E. ของ Variety (S.E._V) ความคลาดเคลื่อนอันเกิดจาก
ระดับความอุดมสมบูรณ์ของดินในแปลงทดลองแต่ละแห่งหรือ S.E. ของ Block (S.E._B)
และยังมี S.E. ของมูลกรณี่อื่น ๆ ที่เราไม่ทราบหรือควบคุมไม่ได้อีกมากมาย เมื่อเราต้อง
การเขียนเป็นสูตรก็จะได้ ดังนี้

$$S.E._T = S.E._V + S.E._B + \dots\dots\dots + S.E._{\text{อื่น ๆ}}$$

ใน S.E. ที่รวมกันหลาย ๆ อย่างนี้ เมื่อเราแยกเอา S.E._V และ S.E._B
ซึ่งเราสามารถจะหาได้ด้วยการคำนวณจากปรากฏการณ์ของสิ่งนั้นออกแล้ว สิ่งที่เหลือเรา
เรียกว่า Remainder (S.E._R) ซึ่งเป็น S.E. ของการทดลองนั้น ๆ โดยแท้
อาจเกิดจากมูลกรณี่อื่นใดที่เราไม่สามารถจะแยกออกได้ เป็นต้นว่าเกิดจากโชค (Chance)
 เป็นต้น แต่ก็ถึงกระนั้นเราก็สามารถที่จะคำนวณหาค่าของความคลาดเคลื่อนที่เหลือนี้ได้โดย
เหตุผลที่ว่า เราหาความคลาดเคลื่อนทั้งหมดได้ หาความคลาดเคลื่อนอันเกิดจากสาเหตุ
ที่เราวัดได้ เพราะฉะนั้นเราก็สามารถจะหาความคลาดเคลื่อนที่เหลือได้โดยนำเอาสิ่งที่
เรารู้ได้นั้นไปหักออกจากความคลาดเคลื่อนทั้งหมด ซึ่งเราสามารถจะเขียนเป็นสูตรได้
ดังนี้

$$S.E._T = S.E._V + S.E._B + (\dots\dots\dots + S.E._R)$$

$$\text{หรือ } S.E._R = S.E._T - (S.E._V + S.E._B)$$

นักถวณแรกที่ได้ใช้วิธีการนี้เป็นนักศึกษา ชื่อ W.S. Gosset ใช้นามปากกาว่า Student ได้เขียนบทความตามความคิดเห็นเกี่ยวกับเรื่องนี้ หลังจากนั้นก็ได้มีผู้ศึกษาตามไปอีกมาก ในปี ค.ศ. ๑๙๒๕ มีอาจารย์ชาวอังกฤษ ชื่อ R.A. Fisher ได้เริ่มศึกษาวิธีใหม่นี้โดยละเอียดและเรียกรวธีนี้ว่า Variance Analysis

เราเคยพบมาจากเบื้องต้นว่า ก่อนที่เราจะหา S.E. ได้เราจะต้องถอด Square root เสียก่อน แลหารหา Variance ต้องเอา S.E. มากกำลังสอง เพราะฉะนั้นถ้าเราเอาผลรวมของ Deviation กำลังสอง หารด้วย degree of freedom แต่ไม่ถอด square root เราก็จะได้ค่า Variance เลยทีเดียว ถ้าเขียนเป็นสูตรก็จะได่ ดังนี้

$$S.E.^2 = \text{Variance}$$

$$S.E. = \sqrt{\frac{\sum d^2}{n-1}} \quad \text{หรือ} \quad \sqrt{\frac{\sum X^2 - (\sum X)^2/n}{n-1}} \quad \text{หรือ} \quad \sqrt{\frac{\text{some of square}}{\text{degree of freedom}}}$$

$$\text{Variance} = \frac{\sum d^2}{n-1} = \frac{\sum X^2 - (\sum X)^2/n}{n-1} = \frac{S.S.}{D.F.}$$

วิธีการ Analysis of Variance นี้มีหลักสำคัญอยู่อย่างหนึ่ง คือหลักการสุ่ม (Random arrangement) คือไม่มีการลำเอียงหรือเลือกที่รักมักที่ชัง หากแต่ใช้วิธีสุ่มลงไปตามแต่จะถูกที่ใดก็ช่าง หลักการนี้มาจากหลักของ Probability หรือเป็นทฤษฎีของโคลด (Law of chance) ซึ่งมี curve เหมือนระฆังจึงใช้วิธีสุ่มเพื่อให้การค้นคว้าทดลองนั้นใกล้เคียงความจริงที่สุด ยิ่งทำถึง infinity ก็จะได้ curve ที่สมบูรณ์ แต่ในทางปฏิบัติของเราไม่สามารถจะกระทำมากมายเช่นนั้นได้ วิธีการสุ่มโดยอาศัยโคลดจึงเป็นวิธีการที่เหมาะสมที่สุดเพื่อจะทำให้ได้ตัวแทนที่สัดส่วนของการแปรปรวนใกล้เคียงกับประชากรที่แท้จริง ส่วนการวางแผนผังทดลองก็ต้องวางแผนสุ่ม คือมีแปลงหมู่หรือ Block ตามจำนวนซ้ำ แล้วใช้วิธี Random arrangement ภายในแต่ละแปลงภู่นั้น ฉะนั้นจะเลือกเอาพันธุ์ไหนหรือ Treatment ใดไว้ตามใจชอบไม่ได้ จำเป็นจะต้องใช้วิธีสุ่ม เพื่อเนี่ยลการเปลี่ยนแปลงความอุดมสมบูรณ์ของดินหรือสิ่งแวดล้อมอื่น ๆ ที่เกิดประจำที่ในทุก ๆ สิ่งทุก ๆ อย่างที่เรา

จะนำมาเปรียบเทียบกันนั้นได้มีโอกาสที่จะได้รับโดยทั่วถึงและใกล้เคียงความจริง ดังโดยยกตัวอย่างอธิบายไว้แล้วในหัวข้อการสุ่มเป็นแปลงหมู่ (Randomized Complete Block System)

วิธีวางแผนยังวิจัยโดยอาศัยหลักการ Analysis of Variance นั้น ประกอบด้วยวิธีการเป็นหลักใหญ่ ๆ และเป็นรากฐานสำคัญอยู่ ๒ อย่าง คือ

๑. Randomized Complete Block Design หรือเรียกสั้น ๆ ว่า

Randomized Block Design

๒. Latin Square Design

RANDOMIZED BLOCK DESIGN

เป็นวิธีการหลักหรือเป็นรากฐานอันสำคัญที่สุดซึ่งสามารถนำไปสู่การศึกษาวิธีการวางแผนเบื้องต้น ในการทดลองจนกว่าใด ๆ ก็ตาม มักจะมีวิธีการนี้เข้าไปเกี่ยวข้องกับ เพราะเป็นวิธีการที่รัดกุมและทั่วไถ่ยอมรับว่าไม่มีวิธีการใดที่จะเหมาะและให้ผลดีเท่าวิธีการนี้ ในการวางแผนการวิจัยเบื้องต้น ก็มักจะหนีวิธีการนี้ไปไม่พ้น และมักจะใช้เป็นหลักสำคัญในการที่จะคิดดัดแปลงวิธีการวางแผนให้ได้นอกกว้างขวางออกไปอีก สรุปแล้วผู้ที่ศึกษาวิชาการแขนงนี้จึงจำเป็นอย่างยิ่งที่จะต้องมีความชำนาญและเข้าใจซาบซึ้งในวิธีการนี้เป็นอย่างดี

วิธีการ Randomized Block คือการใช้ Block เป็นหน่วยสมบูรณ์ ดังที่ได้ อธิบาย และยกตัวอย่างมาครั้งหนึ่งแล้วในหัวข้อ วิธีการสุ่มเป็นแปลงหมู่ สมมติว่าเรามีพันธุ์พืชที่จะเปรียบเทียบกัน ๕ พันธุ์ หรือปุ๋ยที่จะทดลองเปรียบเทียบกัน ๕ ชนิด ทำ ๒ ซ้ำ (๒ Rep. หรือ ๒ Block) ใน Block หนึ่ง ๆ ก็จำเป็นต้องมี ๕ plot, plot ละ ๑ พันธุ์ หรือถ้าเป็นปุ๋ยก็ ๑ ชนิด ภายใน Block หนึ่ง ๆ เป็นหน่วยหนึ่งของการทดลองที่มีจำนวนพันธุ์หรือปุ๋ยสมบูรณ์ครบถ้วนไม่ซ้ำกันและไม่ขาดไป การเรียง Block ถัดข้างไม่สามารถจะจัดเรียงให้เป็นระเบียบได้ ก็จัด Block หนึ่ง ขวางบ้าง หรืออาจแยก Block หนึ่ง ไว้ที่หนึ่ง และอีก ๓ Block ไว้อีกแห่งหนึ่งในบริเวณใกล้เคียงกันสามารถดูแลปฏิบัติให้เป็นเหมือนกัน

โดยทั่วถึงได้ แต่ขนาด Block และขนาด plot จะต้องคงเดิมไม่เปลี่ยนแปลง และจะแยก plot ที่อยู่ใน Block เดียวกันออกจากกันไม่ได้ สิ่งเหล่านี้ได้อธิบายมาบ้างแล้วพอสมควร ในบทที่ว่าด้วยเรื่องการสุ่ม

ตัวอย่าง สมมติว่าเราทำการทดลองเปรียบเทียบพันธุ์พืช ๔ พันธุ์ คือพันธุ์ A, B, C, และ D วางแผนการทดลองแบบ Randomized Block โดยกระทำ ๕ ซ้ำ (5 Replications) ดังนั้นจึงจำเป็นต้องกระทำ ๕ Block, Block ละ ๔ plot ใน ๑ Block ก็มีพันธุ์ A, B, C, D, ครบทุกพันธุ์ ๆ ละ ๑ plot การจัด plot ใน Block แต่ละ Block เพื่อวาง A, B, C, D, ใช้วิธีการสุ่มทุก Block ดังแผนผังต่อไปนี้

Block I				Block II							
C	A	B	D	A	D	B	C	บอกกันไว้ไว้			
16	18	25	23	16	23	28	18				
B	D	C	A	C	A	D	B	D	A	B	C
24	25	17	19	19	21	24	24	22	20	25	18
Block III				Block IV				Block V			

เมื่อเก็บผลผลิตของแต่ละแปลงเรียบร้อยแล้ว ขั้นแรกก็ได้นำที่กลงในแผนผังประจำเป็นแปลงๆ ไปตั้งเห็นปรากฏตัวเลขอยู่ในแผนผังแล้ว ขั้นต่อมาก็ทำการจัดระเบียบใหม่เพื่อนำเอาตัวเลขกรอกลงในตารางเพื่อให้สามารถแสดงยอดของผลผลิตแต่ละพันธุ์ และแสดงยอด Block แต่ละ Block ซึ่งเป็นการแสดงระดับความอุดมสมบูรณ์ของดินหรือความผิดปกติในเรื่องสภาพอันเกิดจากพื้นที่ การจัดตารางจึงกระทำดังต่อไปนี้ ซึ่งเป็นตารางที่แปรมาจากแผนผังในไรทดลอง

Block	Varieties			
	A	B	C	D
I	18	25	16	23
II	16	23	18	23
III	19	24	17	25
IV	21	24	19	24
V	20	25	18	22

การที่จัดรวมเป็นพันธุ์ก็เพราะจะต้องมีการเปรียบเทียบพันธุ์และการวัดความคลาดเคลื่อนอันแสดง
 ออกจากพันธุ์ ส่วนการจัดเข้าเป็น Block ก็เพื่อจะหาความผิดเพี้ยนอันเกิดจากพื้นที่ เพราะการ
 กระจายตัวของแต่ละพันธุ์ออกไปตาม Block ก็เพื่อจะให้ทุก ๆ พันธุ์มีโอกาสสูงความผิดเพี้ยน
 ที่เกิดขึ้นโดยเฉลี่ยกันไป ในการคำนวณหาความคลาดเคลื่อนทั้งหมดและหาความคลาดเคลื่อน
 ปลีกย่อยต่าง ๆ เช่นเรื่องเกี่ยวกับพันธุ์ก็ เราจะคำนวณหาผลรวมของความคลาดเคลื่อนนำดัง
 สอง (Sum of deviation) หรือบางทีก็เรียกว่า Sum of Square ซึ่งใช้สูตรว่า $\sum d^2$
 หรือ $\sum X^2 - C.F.$ ของทุก ๆ รายการ แล้วจึงจะหารด้วย degree of freedom ทุก
 รายการพร้อมกันหมดในตาราง Analysis of Variance ก็จะได้ Variance ของราย
 การนั้น ๆ ดังนั้นจึงขอกล่าวแต่เฉพาะวิธีหา Sum of Square ก่อน

ในเรื่องของ Randomized Block Design นี้เราจำเป็นจะต้องหาความคลาดเคลื่อนใหญ่
 ของการทดลองเสียก่อน แล้วจึงหาความแตกต่างที่แสดงออกจากพันธุ์และความผิดเพี้ยนที่เกิด
 จากพื้นที่ เมื่อหาตามรายการนี้เสร็จแล้วจึงนำไปหักออกจากยอดใหญ่ของการทดลองทั้งหมด
 สิ่งที่เหลือจากการหักออกแล้วเป็นสิ่งที่เราไม่สามารถจะแยกแยะออกเป็นรายการปลีย่อยได้อีก
 จึงรวมเรียกกันว่า สิ่งที่เหลือ (Remainder) ซึ่งเราถือว่าเกิดจากสาเหตุอื่นๆ เล็ก ๆ
 น้อย ๆ อีกมากที่เราไม่รู้ หรือบางทีบางตัวก็ควบคุมไม่ได้ หากจะควบคุมได้ก็ต้องหมดเปลือง
 ทุนรอนมากมายแทนคิไม่คุ้ม เพราะเป็นเรื่องเล็กน้อยเหลือเกิน จึงรวมนามาคิดเป็น
 "สิ่งเหลือ" อันเดียวกันหมด การหา Sum of Square สำหรับทั้งหมด (S.S. Total)

เป็นการหาความคลาดเคลื่อนของตัวเลขทุกจำนวนในการทดลองไม่ว่าจะอยู่ในรายการของอะไรก็ตาม
 แต่องค์เป็นตัวเลขจำนวนน้อยที่สุดของการทดลอง จากความรู้ที่ได้เรียนมาแล้วในตอนนั้น เรา
 อาจทำได้ ๒ วิธีด้วยกันคือ $\sum d^2$ หรือ $\sum X^2 - C.F.$ ก็ได้ ทั้งจะได้เปรียบเทียบวิธีการ
 และผลให้ต่อไปนี้ ทั้งสองวิธีทางก็จำเป็นต้องใช้ตัวเลขน้อยที่สุดของการทดลองทั้งหมดซึ่งเมื่อ
 ตรวจสอบดูแล้วเห็นว่ามีอยู่ทั้งหมด ๒๐ จำนวน เพราะเป็นการหาของทั้งหมด
 ส่วนวิธีการมีดังต่อไปนี้

X	$(X-M)$	d^2
16	-5.0	25.0
18	-3.0	9.0
25	4.0	16.0
23	2.0	4.0
16	-5.0	25.0
23	2.0	4.0
23	2.0	4.0
18	-3.0	9.0
24	3.0	9.0
25	4.0	16.0
17	-4.0	16.0
19	-2.0	4.0
19	-2.0	4.0
21	0.0	0.0
24	3.0	9.0
24	3.0	9.0
20	-1.0	1.0
25	4.0	16.0
18	-3.0	9.0
22	1.0	1.0
420	0.0	190.0

X	X^2
16	256
18	324
25	625
23	529
16	256
23	529
23	529
18	324
24	576
25	625
17	289
19	361
19	361
21	441
24	576
24	576
20	400
25	625
18	324
22	484
420	901.0

$$M = \frac{420}{20} = 21.0$$

$$S.S._{Total} = 190.0 \quad \text{or} = \text{by using X formular}$$

$$= \sum X^2 - (\sum X)^2/n = 9010 - 420^2/20 = 190.0$$

การหา Sum of Square สำหรับพันธุ์ (S.S. Varieties) เพื่อจะวัดความแตกต่างที่แสดงออกโดยพันธุ์ หรือถ้าการทดลองนั้น ๆ เป็นการทดลองเปรียบเทียบปุ๋ยหรือเปรียบเทียบวิธีการปลูกปฏิบัติ

การหา S.S. ในที่นี้เป็นการวัดความแตกต่างซึ่งแสดงออกโดย Treatment ต่าง ๆ กัน แต่ก็มีไทม์หมายควาว่าพันธุ์แต่ละพันธุ์หรือ Treatment แต่ละ Treatment แสดงผลผลิตแตกต่างกันมาก ๆ แล้ว จะถือว่าไทม์แตกต่างกันจริง เพราะความต่างนั้นอาจเป็นผลกระทบกระเทือนมาจากสาเหตุอื่น ๆ ตามธรรมชาติซึ่งเราไม่สามารถจะรู้ได้ก็เป็นได้ ดังนั้นในขั้นแรกนี้เราก็เพียงแต่ทำการคำนวณความแตกต่างทั้งหมดที่แสดงออกโดยพันธุ์ หรือ Treatment เสียก่อน ในตอนหลังเมื่อเราสามารถคำนวณหาความคลาดเคลื่อนที่เกิดจากสาเหตุอื่น ๆ แล้ว จึงจะนำมาเทียบกับค่าสิ่งที่เกิดจากพันธุ์นั้นจะเชื่อถือได้หรือไม่ ต่อไปนี้เป็น การคำนวณขอบเขตของความแตกต่างในเรื่องพันธุ์ เพื่อหาว่าผลผลิตของแต่ละพันธุ์นั้นมีขอบเขตของความแตกต่างกันกว้างเพียงไร จึงจำเป็นต้องเอาผลผลิตของแต่ละพันธุ์ คือ A B C และ D มาเรียงกันแล้วคำนวณหา sum of deviation square ซึ่งมีวิธีทำได้สองวิธี เช่นเดียวกับการคำนวณหา sum of deviation square

การหา S.S. varieties แบบใช้สูตร $\sum d^2$

พันธุ์						X _{Var.}	d	d ²
A	18	16	19	21	20	94	-11	121
B	25	23	24	24	25	121	16	256
C	16	18	17	19	18	88	-17	289
D	23	23	25	24	22	117	12	144
						420	$\sum d = 0$	$\sum d^2 = 810$

รวมทั้ง 4 varieties = 420

ผลเฉลี่ยต่อ 1 variety $M = \frac{420}{4} = 105$

แต่เนื่องจากค่าของ X ของแต่ละพันธุ์เป็นค่าของผลรวมพันธุ์ & ถ้า มีค่าจากหน่วยเดียว ดังนั้นจึงต้องการ $\sum d^2$ ด้วย & เพื่อเฉลี่ยให้เป็นหน่วยเดียวเสียก่อน

เนื่องจาก S.S. Total ที่เรากำหนดไว้แล้วเป็นความคลาดเคลื่อนของตัวเลขแต่ละจำนวน ถ้าเราไม่ทำ S.S. อื่น ๆ ให้เป็นหน่วยเดียวก็ยากที่จะไปสามารถนำไปใช้ร่วมกันได้ จึงต้องทำ S.S. var. ให้เป็นหน่วยเดียวโดยหารด้วยจำนวนซ้ำหรือ ๕ ดังนั้น

$$\frac{\sum x_{var.}^2}{5} = \frac{810}{5} = 162$$

การหา S.S. varieties แบบใช้สูตร $\sum X^2 - C.F.$

พันธุ์						$X_{var.}$	$X_{var.}^2$
A	18	16	19	21	20	94	8836
B	25	23	24	24	25	121	14641
C	16	18	17	19	18	88	7744
D	23	23	25	24	22	117	13689
						420	44910

$$C.F. = \left(\sum X \right)^2 / n = 420^2 / 20 = \frac{176400}{20} = 8820$$

สำหรับ $\sum X^2$ นั้น เนื่องจาก X แต่ละตัว คือ ๑๘, ๑๖, ๑๙, ๒๑, ๒๐ ที่นำมายกกำลังสองนั้น มีใช้ X เดียว ๆ แต่ที่จริงแต่ละตัวมีค่ามาจาก 5X เช่น ๑๘ มาจาก ๑๘ + ๑๘ + ๑๘ + ๑๘ + ๑๘ ดังนั้น $\sum X^2$ ที่คำนวณได้จึงจำเป็นต้องหารด้วย ๕ เพื่อให้เป็นขอบเขตของความผิดพลาดระหว่างเลขเดียว ๕ จำนวนโดยแท้ หรืออาจกล่าวได้ว่าเป็นการวัดความผิดพลาดที่เกิดจากพันธุ์ซึ่งเฉลี่ยต่อ ๑ Rep. ส่วนค่าของ C.F. นั้นเป็นค่าเดียวอยู่แล้วเพราะเฉลี่ยด้วยจำนวนตัวเลขทั้งหมด คือ ๒๐

$$S.S. var. = \frac{\sum x_{var.}^2}{n_{block}} - C.F. = \frac{44910}{5} - 8820 = 162.0$$

การที่ตัวเลขในทำนองเดียวกันถึง ๕ ชุด และในการหารด้วย ๕ เพื่อเฉลี่ยให้เป็นความผิดพลาดในระหว่างตัวเลขชุดเดียวกัน (๕ จำนวน) นั้น เราจะหารเฉลี่ยตัวเลขแต่ละจำนวนก่อนหา S.S. ไม่ได้เพราะตัวเลขย่อยทั้ง ๕ จำนวนที่ประกอบเป็นยอดของแต่ละพันธุ์นั้นต่างก็มีความคลาดเคลื่อนในตัวเอง ถ้าเราทำการเฉลี่ยเสียก่อน ความคลาดเคลื่อนเหล่านั้นก็จะ

ไม่แสดงออกมา คงแสดงเฉพาะความคลาดเคลื่อนของผลเฉลี่ยทั้ง ๔ จำนวนเท่านั้นเอง ดังนั้น การเฉลี่ยความคลาดเคลื่อนให้ทั้งหมดจึงรวมเอาความคลาดเคลื่อนระหว่างจำนวนตัวเลขย่อยแต่ละจำนวน กับความคลาดเคลื่อนซึ่งเกิดจากความแตกต่างระหว่างพันธุ์ ซึ่งแสดงออกโดยยอดผลผลิตแต่ละพันธุ์เข้าด้วยกัน $S.S._{var.}$ ที่ได้จึงไม่ใช่ความคลาดเคลื่อนซึ่งเกิดจากความแตกต่างระหว่างพันธุ์อย่างเดียว แต่มีอำนาจของแฟกเตอร์อื่น ๆ ที่ทำให้ความคลาดเคลื่อนในระหว่างต้นระหว่าง plot หนึ่งอยู่ภายใน นั่นเป็นเพราะหากของ $S.S._{var.}$ ไปแล้วจึงยังไม่ใช่สามารถสรุปความแตกต่างระหว่างพันธุ์ใด จนกว่าเราหา $S.S._{}$ ของแฟกเตอร์อื่น ๆ ที่เราไม่รู้ได้เสียก่อน แล้วจึงนำมาเทียบกับ $S.S._{}$ ของแฟกเตอร์อื่น ๆ ที่กล่าวนี้เราเรียกว่า $S.S._{error}$

การหา Sum of square สำหรับระดับความอุดมสมบูรณ์ของดิน เนื่องจากการที่เราวางแผนออกเป็น Block นั้น ทิศทางของ Block เป็นทิศทางที่จะแสดงความแตกต่างของความอุดมสมบูรณ์ของดินในแต่ละ Block ได้ การทำลาย ๆ Block จึงเป็นวิธีหนึ่งที่เราจะใช้ชนิดใดก็ตามหาขอบเขตความคลาดเคลื่อนของความอุดมสมบูรณ์ของดินแต่ละจุด มิใช่แต่เรื่องดินเท่านั้น ถ้าเราทดลองคนควาของเราเป็นที่เกี่ยวของกับสิ่งที่จะต้องการเนื้อที่และความแตกต่างระหว่างที่หรือจุดที่ตั้งมีผลกระทบกระเทือนการทดลองแล้ว การกระทำเป็น Block ก็เป็นสิ่งจำเป็นอย่างยิ่งเช่นการทดลองในเรือนกระจก บางแห่งอยู่ใกล้ทางระบายอากาศ บางแห่งอยู่ไกล บางแห่งได้แสงสว่างน้อย บางแห่งได้แสงสว่างมาก ฯลฯ หรือการทดลองในเรือนเรือนเพาะชำ บางตอนลมโกรกมาก บางตอนลมโกรกน้อย บางตอนแสงมาก แสงน้อย หรือบางตอนได้แดดเช้า บางตอนได้บ่าย ฯลฯ เป็นต้น ดังนั้นผู้ดำเนินการวางแผนการทดลองคนควาจึงจำเป็นต้องได้ศึกษาหรือมีความรู้ถึงแฟกเตอร์ที่อาจจะกระทบกระเทือนการทดลองเรื่องนั้น ๆ ให้แจ่มแจ้งเสียก่อน จึงสามารถวางแผนการทดลองคนควาได้รัดกุม

การหาความคลาดเคลื่อนของระดับ Block ก็จำเป็นจะต้องวางตัวเลขผลผลิตหรือพิจารณาตัวเลขผลผลิตไปตามทิศทางระหว่าง Block เพื่อจะได้ทราบว่าผลผลิตของแต่ละ Block นั้นมีเปลี่ยนแปลงไปอย่างไรบ้าง ทั้ง ๆ ที่ทุก ๆ Block ก็มีครบทุกพันธุ์ ๆ ละ ๑ แปลง และจำนวนต้นก็เท่ากัน

วิธีการคำนวณก็กระทำได้อีกสองสูตรเช่นเดียวกับการหา S.S.varieties และถ้อยคำก็เหมือนกัน ดังได้อธิบายไว้ในเรื่องการ S.S.varieties แล้ว เพียงแต่กระทำไปคนละทิศทางหรือเป็นทิศทางตัดกับ varieties เท่านั้นเอง

การหา S.S.Block แบบใช้สูตร $\sum d^2$

Block					$\sum X_{Block}$	d	d^2
I	18	25	16	23	82	-2	4
II	16	23	18	23	80	-4	16
III	19	24	17	25	85	1	1
IV	21	24	19	24	88	4	16
V	20	25	18	22	85	1	1
					420	0	38

$$M_{Block} = \frac{420}{5} = 84$$

แต่ $\sum X_{Block}$ ใน ๕ จำนวนนั้นไม่ใช่จำนวนเดียว แต่ละจำนวนมาจากเลขย่อยอีก ๔ จำนวน ดังนั้นค่า $\sum d^2_{Block}$ จึงเท่ากับ ๓๘ จึงไม่ใช่แสดงความคลาดเคลื่อนของเลขชุดเดียว แต่เป็นเครื่องแสดงทั้ง ๔ ชุด ชุดละ ๕ จำนวน จึงต้องหารเฉลี่ยโดยหารด้วยจำนวนพันธุ์คือ ๔ ($n_{var.} = 4$) ก็เท่ากับต้องการทราบความคลาดเคลื่อนในเรื่องดินเฉลี่ยต่อหนึ่งพันธุ์นั่นเอง

$$\text{ดังนั้น } S.S.Block = \frac{\sum d^2_{Block}}{n_{var.}} = \frac{38}{4} = 9.5$$

การหา S.S.Block แบบใช้สูตร $\sum X^2 - C.F.$

Block					$\sum X_{Block}$	$\sum X^2_{Block}$
I	18	25	16	23	82	6724
II	16	23	18	23	80	6400
III	19	24	17	25	85	7225
IV	21	24	19	24	88	7744
V	20	25	18	22	85	7225
					420	$\sum X^2 = 35318$

$$C.F. = (\sum X)^2 / n = \frac{420^2}{20} = 8820$$

สำหรับ $\sum X^2_B$ ก็จำเป็นต้องหารเฉลี่ยด้วย $n_{var.}$ โดยมีเหตุผลเช่นเดียวกันกับในเรื่องการหา S.S.var. ซึ่งได้อธิบายมาแล้ว ฉะนั้น

$$S.S. \text{ Block} = \frac{(\sum \text{Block})^2}{n \text{ var.}} - C.F. = \frac{35318}{4} - C.F.$$

$$= 8829.5 - 8820.0 = 9.5$$

การหา Sum of Square สิ่งที่เกิดจากแฟกเตอร์อื่น ๆ ซึ่งเราวัดไม่ได้ หรือ S.S. error ขณะนี้เราได้อะไรมา S.S. Total ซึ่งเป็นยอดใหญ่ทั้งหมดของการทดลองแล้ว นอกจากนั้นเรายังได้ทราบส่วนประกอบสำคัญอีกสองส่วน ก็คือความคลาดเคลื่อน ซึ่งแสดงออกทางพันธุ์ (S.S. var.) และที่แสดงทาง Block (S.S. Block) ดังนั้นในการวางแผนการทดลองแบบ Randomized Block นี้ ก็มีสิ่งอื่นที่เราสามารถวัดได้อีก ดังนั้นสิ่งที่เหลือจากการเอาสิ่งที่วัดได้ไปหักออกจากยอดใหญ่ จึงถือว่าเป็นสิ่งที่เกิดจากแฟกเตอร์อื่น ๆ ซึ่งเราวัดไม่ได้ หรือ S.S. error ทั้งหมด เราจึงอาจเขียนสูตรได้ดังนี้

$$S.S. \text{ Total} = S.S. \text{ var.} + S.S. \text{ Block} + S.S. \text{ error}$$

$$\text{หรือ } S.S. \text{ error} = S.S. \text{ Total} - (S.S. \text{ var.} + S.S. \text{ Block})$$

$$\text{แทนค่า } S.S. \text{ error} = 190.0 - (162.0 + 9.5) = 18.5$$

ตามที่ได้อธิบายวิธีการคำนวณที่แล้ว ๆ มาให้เห็นทั้งสองวิธี และได้อธิบายเท่ากันทั้งสองวิธีด้วย ซึ่งจะทำให้ผู้ศึกษาใหม่ได้มีความกระจ่างแจ้งและเข้าใจเหตุผลไม่มีอะไรใหม่แปลกประหลาดหรือยากไปกว่าเรื่องอื่น ๆ ที่แล้วมา วิธีการที่ใช้สูตร $\leq d^2$ ซึ่งเป็นวิธีการหลักซึ่งจะทำให้เห็นประโยชน์ และเข้าใจเหตุผลของการคำนวณได้อย่างแจ่มแจ้ง แต่ก็เป็วิธีการที่ยากยาวและสับสน เพราะจะต้องหา mean และหา deviation ทุกครั้ง ส่วนวิธีการใช้สูตร $\leq X^2 - C.F.$ นั้น แม้ว่าจะไม่ใช่สูตรตรงแต่เป็นสูตรที่ได้อะไรมา $\leq d^2$ เหมือนกัน (ดังนั้นพิสูจน์มาแล้วในตอนต้นในเรื่องการหา mean ค่า) แต่วิธีการใช้สูตร $\leq X^2 - C.F.$ เป็นวิธีการสั้นกว่ามาก เพียงแต่หา S.S. ของเรื่องอะไรก็ได้เอายอดใหญ่ของเรื่องนั้นยกกำลังสองแล้วรวมกัน ถ้ายอดนั้นมาจากเลขหลาย ๆ จำนวน เมื่อรวมกำลังสองแล้วก็หารเฉลี่ยด้วยจำนวนเลขในรายการหนึ่งรายการใดของยอดนั้น ส่วน C.F. ก็หาเพียงทีแรกครั้งเดียวแล้วก็นำไปใช้ตลอดในการทดลองเรื่อย ๆ เพราะยอดใหญ่ของการทดลองมีอยู่เดียวเท่านั้น ดังนั้นการที่นำเอาวิธี

คำนวณแบบไขว้สูตร $\leq d^2$ มาแสดงก็เพื่อให้เข้าใจเหตุผล แต่ในทางปฏิบัติจริง ๆ เราใช้วิธีคำนวณโดยไขว้สูตร $\leq X - C.F.$ เพราะสะดวกกว่าและสั้นกว่ามาก เนื่องจากสูตรนี้เหมาะเรื่องของการยกกำลังสอง ดังนั้นจึงหาวิธีทำตารางซึ่งสามารถยกกำลังสองไว้ให้เสร็จทุกคานทุกมุม เมื่อจะต้องการใช้ในเรื่องใดก็สามารถหยิบยกตัวเลขนั้นมาคำนวณต่อได้โดย นอกจากนั้นการสร้างตารางสำเร็จรูปซึ่งยกกำลังสองไว้เป็นระเบียบเรียบร้อย ยังเป็นเครื่องมือสำหรับการตรวจสอบตัวเลขง่ายและสะดวกอีกด้วย หากมีการผิดพลาดขึ้นในระยะใดก็สามารถตรวจพบได้ง่าย ส่วนวิธีการคำนวณก็กระทำวนรัศไปเลย ดังจะได้แสดงให้เห็นต่อไปนี้

หมายเหตุ วิธีทำ Code ควบไปด้วยโดยการ code ด้วย -๓๐ ดังนั้นในช่องเดียวกันเลขจำนวนบนเป็นวิธีธรรมดา เลขจำนวนล่างเป็นวิธี code

Block	Varieties								X_B	X_B^2
	A		B		C		D			
	X	X^2	X	X^2	X	X^2	X	X^2		
I	18		25		16		23		82	6724
	-2	4	5	25	-4	16	3	9	2	4
II	16		23		18		23		80	6400
	-4	16	3	9	-2	4	3	9	0	0
III	19		24		17		25		85	7225
	-1	1	4	16	-3	9	5	25	8	64
IV	21		24		19		24		88	7744
	1	1	4	16	-1	1	4	16	8	64
V	20		25		18		22		85	7225
	0	0	5	25	-2	4	2	4	5	25
$\Sigma X_{var.}$	94		121		88		117		420	35318
	-6		21		-12		17		20	118
$\Sigma X_{var.}^2$	8836		14641		7744		13689			
	36		441		144		289			

ต่อไปนี้จะแสดงการเปรียบเทียบวิธีการทั้งสองวิธีเพื่อให้เห็นกระจ่างยิ่งขึ้น

รายการ	สูตรที่ใช้	วิธีการคำนวณธรรมดา	Code เพื่อใช้หาค่าเฉลี่ย
C.F. Total	$\frac{(\sum X)^2}{n}$	$\frac{420^2}{20} = 8820$	$\frac{20^2}{20} = 20$
S.S. Total	$\sum X^2 - C.F.$	$(18^2 + 16^2 + 19^2 + 21^2 + 20^2 +$ $+ 25^2 + 23^2 + 24^2 + 21^2 + 25^2 +$ $+ 16^2 + 18^2 + 17^2 + 19^2 + 18^2 +$ $+ 23^2 + 23^2 + 25^2 + 24^2 + 22^2) -$ $- 8820 = 190.0$	$(-2^2 + -1^2 + -1^2 + 1^2 + 0^2 +$ $+ 5^2 + 3^2 + 4^2 + 4^2 + 5^2 + 4^2 +$ $+ -2^2 + -3^2 + -1^2 + -2^2 + 3^2 +$ $+ 3^2 + 5^2 + 4^2 + 2^2) - 20$ $= 190.0$
S.S. var.	$\sum \frac{X^2}{n_{Block}} - C.F.$	$\frac{94^2 + 121^2 + 88^2 + 117^2}{5} - 8820$ $= 162.0$	$\frac{-6^2 + 21^2 + -12^2 + 17^2}{5} - 20$ $= 162.0$
S.S. Block	$\sum \frac{X^2}{B} - C.F.$	$\frac{82^2 + 80^2 + 85^2 + 88^2 + 85^2}{4} -$ $- 8820 = 9.5$	$\frac{2^2 + 0^2 + 5^2 + 8^2 + 5^2}{4} - 20$ $= 9.5$

$$\begin{aligned}
 S.S. error &= S.S. Total - (S.S. var. + S.S. Block) \\
 &= 190.0 - (162.0 + 9.5) \\
 &= 18.5
 \end{aligned}$$

จะเห็นได้ว่าทั้งวิธีการคำนวณแบบธรรมดาและวิธีการใช้ Code ไขว่ขว้าง
 เท่ากัน แต่เราไม่มีเครื่องคิดเลข วิธีการ Code นี้ว่าเหมาะสมกว่า เพราะทำให้
 ตัวเลขลดน้อยลงไปมากสะดวกแก่การคิด ทำให้มีโอกาสผิดได้น้อย นอกจากนั้นยังรวดเร็ว
 ยิ่งขึ้น

สำหรับสูตรที่ใช้ในการคำนวณแต่ละรายนี้ เราใช้สูตร $\sum X^2 - C.F.$
 เป็นพื้นฐานอยู่แล้ว หากแต่นำไปใช้ในรายการของอะไร X ก็หมายความว่าถึงยอดแต่ละชนิด
 ของรายการนั้น เช่นนำไปใช้ในรายการของพันธุ์หรือ varieties เราใช้นับอยู่ ๔ พันธุ์

คือ A, B, C, D X จึงหมายถึงยอดหรือผลรวมของยอดดิบของแต่ละพันธุ์ใน Experiment
นี้ ดังที่แสดงข้างต้น

$$X_{\text{var. A}}^2 + X_{\text{var. B}}^2 + X_{\text{var. C}}^2 + X_{\text{var. D}}^2 = \sum X_{\text{var.}}^2$$

เช่น $94^2 + 121^2 + 88^2 + 117^2 = 44910$ เช่นนี้เป็นต้น แต่บางทีอาจหา

ให้บางคนเกิดความสับสนในสูตรเหล่านี้ เพราะจะต้องกังวลเรื่องพิจารณาว่า X ของอะไร
เพื่อจะจัดอุปสรรคเรื่องนี้จึงใคร่จะใช้สัญลักษณ์ในสูตรเสียใหม่ โดยตั้งหลักเกณฑ์ดังต่อไปนี้

(๑) ใช้ในอักษรย่อจากชื่อของรายการนั้นโดยตรง เช่น Block ก็ใช้ B
หรือพันธุ์ (Variety) ก็ใช้ V แต่บางทีอนุโลมได้ในเมื่อการทดลองคนคว้านนั้นไม่เกี่ยวกับพันธุ์ แต่เกี่ยวกับวิธีการปลูก ปฏิบัติ ไร่ หรือการไถยา ก็เป็นการทดลองเกี่ยวกับ Treatment
อาจใช้ V แทนได้ เพราะอยู่ในสภาพเดียวกัน

(๒) ลักษณะตัวเลขในสูตรมีอยู่ ๒ ลักษณะคือ เป็นยอดหรือผลรวมลักษณะหนึ่ง
กับแสดงจำนวนอีกลักษณะหนึ่ง ดังนั้นจึงใคร่วางหลักเกณฑ์ไว้ด้วยว่า อักษรย่อตามข้อหนึ่งนั้น
จะหมายถึงยอดหรือผลรวมของสิ่งนี้แก้ไขตัวใหญ่ เช่นผลรวมของ Variety A ก็ใช้ V_A
และจะใช้ $X_{\text{var. A}}$ ถ้าหากตัวเลขนั้นหมายถึงจำนวนของอะไรก็แก้ไขตัวย่อของสิ่งนั้น
เป็นว่าแน่ เช่นจำนวน Variety เราแก้ไข ๗๖ เป็นต้น จากหลักเกณฑ์ที่วางไว้นี้
เราก็เขียนสัญลักษณ์ของสูตรได้อีกแบบหนึ่ง จากสูตรที่เราคำนวณกันมาแล้วคือ

$$S.S._{\text{var.}} = \frac{\sum X_{\text{var.}}^2}{n_{\text{Block}}} - \text{C.F. or } \frac{\sum V^2}{b} - \text{C.F.}$$

$$S.S._{\text{Block}} = \frac{\sum X_{\text{Block}}^2}{n_{\text{var.}}} - \text{C.F. or } \frac{\sum B^2}{v} - \text{C.F.}$$

ส่วน $S.S._{\text{Total}}$ นั้นก็แก้ไขตามเดิม เพราะคิดจาก X เป็นเลขยออยู่แล้ว มิได้
เป็นผลรวมอย่าง $S.S._{\text{var.}}$ และ $S.S._{\text{Block}}$ และ n ของ Total ก็เป็นจำนวน
ตัวเลขทั้ง Experiment อยู่แล้ว

การใช้สูตรไว้อีกแบบหนึ่งก็ เพื่อให้ให้นักศึกษาได้เลือกใช้อาแบบใดแบบหนึ่งตาม
ถนัด แล้วแบบใดจะทำให้เข้าใจแก่ตนเองได้ง่ายกว่า เพื่อให้ผลการคำนวณเป็นไปโดยถูกต้อง
ต่อไปนี้จะขอกล่าวถึงวิธีการคำนวณต่อไป

เราได้นำมาหา S.S. Total, S.S. Block, S.S. var., & S.S. error
ไว้แล้ว ในตาราง Analysis of Variance (ย่อว่า A. of V.) เพื่อ
ประโยชน์ในการคำนวณวัดระดับความสำคัญให้ผลงานนั้นเป็นที่เชื่อถือได้หรือไม่ เช่นในการ
เปรียบเทียบพันธุ์ตามตัวอย่างนี้ เราก็จะต้องการหาว่าพันธุ์หนึ่งให้ผลดีกว่าอีกพันธุ์หนึ่งนั้น
ตัวเองต้องนับคุณสมบัติดีกว่าจริง ๆ หรือว่าดีกว่าเพราะมีเหตุอื่น ๆ มากกระทบกระเทือนทำให้
ดังนี้

ตาราง Analysis of Variance

Sources of Difference	D.F. (n-1)	S.S.	Variance (M.S.) ⁺¹	F Ratio	
				Calculated	Table
Total	19	190.0	S.S./D.F.		
Block	4	9.5	2.375		
Varieties	3	162.0	54.00	35.018	
Error	12	18.5	1.542		

๑. ของ Sources of Difference เป็นของแสดงรายการ
สาเหตุที่ทำให้เกิดความแตกต่าง" ซึ่งมีสาเหตุใหญ่ทั้งหมดเป็นรายการแรก รายการที่
สองก็คือสาเหตุในเรื่องของสิ่งปลูกซึ่งอาจเป็นดิน หรือแฟกเตอร์อื่นเกี่ยวกับพื้นที่ หรือความ
แตกต่างตามผลประจำที่ รายการที่สามก็คือสาเหตุจากสิ่งที่เราต้องการค้นคว้า และรายการ
สุดท้ายคือ Error ก็รวมเอาสาเหตุอื่น ๆ ที่เรารู้ไม่ได้อาไว้วางกันหมด

+ 1 = Variance or M.S. ย่อมาจาก Mean Square

๒. ของ Degree of Freedom หรือ D.F. ถ้าเป็น D.F. ของรายการใด ก็เท่ากับจำนวนรายการนั้นลบด้วย ๑ เช่น Total หรือรวมหมด ก็คือจำนวนรายการย่อยที่สุดใน Experiment คือ ๒๐ แลวลบด้วย ๑ เหลือ ๑๙, Block มีอยู่ ๕ Block ลบ ๑ D.F. ของ Block จึงเป็น ๔, พันธุ์หรือ varieties มี ๔ พันธุ์ ลบด้วย ๑ ดังนั้น D.F. varieties จึงเป็น ๓ เมื่อเราทราบ D.F. ของรายการใดก็คือ ๑๙ และทราบ D.F. รายการ Block และ D.F. Varieties แล้ว เราก็ทราบสิ่งที่เหลือ D.F. ของ error ได้คือ

$$D.F. \text{ error} = D.F. \text{ Total} - (D.F. \text{ Block} + D.F. \text{ var.})$$

$$= 19 - (4 + 3)$$

$$\text{or } D.F. \text{ error} = D.F. \text{ Block} \times D.F. \text{ var.}$$

$$= 4 \times 3 = 12$$

๓. S.S. นั้น เราคำนวณได้ไว้แล้วจึงนำมากรอกลงตามรายการใหญ่คือ
๔. ของ Variance หรือบางตำราเรียกว่า Mean Square หรือใช้ตัวย่อ M.S. นั้น เราได้ทราบตั้งแต่ตอนเริ่มต้นบทที่ว่าด้วย Analysis of Variance แล้วว่า Variance ก็คือ S.E. ยกกำลังสอง ดังนั้นสูตรคำนวณหา S.E. ซึ่งต้องถอด Square root แต่เราไม่ถอดก็ยังคงเป็น Variance ซึ่งที่จริง Variance ก็ใช้สูตร $\frac{\sum X^2 - C.F.}{D.F.}$ หรือ S.S./D.F. เอง ดังนั้นในของ Variance เราก็สามารถคำนวณหา Variance ของแต่ละรายการได้โดยเอา S.S. ของรายการนั้นหารด้วย D.F. ของรายการเดียวกัน เช่น Variance ของรายการ Block ก็เท่ากับ $\frac{9.5}{4} = 2.375$ หรือ Variance ของ Varieties ก็เท่ากับ $\frac{162.0}{3} = 54.00$ และ Variance ของ error ก็คงหาเช่นเดียวกันคือ เอา $\frac{18.5}{12} = 1.542$ ซึ่งความจริงเป็น ๑.๕๔๑ กว่า ๆ แต่ปัดเศษเป็น ๑.๕๔๒ ส่วน Variance ของ Total ไม่จำเป็นต้องหาเนื่องจากการหา D.F. หรือ S.S. ของ Total ก็ได้ ก็เพื่อจะนำไปหา D.F. และ S.S. ของ error

เพราะ Total เป็นยอดใหญ่ ถ้าหากเราไม่รู้อยกใหญ่แม่เราจะรู้ S.S., D.F. ของ Block และ Varieties ซึ่งประกอบอยู่ในยอดใหญ่นั้นก็จริง แต่เราก็ไม่ทางจะรู้สิ่งที่เหลืออื่น ๆ คือ Error ได้จนกว่าจะได้รู้อยกใหญ่และนำรายการ Block และ Varieties ไปหักออก ดังนั้นเมื่อเรารู้ D.F. และ S.S. ของ error แล้วก็หมดภาระที่จะต้องเกี่ยวข้องกับเรื่องของ Total อีกต่อไป

๕. ของ "F" ratio ถ้าพลิกกลับไปดูเรื่องการหา Sum of Square ของ Varieties ในหน้าก่อน ๆ จะพบว่า ได้อธิบายไว้ตอนหนึ่งว่า S.S. varieties หรือขอบเขตของความแตกต่างที่แสดงออกทางพันธุ์ เช่นผลรวมของพันธุ์ A ต่างกับผลรวมของพันธุ์ B เหล่านั้น มิใช่เกิดจากความแตกต่างระหว่างพันธุ์เองแต่อย่างเดียว แต่มีสาเหตุอื่น ๆ ที่เราไม่รู้ปะปนอยู่ด้วย จึงไม่อาจทำให้เราสรุปลงไปได้ว่าพันธุ์นั้นดีกว่าพันธุ์อื่น ทั้ง ๆ ที่ผลผลิตของพันธุ์นั้นสูงกว่าพันธุ์อื่นมาก โดยเหตุที่เรายังไม่ทราบถึงความแตกต่างที่เกิดจากสาเหตุอื่น ๆ แต่แสดงออกได้ทางพันธุ์นั้นมีมากน้อยเพียงใด ถ้าพันธุ์แสดงความแตกต่างมาก แต่ความแตกต่างในเรื่องพันธุ์จริง ๆ เราก็จะเชื่อถือและยึดเป็นหลักต่อไปได้ว่าพันธุ์นั้นดีกว่าพันธุ์จริง ปุ๋ยนั้นดีกว่าปุ๋ยจริง หรือวิธีปลูกแบบนี้ดีกว่าแบบนั้นจริง แล้วแต่เรื่องที่ทำการค้นคว้าทดลอง ส่วนความแตกต่างที่เกิดจากสาเหตุอื่น ๆ นอกเหนือไปจากเรื่องดินและเรื่องพันธุ์จริง ๆ เราก็หาได้แล้วคือเรื่องของ Error แต่ก็มีปัญหาอยู่ว่าเราจะยึดถือหลักเกณฑ์อันใดในกรณีที่ตัดสินว่า ในการพิจารณาความแตกต่างในเรื่องพันธุ์และความแตกต่างอันเกิดจาก Error อันใดควรจะถือหลักเกณฑ์อย่างไรในการพิจารณาตัดสินว่า "มาก" หรือ "น้อย" ในการพิจารณาความมากน้อยของสองปัญหารวมกันเอง วิธีที่รัดกุมก็คือหาทางรวมเป็นอันเดียวกันเพื่อจะได้ตัดสินลงไปได้ การรวมความสัมพันธ์แบบนี้กระทำได้โดยกรณีคิดเป็นอัตราส่วนหรือปฏิภาค (Ratio) เพื่อจะได้กล่าวสรุปได้ว่า ความแตกต่างที่แสดงออกโดยสิ่งซึ่งเรานำมาเปรียบเทียบกับมันใหญ่กว่าหรือมากกว่าความแตกต่างที่เกิดจากสิ่งรบกวนอื่น ๆ ก็เท่าตัว ถ้ายิ่งมากกว่าหลาย ๆ เท่าก็แสดงว่าพันธุ์หรือสิ่งที่เรานำมาเปรียบเทียบกับมันนั้นดีกว่าหรือเลวกว่ากันจริง แต่ถ้ามหาความแตกต่างอันเกิดจากสิ่งรบกวนอื่น ๆ มากตามตัวขึ้นไปด้วย ทำให้ค่าของ Ratio น้อยหรือแคบลง ความแตกต่างที่พันธุ์แสดงออกทั้งหมดอาจเกิดจาก

สิ่งรวมกันอื่น ๆ ทั้งหมดก็ได้ ดังนั้นปฏิภาคหรือ ratio ระหว่างตัวเลขที่วัดความแตกต่างซึ่งแสดงโดยสิ่งที่ต้องการเปรียบเทียบกับตัวเลขที่วัดความแตกต่างซึ่งแสดงโดยแฟกเตอร์อื่น ๆ ที่เราแยกไม่ได้ เราจึงเรียก ratio นี้ว่า "F ratio" ดังตัวอย่างเช่น

$$\text{F ratio for Varieties} = \frac{\text{Variance of varieties}}{\text{Variance of error}} \quad \checkmark$$

หรือ Variance ของ varieties ใหญ่เป็นกี่เท่าของ Variance ของ error
ใช้สำหรับวัดความแน่นอนของความแตกต่างระหว่างพันธุ์

$$\text{F ratio ของ Block} = \frac{\text{Variance ของ Block}}{\text{Variance of error}}$$

ใช้สำหรับวัดความแน่นอนของความแตกต่างเรื่องระดับความสมบูรณ์ของดินหรือ
สิ่งปลูก ดังตัวอย่างการทดลองที่ยกตัวอย่างไว้

$$\text{F. ratio ของ Varieties} = \frac{54.00}{1.542} = 35.018$$

F ratio ที่เราคำนวณได้จากการทดลองครั้งนี้จริง ๆ นี้เราเรียกว่า

Calculated F ratio

ถ้า F ratio ที่เราคำนวณได้ครั้งนี้มากเกินเท่าใด ผลงานของเรา ก็จะแสดงออกให้เป็นที่เชื่อถือได้เด่นชัดยิ่งขึ้นเท่านั้น เพราะเหตุว่าตัวเลขที่เป็นเครื่องวัดสิ่งก่อกวนนั้นเป็นตัวหาร ดังนั้น ถ้าค่าของ F ratio สูงขึ้น แสดงว่าตัวตั้งมากกว่าตัวหารหลาย ๆ เท่า ตัวตั้งก็คือสิ่งที่เราต้องการทราบความแตกต่างนั่นเอง เมื่อเหตุผลเป็นดังนี้ จึงเกิดมีข้อเตือนใจที่พึงระวังอีกอย่างหนึ่ง คือการปัดเศษในวิธีการคำนวณของเรา เนื่องจากเราต้องการผลงานที่มีความแน่นอนเป็นที่เชื่อถือได้ ถ้าจำเป็นจะต้องมีการปัดเศษ ผลของการปัดเศษจะต้องไม่ทำให้ค่าของ F ratio สูงขึ้น เพราะเท่ากับเราช่วยให้ผลงานลดความเป็นจริงให้น้อยลง แต่ควรปัดเศษให้ค่าของ F ratio อยู่ข้างต่ำ ผลงานของเรา ก็จะเป็นที่เชื่อถือได้จริง ๆ เพราะเท่ากับเราเข้มงวดในการตรวจความแน่นอนยิ่งขึ้น วิธีปัดเศษเพื่อได้ค่าของ F ratio

อยู่คนค่าก็คือ ถ้าตัวตั้งจะตองปัดเศษก็ขอให้ปัดลงต่ำ แต่ถ้าวหารจะตองปัดเศษก็ให้ปัดขึ้นทางสูง การปัดเศษนี้ ถ้าหากค่าของ F ratio ที่คำนวณได้ ได้เรียหรือค่าบเส้นกับค่าของ F ratio มาตรฐานจากตารางแล้ว ต้องใช้ความระมัดระวังและเข้มงวดการปัดเศษอย่างที่สุด ถ้าหากปัดเศษให้ F ratio ของการค้นคว้าสูงขึ้นเพียงนิดเดียว ก็จะสรุปไควความลงงานนั้นแน่นอนเชื่อถือได้ ซึ่งความจริงอยู่ในระหว่างกำกั่งกันจึงทำให้มีการสรุปผิดไปจากความจริงในด้านที่อาจทำให้เกิดความเสียหายแก่ผู้เชื่อถือหรือนำไปปฏิบัติได้ ดังนั้น ถ้าผลการวิเคราะห์แสดงคลุมเครือให้สงสัย ควรตัดสินใจไปในทางไม่เชื่อถือไว้เป็นปลอดภัยที่สุด เพราะไม่เกิดความเสียหายแต่ประการใด ดีกว่าเชื่อถือสิ่งที่ยังไม่เป็นจริงเช่นนั้น จะทำให้เกิดความเสียหายขึ้นได้

การตรวจความแน่นอนของผลการค้นคว้า ในเรื่องของ Analysis of Variance นี้ หลังจากได้ทำการวิเคราะห์ตัวเลขมาจนถึงระยะนี้แล้ว เราก็เริ่มดำเนินการตรวจความแน่นอนของผลงานได้ การตรวจแบ่งออกเป็น ๒ ระยะด้วยกัน ระยะแรกเป็นการตรวจอย่างหยาบ ๆ หรือเรียกว่าเป็นการตรวจขั้นต้น กระทำโดยอาศัย F ratio เป็นเครื่องมือ เพื่อต้องการจะทราบว่าการค้นคว้าของเรานั้น มีสาระสำคัญตามความมุ่งหมายเป็นที่เชื่อถือได้หรือไม่ ถ้าหากผลการตรวจระยะนี้แสดงว่าเชื่อถือได้ เราก็จะไคทำการตรวจโดยละเอียดต่อไปเพื่อสรุปว่าปัญหาใดสิ่งใดแสดงผลให้เชื่อถือได้มากน้อยเพียงใด แต่ถ้าการตรวจขั้นต้นแสดงว่าการค้นคว่านั้นไม่เกิดผลให้เป็นที่เชื่อถือ ก็ไม่จำเป็นต้องเสียเวลาไปตรวจให้ละเอียดต่อไปอีก ดังตัวอย่างเช่น การค้นคว้าเปรียบเทียบพันธุ์พืช ๔ พันธุ์ดังกล่าวแล้ว ถ้า F ratio ของ Variety แสดงว่าพันธุ์ไหนแตกต่างกันจริงเราก็หาต่อไปอีกว่าพันธุ์ไหนต่างกับพันธุ์ไหน พันธุ์ไหนดีที่สุดและเลวที่สุด แต่ก็มีไคหมายความว่าต้องต่างกันทุกพันธุ์ อาจมีบางพันธุ์ไม่ต่างกันได้ แต่ถาการตรวจด้วย F ratio แสดงว่าพันธุ์ไม่ต่างกันแล้ว แม้ว่าจะคำนวณต่อไปก็จะพบว่าไม่มีพันธุ์ไหนต่างกันเลย

การตรวจด้วย F ratio เป็นการตรวจสอบขั้นแรกดังไคกล่าวมาแล้ว ในการเปรียบเทียบผลผลิตของพันธุ์พืช ๔ พันธุ์ตามตัวอย่างเรากำนวน F ratio ของพันธุ์ไคแล้วคือ $= ๓๕.๐๑๘$ เราได้จากเหตุผลที่ไคอธิบายมาแล้วว่าค่าของ F ratio มากขึ้นเท่าไค ความแตกต่างระหว่างพันธุ์จะยิ่งเป็นที่เชื่อถือได้มากเท่านั้น แต่เราจะมีระดับอะไรที่จะมาตัดสินไควว่า

F ratio จะต้องสูงเท่านั้นเท่านั้นผลงานนั้นจึงจะเชื่อถือได้ นักคำนวณทางสถิติได้ทำการตรวจสอบความแปรปรวนของตัวเลขจากปรากฏการณ์ต่าง ๆ และได้วัดขอบเขตของความแปรปรวน จึงได้คำนวณออกมาเป็นระดับความเชื่อถือ และตั้งไว้เป็นตารางมาตรฐาน เพื่อประโยชน์ให้นักคนอื่น ๆ ได้ใช้วัดผลงานของตน ระดับมาตรฐานของ F ratio ก็มีปรากฏในตาราง F ratio เช่นเดียวกับ "t" ratio ซึ่งได้กล่าวมาแล้วในบทอื่น ๆ หากแต่ตาราง F ratio มีอยู่ ๒ ทิศทาง (2 direction) เนื่องจาก Degree of freedom มี ๒ รายการควบกัน เพราะจะต้องคิดเทียบความแตกต่างที่เกิดจากสาเหตุภายนอกอื่น ๆ (error) กับความแตกต่างที่แสดงออกทางพันธุ์ การหาระดับมาตรฐานจากตารางก็กระทำได้โดยให้สังเกตดูว่าด้านหนึ่งของตารางเขียนไว้ว่า degree of freedom for smaller meansquare ซึ่งหมายความว่า D.F. ของ error (ตามตัวอย่าง == ๑๒) อีกด้านหนึ่งคือ degree of freedom for greater mean square ซึ่งหมายความว่า D.F. ของสิ่งที่เราต้องการเปรียบเทียบ (ในตัวอย่างคือ พันธุ์ซึ่ง D.F. = ๓) เมื่อดูตรงเข้าไปในตาราง ตั๊กก้นที่ช่องใด ค่าของ F ratio มาตรฐานสำหรับเปรียบเทียบพันธุ์ในการทดลองนั้นก็จะอยู่ในช่องนั้น

Degree of Freedom for Greater Mean Squar

	1	2	3	4	5	6	7	8	9	10	11	12	...
Degree of	1												
Freedom for	2												
Smaller	3												
Mean	4												
Square	5												
	6												
	7												
	8												
	9												

D.F. varieties
= 3

	10	
	11	
D.F. error	12	→ 3.49 5.95
	13	
	14	F value
	15	
	16	
	17	
	18	
	19	
	20	

การที่กำหนดกฎเกณฑ์ตายตัวลงไปว่า D.F. ของ Mean Square ที่เล็กกว่าคือ D.F. ของ error และ D.F. ของ M.S. ที่ใหญ่กว่าคือ D.F. ของสิ่งที่ต้องการ นั่น ความจริงมีเหตุผลอยู่ว่าเราจะหา F ratio ก็ต่อเมื่อ M.S. ของ error เล็กกว่า และจะเป็นตัวหาร ฉะนั้น M.S. ของ error ใหญ่กว่า ค่าของ F ratio ย่อมจะไม่ถึง ๑ หรือกล่าวอีกนัยหนึ่งว่า ความแตกต่างที่เกิดจากการรบกวนของสาเหตุอื่น ๆ ย่อมจะมีผลทั้งลดความแตกต่างที่แสดงออกทางพันธุ์จึงไม่สามารถจะเชื่อได้ว่าพันธุ์ต่างกันจริง แม้ผลผลิตประจำพันธุ์จะแตกต่างกันไปปรากฏแก่ตาก็ตาม จึงสรุปได้ว่า F ratio มีค่าไม่ถึง ๑ ก็ไม่ต้องเปิด F table ให้เสียเวลา คงสรุปได้เลยว่าสิ่งที่เราเปรียบเทียบกันนั้นไม่มีความแตกต่างให้เชื่อถือได้ แต่เมื่อใดที่ F ratio มีค่าเกิน ๑ ขึ้นไป M.S. ของ error ก็จะคงเล็กกว่า M.S. ของสิ่งที่เราเปรียบเทียบเสมอ ด้วยเหตุนี้การเปิด F ratio จึงได้ระบบเป็นกฎเกณฑ์ไว้อย่างถาวรแล้ว

จากการค้นคว้าเปรียบเทียบพันธุ์ตามตัวอย่าง เราเปิดค่าของ F ratio จาก table ได้ $\frac{3.49}{2.52}$ เป็น ๒ ค่าเช่นเดียวกับ t ratio ในตอนต้นคือ ๓.๔๕ เป็นระดับ

๘๘:๕ หรือ ๘๘:๖ ถ้ากล่าวเป็นร้อยละก็คือ ๕ % และค่าสูงเป็นระดับ ๘๘:๑ หรือ ๑ %

(เหตุผลที่ง่ายและวิธีที่ง่ายก็กระทำเช่นเดียวกันกับ ("t" ratio) เมื่อเทียบกับ F ratio ที่คำนวณได้จากการทดลองปรากฏว่า F ratio ที่คำนวณได้จากการทดลองคือ ๓๕.๐๑๘ สูงกว่าระดับมาตรฐาน ๑ % คือ ๕.๘๕ มาก แสดงว่าพันธุ์พืชที่ทำการค้นคว้าเปรียบเทียบนี้ไม่ผลผลิตแตกต่างกันจริงโดยถือเป็นนัยสำคัญทางสถิติได้

การหา Least Significant Difference หรือ L.S.D.

เมื่อเราทราบจากการตรวจสอบด้วย F. ratio แล้วว่ามีความแตกต่างในระหว่างพันธุ์พืชที่ทำการค้นคว้าเปรียบเทียบจริง จึงดำเนินการตรวจสอบต่อไปอีกเพื่อจะหาว่าพันธุ์ใดแตกต่างกับพันธุ์ใดบ้าง และพันธุ์ไหนให้ผลดีที่สุด การตรวจสอบในขั้นนี้เป็นการตรวจสอบโดยละเอียด เพื่อจะสรุปผลการค้นคว้าให้แน่นอนลงไปโดยเด็ดขาด ดังนั้น L.S.D. จึงเป็นระดับมาตรฐานที่จะนำมาใช้ตัดสินความแตกต่างอันได้จากการเปรียบเทียบที่แท้จริง การเปรียบเทียบกับระดับ L.S.D. จึงมีความหมายว่า ความแตกต่างที่เกิดขึ้นจากการเปรียบเทียบนั้นอย่างน้อยที่สุดจะต้องเท่ากับระดับ L.S.D. จึงจะเชื่อถือได้ว่าแตกต่างกันจริง ถ้าหากความแตกต่างน้อยกว่าระดับ L.S.D. การที่สิ่งหนึ่งให้ผลดีกว่าอีกสิ่งหนึ่ง ก็ไม่เป็นที่เชื่อถือได้ ระดับ L.S.D. จะสูงหรือต่ำขึ้นอยู่กับสาเหตุหลายประการเช่น ถ้ากระทำน้อยซ้ำค่าของ L.S.D. ก็จะสูงหรือถ้าสาเหตุอื่น ๆ ที่เราควบคุมผลของเรามีมากทำให้ error สูงขึ้น ค่าของ L.S.D. ก็จะต้องสูงขึ้นด้วย เพราะแสดงว่าผลงานนั้นมีความแปรปรวนมาก ต้องแตกต่างกันมาก ๆ จึงจะเชื่อถือได้ การเปลี่ยนแปลงสูงต่ำของระดับ L.S.D. นั้นเป็นไปตามวิธีการคำนวณอยู่แล้ว เนื่องจากการคำนวณหาระดับ L.S.D. จำเป็นต้องอาศัยตาราง "t" ประกอบ Variance ของ error ถ้ากระทำมากซ้ำค่าของ t ในตารางก็จะลดลงเอง นอกจากนั้น ถ้าสิ่งควบคุมมีน้อย Variance ของ error ก็จะลดลงด้วยระดับ L.S.D. จึงสูงหรือต่ำได้แล้วแต่แปรเทอว์ หรืออาจกล่าวอีกนัยหนึ่งได้ว่า การตัดสินความแตกต่างจะเข้มงวดมากหรือน้อย ย่อมแล้วแต่การสร้างความเชื่อมั่นของผลงานนั้น

การกำหนดหารระดับ L.S.D. อาจทำได้สองวิธี คือ.—

๑. L.S.D. สำหรับตรวจสอบความแตกต่างของผลเฉลี่ย เช่นในการตรวจสอบความแตกต่างระหว่างพันธุ์ ถ้าต้องการเปรียบเทียบผลเฉลี่ย (Mean) ของแต่ละพันธุ์ ก็ใช้สูตร ดังนี้

$$L.S.D. = t \sqrt{\frac{2 \text{ Variance of Error}}{n}}$$

๒. สำหรับตรวจสอบความแตกต่างของผลรวม ถ้าเราต้องการตรวจสอบความแตกต่างระหว่างพันธุ์โดยทำการเปรียบเทียบผลรวมจากทุก ๆ ค่าของแต่ละพันธุ์ ก็ใช้สูตรนี้

$$L.S.D. = t \sqrt{2 \text{ Variance of Error} \cdot n}$$

ส่วนปัญหาที่เราจะทำการตรวจสอบวิธีที่ ๑ หรือวิธีที่ ๒ นั้นไม่เป็นสิ่งสำคัญ เพราะผลที่ได้จะเหมือนกันทั้งสองวิธี

* ค่าของ "t" นั้นได้จากตาราง "t" โดยเปิดดูที่ D.F. ของ error ค่าของ n ก็คือจำนวนค่าของแต่ละพันธุ์หรือคือค่าของตัวเลขตัวเดียวกันกับเรานำมาหารเฉลี่ยนั่นเอง ดังนั้นความหมายของการใช้ n ในสูตรก็เพื่อลดหรือเพิ่มระดับให้สอดคล้องกับการเลือกเปรียบเทียบผลเฉลี่ยหรือผลรวมนั่นเอง สมมติว่าถ้าเราหาผลเฉลี่ยของแต่ละพันธุ์โดยหารด้วย ๕ (เพราะการค้นคว้าเรื่องนี้มี ๕ ค่า) ค่าของผลเฉลี่ยก็จะน้อยกว่าผลรวม ๕ เท่า ระดับ L.S.D. ที่ใช้เทียบผลเฉลี่ยก็จะน้อยกว่าระดับ L.S.D. ที่ใช้เทียบผลรวม ๕ เท่า ด้วยเหมือนกัน

* สำหรับค่า "t" เราจะเห็นได้ว่ามีอยู่ ๒ ระดับ ดังนั้นการแทนค่า "t" ในสูตรจึงต้องแทน ๒ ครั้ง ดังนั้นค่าของระดับ L.S.D. จึงคำนวณได้ ๒ ระดับ คือระดับ ๕ % และระดับ ๑ % เช่นเดียวกัน

ส่วนการเปรียบเทียบผลต่างกับระดับ L.S.D. เพื่อสรุปผลการค้นคว้าก็ใช้หลักการที่เหมือนกันกับการเทียบ t ratio หรือ F ratio นั่นเอง ถ้าผลต่างระหว่างพันธุ์ไหนเท่าหรือสูงกว่าระดับ L.S.D. ๕ % เราก็สรุปว่าพันธุ์นั้นแตกต่างกันพอเชื่อถือได้ แต่ถ้าผลต่างเท่ากันหรือสูงกว่า L.S.D. ระดับ ๑ % ก็สรุปได้ว่าสองพันธุ์นั้นแตกต่างกันอย่าง

แน่นอนโดยถือเป็นนัยสำคัญยิ่งในทางสถิติ (Highly Significant Different)

ถ้าพันธุ์ไหนให้ผลผลิตสูงกว่าเราก็ก้าวสรุปว่าพันธุ์ไหนผลดีกว่าพันธุ์จริง

ตามตัวอย่าง สมมติว่าเราต้องการเปรียบเทียบผลเฉลี่ยของผลผลิตระหว่าง ๔ พันธุ์
นั้น ค่าของ "t" ก็เปิดตาราง t ที่ D.F. ของ error คือ ๑๒ "t" = ๒.๑๗๐ ๕ %
๓.๐๕๕ ๑ %

Variance ของ error จากตาราง Analysis of Variance = ๑.๕๔๓

$$\text{ผลเฉลี่ยของพันธุ์ A} = \frac{94}{5} = 18.8$$

$$\text{" B} = \frac{121}{5} = 24.2$$

$$\text{" C} = \frac{88}{5} = 17.6$$

$$\text{" D} = \frac{117}{5} = 23.4$$

$$\text{L.S.D.} = t \sqrt{\frac{2 \times 1.542}{5}}$$

$$= t \sqrt{0.771}$$

$$= 0.878 t$$

$$\text{L.S.D. } 5 \% = 0.878 \times 2.170$$

$$= 1.913$$

$$\text{L.S.D. } 1 \% = 0.878 \times 3.055$$

$$= 2.682$$

เรานำเอาผลผลิตของแต่ละพันธุ์มาเรียงกันจากผลผลิตสูงมาหาผลผลิตต่ำ

แล้วเปรียบเทียบความแตกต่างกันตามลำดับดังนี้คือ

$$\text{พันธุ์ B} = 24.2$$

$$\text{พันธุ์ D} = 23.4$$

$$\text{พันธุ์ A} = 18.8$$

$$\text{พันธุ์ C} = 17.6$$

$$B - D = 0.8 \text{ insignificant}$$

$$D - A = 4.6 \text{ highly significant}$$

$$A - C = 1.2 \text{ insignificant}$$

$$D - C = 8.8 \text{ highly significant}$$

เมื่อนำผลต่างระหว่างพันธุ์ กับพันธุ์ ซึ่งต่างกันเพียง ๑.๐๐ ไปเทียบกับระดับ L.S.D. ปรากฏว่าต่ำกว่าระดับ ๕ % ซึ่งมีค่าระดับเป็น ๑.๕๓๓ จึงกล่าวได้ว่าพันธุ์ B กับพันธุ์ D นั้นมีความแตกต่างไม่ถึงเป็นนัยสำคัญในทางสถิติ หรือกล่าวอย่างธรรมดาว่า B กับ D ให้ผลไม่แตกต่างกัน (ส่วนที่แสดงความแตกต่างกัน ๑.๐๐ นั้นเป็นส่วนที่เชื่อถือไม่ได้) / แต่พันธุ์ D ให้ผลสูงกว่าพันธุ์ A ถึง ๕.๓๕ สูงกว่าระดับ L.S.D. ๑ % ซึ่งมีค่า ๒.๖๘๒ แสดงว่าพันธุ์ D ดีกว่าพันธุ์ A อย่างแน่นอน หรือความแตกต่างระหว่างผลผลิตของพันธุ์ D และ A นี้ถือเป็นนัยสำคัญยิ่งในทางสถิติ จึงใส่เครื่องหมายดอกจัน ๒ ดอกไว้ที่หลัง ๕.๓๕ อันเป็นเครื่องหมายแสดงว่าแตกต่างกันอย่างแท้จริงในระดับเกินกว่า ๑ % (Highly Significant) แต่ถามว่าเกิดความแตกต่างระหว่างพันธุ์สูงกว่าระดับ ๕ % คือ ๑.๕๓๓ เกยยังต่ำกว่าระดับ ๑๐ คือ ๒.๖๘๒ ก็ถือว่าแตกต่างกันพอเชื่อถือได้ (Significant) และได้เครื่องหมายดอกจัน ๑ ดอกไว้ที่หลังผลต่างนั้น ถ้าความแตกต่างไม่ถึงแม้แต่ระดับ ๕ % ก็ถือว่าไม่มีความแตกต่างกัน (unsignificant) ไม่ใส่เครื่องหมายดอกจันเลย เช่นความแตกต่างระหว่างพันธุ์ B กับพันธุ์ D และระหว่างพันธุ์ A กับพันธุ์ C เป็นต้น แต่เมื่อตรวจสอบเช่นนี้แล้วเราก็พิศุขต่อไปได้ว่า เมื่อพันธุ์ D ดีกว่า A อย่างแน่นอนแล้ว พันธุ์ B ก็ย่อมดีกว่าพันธุ์ A อย่างแน่นอนด้วย จึงกล่าวสรุปการทดลองได้ว่า ในการทดลองเปรียบเทียบพันธุ์ A, B, C, D, นี้ ปรากฏผลว่า พันธุ์ B และพันธุ์ D ให้ผลผลิตสูงสุด ส่วนพันธุ์ A และพันธุ์ C ให้ผลผลิตต่ำที่สุด นอกจากการค้นคว้าเปรียบเทียบพันธุ์ดังโดยทั่วๆ ไปมาแล้ว แม้จะเป็นการค้นคว้าเปรียบเทียบปุ๋ย วิธีปลูกปฏิบัติ เช่นเปรียบเทียบวิธีการไถ วิธีการปลูกแบบต่าง ๆ ระยะเวลาในการใช้ปุ๋ย ความเข้มข้นของปุ๋ย อันตรายของยาฆ่าแมลง จำนวนน้ำ ฯลฯ ก็กระทำแบบเดียวกันนี้ได้ หากแต่เปลี่ยนจากเรื่องพันธุ์มาเป็นเรื่องของ Treatments เท่านั้น ส่วนรายการอื่น ๆ เช่น Block ก็คงเดิม

ตามที่ได้อธิบายในเรื่องของ Randomized Block Design มาแล้วทั้งหมดนี้ นิยมที่จะอธิบายวิธีคำนวณอย่างเดียว หากแต่ได้แสดงเหตุผล และแสดงที่มาของวิธีการคำนวณให้เห็นกระจ่างแจ้งด้วยว่า เป็นวิธีการที่เกี่ยวข้องเนื่องมาจากเบื้องต้น ๆ ทั้งสิ้น ไม่มีอะไรเป็นสิ่งที่

ยากหรือแปลกออกไปอีก แต่ในทางปฏิบัติเท่า ๆ นี้ใช้วิธีการยกกำลังสองกันโดยทั่วไป คือใช้สูตร

$$\sum x^2 - \text{C.F.} \text{ แทนที่จะใช้ } \sum d^2 \text{ เพราะการใช้วิธียกกำลังสอง}$$

สามารถเตรียมตารางยกกำลังสองให้สำเร็จรูปไว้ล่วงหน้าได้ เมื่อจะใช้คำนวณหา S.S. ของอะไรก็เอากำลังสองเหล่านั้นมาคิดได้เลย จึงนับว่าสะดวกกว่า คงจะยกตัวอย่างวิธีการคำนวณโดยตรงดังต่อไปนี้

ตัวอย่าง ในการค้นคว้าเพื่อเปรียบเทียบปุ๋ยผสม ๕ อัตราส่วน ใช้กับพืชชนิดหนึ่ง กระทำการวางแผนแบบ Randomized Block Design โดยกระทำ ๖ ซ้ำ หรือ 6 Block เมื่อเก็บสถิติตัวเลขจากผลผลิตของแต่ละแปลงแล้ว จึงนำมาวิเคราะห์ทางสถิติ ดังต่อไปนี้

ตารางยกกำลังสองและออกยอตาราง (Coded - 50)

Block	Treatments										B	B ²
	1		2		3		4		5			
	X	X ²	X	X ²	X	X ²	X	X ²	X	X ²		
I	41		44		51		55		53			
	-9	81	-6	36	1	1	5	25	3	9	-6	36
II	43		44		49		57		54			
	-7	49	-6	36	-1	1	7	49	4	16	-3	9
III	42		46		48		60		52			
	-8	64	-4	16	-2	4	10	100	2	4	-2	4
IV	44		45		50		58		54			
	-6	36	-5	25	0	0	8	64	4	16	1	1
V	40		43		52		59		55			
	-10	100	-7	49	2	4	9	81	5	25	-1	1
VI	43		45		48		56		55			
	-7	49	-5	25	-2	4	6	36	5	25	-3	9
Tr.	-47	379	-33	187	-2	14	45	357	23	95	-14	
Tr ²	2209		1089		4		2025		529			

$$\begin{aligned} \text{C.F.} &= (\sum X)^2 / n = -14^2 / 30 \\ &= 196 / 30 = 6.533 \end{aligned}$$

$$\begin{aligned} \text{S.S.Total} &= \sum X^2 - \text{C.F.} \\ &= (-9^2 + -7^2 + -8^2 + \dots + 4^2 + 5^2 + 5^2) - 6.553 \\ &= (81 + 49 + 64 + \dots + 16 + 25 + 25) - 6.553 \\ &= 1032 - 6.533 \\ &= 1025.467 \end{aligned}$$

$$\begin{aligned} \text{S.S.Block} &= \sum B^2 / \text{tr.} - \text{C.F.} \\ &= \frac{-6^2 + -3^2 + -2^2 + 1^2 + -1^2 + -3^2}{5} - 6.553 \\ &= 5.467 \end{aligned}$$

$$\begin{aligned} \text{S.S.Treatments} &= \sum \text{Tr}^2 / b - \text{C.F.} \\ &= \frac{-47^2 + -33^2 + -2^2 + 45^2 + 23^2}{6} - 6.533 \\ &= 969.467 \end{aligned}$$

$$\begin{aligned} \text{S.S.error} &= 1025.467 - (5.467 + 969.467) \\ &= 50.533 \end{aligned}$$

Analysis of Variance

Sources of Difference	D.F.	S.S.	Variance (S.S./D.F.)	F ratio	
				Calculated	Table
Total	29	1025.467			
Block	5	5.467	1.0934		
Treatment	4	969.467	242.36675	95.92 ***	2.87 5 % 4.43 1 %
Error	20	50.533	2.52665		

$$\begin{aligned}
 \text{Calculated F ratio for Treatment} &= \frac{\text{Variance of Treatment}}{\text{Variance of Error}} \\
 &= \frac{242.36675}{2.52665} \\
 &= 95.92
 \end{aligned}$$

$$\text{F value from table at } \begin{cases} \text{D.F. Treatment} = 4 \\ \text{D.F. Error} = 20 \end{cases}$$

$$\begin{aligned}
 \text{ระดับ } 5\% &= 2.85 \quad \begin{matrix} 3.26 - 14 \\ 2.76 - 25 \end{matrix} \\
 \text{ระดับ } 1\% &= 4.43
 \end{aligned}$$

F ratio ที่คำนวณได้สูงกว่าระดับ 1% ของ Table มาก แสดงว่าปุ๋ยทั้ง ๕ อัตราส่วนผสมนั้น ทำให้ผลผลิตของพืชชนิดนั้นแตกต่างกันอย่างแท้จริง ต่อไปจึงหาระดับ L.S.D. เพื่อใช้เปรียบเทียบว่าปุ๋ยผสมใดที่ให้ผลดีที่สุด

$$\text{L.S.D.} = t \cdot \sqrt{\frac{2 \text{ Variance of Error}}{n}}$$

เปรียบเทียบผลเฉลี่ยของผลผลิตของพืชจากปุ๋ยผสมอัตราส่วนต่าง ๆ กัน

ปุ๋ยส่วนผสมที่ ๑	ให้ผลผลิตเฉลี่ย	$\frac{41+43+42+44+40+43}{6} = \frac{253}{6} = 42.166$
ปุ๋ยส่วนผสมที่ ๒	ให้ผลผลิตเฉลี่ย	$\frac{44+44+46+45+43+45}{6} = \frac{267}{6} = 44.500$
ปุ๋ยส่วนผสมที่ ๓	ให้ผลผลิตเฉลี่ย	$\frac{51+49+48+50+52+48}{6} = \frac{298}{6} = 49.666$

ปุ๋ยส่วนผสมที่ ๔ ไนโตรเจนเฉลี่ย $\frac{55+57+60+58+59+56}{6} = 57.500$
 ปุ๋ยส่วนผสมที่ ๕ ไนโตรเจนเฉลี่ย $\frac{53+54+52+54+55+55}{6} = 53.833$

L.S.D. ของปุ๋ย = $t \cdot \sqrt{\frac{2 \times 2.52665}{6}}$
 = $t \cdot 0.917$

"t" Table D.F. error = 20

$t_{5\%} = 2.086$
 $t_{1\%} = 2.845$

L.S.D. 5% = 2.086×0.917
 = 1.913

L.S.D. 1% = 2.845×0.917
 = 2.609

Tr. No.	Mean	Diff.
Treatment 4.	57.500	
" 5.	53.833	3.667++
" 3.	49.666	4.167**
" 2.	44.500	5.166**
" 1.	42.166	2.334**

LATIN SQUARE DESIGN

เป็นวิธีการวางแผนอีกแบบหนึ่งที่คิดไปจากวิธีการ Randomized Block Design
 ดังกล่าวแล้ว แต่หลักเกณฑ์ในการคำนวณวิเคราะห์ผลก็จะเป็นอย่างเดียวกัน หากแต่วิธีการ
 วางแผนนั้นเองที่ช่วยให้เราสามารถแยกการเมื่อแสดงความแตกต่างออกได้อีกการหนึ่ง
 เปรียบจากวิธีการ Randomized Block Design ที่ได้มานี้แล้วภายในส่วนใหญ่ทั้งหมด
 เราสามารถแยกเอาการปลูกย่อยออกได้ ๒ รายการคือค่าหนึ่งเป็นสิ่งที่เราต้องการค้นหา
 เช่นพันธุ์หรือปุ๋ย ฯลฯ อีกค่าหนึ่งเป็นเรื่องดินหรือแฟกเตอร์ที่เราได้แต่ในเรื่องของ Latin
 Square นั้น เราสามารถตรวจหรือวัดความไม่สม่ำเสมอของดินได้ทั้งสองทางใน
 ทิศทางตั้งกัน

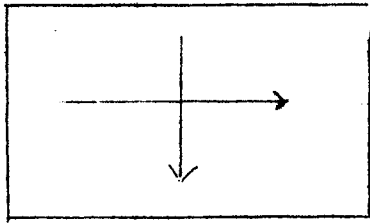
Randomized Block Design

χ^2 Total = Blocks + Varieties หรือ Treatments + สิ่งอื่น ๆ (Error)

Latin Square Design

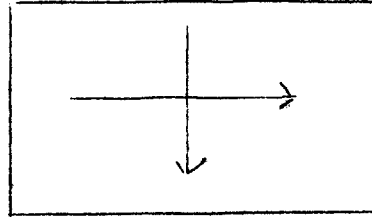
χ^2 Total = Blocks ตั้ง + Blocks นอน + Varieties หรือ Treatments +
 สิ่งอื่น (Error) (row)

Randomized Block



Varieties หรือ Treatments

Block

Latin Square

Block ตั้งเรียกว่า Column

Block นอน
เรียกว่า
Row

เมื่อปรากฏว่าแผนผังของ Latin Square นั้นแสดงยอดของ Block ทั้งสองทิศทางตัดกัน แต่บริเวณทำการทดลองอยู่ในสภาพที่เป็นพื้นที่ (Area) ไม่มีส่วนลึก มีแต่ความแบน เมื่อทั้งสองทิศเป็น Block ทั้งหมด จึงไม่แสดงความแตกต่างในเรื่องสิ่งที่ต้องการค้นคว้า เช่น Varieties หรือ Treatments เพราะไม่มีทิศทางที่จะปรากฏได้ นอกจากนั้น เมื่อทุก ๆ แถวไม่ว่าจะเป็นทางตั้งหรือนอนมีสภาพเป็น Block ทั้งนั้น จึงต้องคำนึงตามกฎเกณฑ์ของ Block กล่าวคือทุก ๆ แถวไม่ว่าจะเป็นทางตั้งหรือทางนอนก็ตาม แต่ละแถวจะต้องมีครบทุกชนิด ๆ ละ ๑ แปลง โดยไม่ซ้ำกัน เช่นต้องการทดลองเปรียบเทียบพันธุ์ แต่ละแถวไม่ว่าจะเป็นทางตั้งหรือทางนอนแต่ละแถวจะต้องมีครบทุกพันธุ์ ๆ ละ ๑ แปลง หรือ ๑ plot แถวนั้นจึงจะเรียกว่าเป็น Block ที่สมบูรณ์ได้ ดังนั้นเมื่อเวลาเก็บสถิติตัวเลขผลได้มาทำการวิเคราะห์ เราจะนำมาจัดเรียงเป็นพันธุ์ ๆ อย่างวิธีการ Randomized Block ไม่ได้ จำเป็นจะต้องรักษารูปให้คงเดิมอยู่ เพื่อรักษาสภาพของทั้งสองทิศทางให้คงเป็น Block อยู่ จะทำการคำนวณการเปลี่ยนแปลงของ Block ทั้งสองทางเมื่อแต่ละพันธุ์จำเป็นต้องกระจายกันอยู่อย่างไม่มีแถวไม่มีแนวเช่นนี้ จึงจำเป็นอย่างยิ่งที่จะต้องใส่ชื่อหรืออักษรประจำพันธุ์หรือเลขประจำ Treatment ความกับผลผลิตของแปลงนั้น ๆ ไว้โดยเฉพาะทุกแปลงไป เพราะเราไม่สามารถจะใส่อักษรประจำพันธุ์ไว้ตามหัวแถวของตารางอย่างตาราง Randomized Block ได้ ส่วนวิธีการคำนวณก็มีหลักเกณฑ์เช่นเดียวกันกับ Randomized Block ผิดกันแต่ส่วนของ Randomized Block ทิศทางหนึ่งหา S.S. ของ Block และอีกทิศทางหนึ่งหา S.S. ของ Varieties หรือ Treatments ส่วนของ Latin Square ทั้งสองทิศทางเป็น Block ด้วยกัน จึงต้องหา S.S. Block ทั้งสองทิศทางซึ่งทางคานตั้งเราเรียกว่า S.S. Column และทางคานนอนเราเรียกว่า S.S. Row ดังตัวอย่างต่อไปนี้

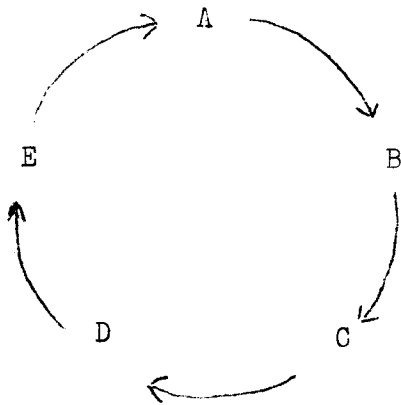
ตัวอย่างการค้นคว้าแบบ Latin Square

เราอาจคำนวณหา S.S. ของรายการต่าง ๆ โดยใช้วิธีหา deviation $\sum d^2$ หรือวิธียกกำลังสอง ($\sum x^2 - C.F. \times$) ได้เช่นเดียวกันกับการคำนวณในวิธี Randomized Block ซึ่งได้ยกตัวอย่างมาแล้ว แต่จะขอยกตัวอย่างด้วยวิธียกกำลังสอง ซึ่งใช้กันอยู่ทั่วไป เพราะสะดวกมาก

ต่อไปนี้จะกล่าวถึงวิธีการวางแผนผังการค้นคว้าเสียก่อน โดยเหตุที่วิธีการนี้ปฏิบัติ จะต้องทำทางวางแผนผังให้มันเป็น Block ได้ทั้งสองทิศทางไม่ว่าจะเป็นแถวใดจะต้องมีครบทุกพันธุ์หรือทุก Treatments และจะต้องไม่ซ้ำกันด้วย จึงเป็นการยากที่จะใช้วิธีสุ่มแบบธรรมดา แต่อย่างไรก็ตามเราก็ต้องมีวิธีการสุ่มเซารวมด้วยเพื่อกันมิให้เกิดการลำเอียง แต่ให้กระจายไปคามแต่โลกของธรรมชาติ

สมมติว่าเราต้องการเปรียบเทียบผลผลิตของพืช ๕ พันธุ์ด้วยกัน โดยการวางแผนแบบ Latin Square กำหนดให้ใช้อักษร A,B,C,D,E แทนพืชทั้ง ๕ พันธุ์ วิธีการวางแผนผังแบ่งออกได้เป็นชั้น ๆ ดังต่อไปนี้ คือ

X (๑) ทำแบบมาตรฐาน โดยจัดเรียงตัวอักษรมิให้ซ้ำกันทุก ๆ แถวไม่ว่าจะเป็นแถวตั้งหรือแถวนอน โดยการเขียนวงอักษรเวียนทางหนึ่งทางใดก็ได้ ถ้าเวียนทางใด ก็ให้ใช้เวียนทางนั้นตลอดไป



การเขียนให้เรียงตามนอน บรรทัดแรกมีครบ ๕ ตัวรวม ๑ ชุด เพื่อให้แถวตั้งเรียงครบชุดจึงเขียน ๕ บรรทัด ถ้าบรรทัดแรกขึ้นด้วยตัวอะไร บรรทัดที่สองขึ้นด้วยตัวที่สองและบรรทัดที่สามก็ขึ้นด้วยตัวเลขที่สามของบรรทัดแรก บรรทัดที่สี่ ก็ขึ้นด้วยตัวเลขที่สี่ บรรทัดที่ห้า ก็ขึ้นด้วยตัวเลขที่ห้า ตามลำดับ

		๑	๒	๓	๔	๕
Row	ที่ ๑	A	B	C	D	E
Row	ที่ ๒	B	C	D	E	A
Row	ที่ ๓	C	D	E	A	B
Row	ที่ ๔	D	E	A	B	C
Row	ที่ ๕	E	A	B	C	D

จะสังเกตเห็นได้ว่าทุกแถวไม่ว่าแถวตั้งหรือแถวอื่นต่างก็มีครบทุกพันธุ์และไม่ซ้ำกัน แต่เป็นอักษรที่เรียงกันเป็นระเบียบ ไม่มีการสลับให้เป็นไปตามโฉลกของธรรมชาติ จึงจำเป็นต้องทำการสลับ วิธีการสลับกระทำเป็นรายแถวไม่ได้ เนื่องจากจะทำให้ไขว้เขวและรูป Block เสียไปจึงต้องสลับทั้ง Block

× (๒) การสลับ Row หรือสลับ Block ตามอน เรากำหนดเลขประจำ Row เป็น Row ที่ ๑ ที่ ๒ ที่ ๓ ที่ ๔ ที่ ๕ สมมติว่า Random โดยวิธีจับฉลากปรากฏว่าได้ลำดับตัวเลข ดังนี้ ๓, ๕, ๒, ๑, ๔ เราได้สลับ Row โดยเอา Row ที่ ๓ ขึ้นต้นแล้วรองลงมาคือ Row ที่ ๕, ๒, ๑ และ ๔ ตามลำดับ ดังนี้

Column				
1	2	3	4	5
C	D	E	A	B
E	A	B	C	D
B	C	D	E	A
A	B	C	D	E
D	E	A	B	C

× (๓) ทำการสลับ Column ตามตั้งเช่นเดียวกันกับวิธีการสลับ Row สมมติว่า สลับโดยวิธีจับฉลากได้เลขเรียงลำดับก่อนหลัง ดังนี้ ๒, ๑, ๔, ๓, ๕, ให้กำหนดเลขประจำ Column จากแผนผังในข้อ ๒ (มีใช้แผนผังมาตรฐานเดิม) แล้วการดำเนินการเรื่อง Column ใหม่ ต่อจากข้อ ๒ ดังนี้ ขึ้นต้นด้วย Column ที่ ๒ แล้วตามด้วย Column ที่ ๑, ๔, ๓ และ ๕

D C A E B
 A E C B D
 C B E D A
 B A D C E
 E D B A C

ก็เป็นอันเสร็จพิธี นำแผนผังนี้ไปใช้ได้ และในเวลานั้นที่กสอพิจารณาการทดลอง เราจะนำตัวเลขมาจัดเป็นหมู่ เป็นพันธุ์ไม่ได้จะต้องวางไว้ตามแผนผังที่ทำการทดลองกันกว่าจริง ๆ มิฉะนั้นเราจะไม่สามารถตรวจสอบการเปลี่ยนแปลงของกินทั้งสองทิศทางได้ ดังนั้นการวางกำลังสองซึ่งใช้ในการคำนวณก็ต้องมีลักษณะประจำพันธุ์กำกับผลผลิตทุกรายไป

สมมติว่าได้เก็บผลผลิตมาซึ่งเปลี่ยนแปลง ๆ และจดน้ำหนักเป็นกิโลกรัม ดังต่อไปนี้

ตารางกำลังสอง

(คำนวณโดย Coded - 20)

Row	Column										R	R ²
	1		2		3		4		5			
	X	X ²	X	X ²	X	X ²	X	X ²	X	X ²		
1	E/21 1	1	C/24 4	16	A/20 0	0	E/26 6	36	B/16 -4	16	R ₁ ₇	49
2	A/18 -2	4	E/26 6	36	C/21 1	1	B/15 -5	25	D/19 -1	1	R ₂ ₋₁	1
3	C/23 3	9	B/18 -2	4	E/27 7	49	D/20 0	0	A/21 1	1	R ₃ ₉	81
4	B/15 -5	25	A/19 -1	1	D/18 -2	4	C/24 4	16	E/24 4	16	R ₄ ₀	0
5	E/25 5	25	D/20 0	0	B/13 -7	49	A/17 -3	9	C/23 3	9	R ₅ ₋₂	4
C	2	64	7	57	-1	103	2	86	3	43	= 13 ^X	
C ²	4		49		1		4		9			

จะเห็นได้ว่าทั้งสองทิศทางไม่มีผลรายการของพันธุ์ที่เราต้องการเปรียบเทียบเหมือน
 อย่างตาราง Randomized Block ดังนั้นจากตารางนี้จึงไม่ช่วยให้เราหา S.S. for Varie-
 ties ได้เลย คงหาได้เฉพาะรายการที่มีผลปรากฏอยู่ในตารางนี้เท่านั้น

$$\begin{aligned}\text{การหา C.F.} &= (\sum X)^2 / n \\ &= 13^2 / 25 = 6.76\end{aligned}$$

$$\begin{aligned}\text{S.S. Total} &= (X_1^2 + X_2^2 + X_3^2 + X_4^2 + \dots + X_{25}^2) - \text{C.F.} \\ &= (1 + 4 + 9 + 25 + 25 + \dots + 9) - 6.76 \\ &= (64 + 57 + 103 + 86 + 43) - 6.76 \\ &= 346.24\end{aligned}$$

$$\begin{aligned}\text{S.S. Row} &= \frac{R_1^2 + R_2^2 + R_3^2 + R_4^2 + R_5^2}{c} - \text{C.F.} \\ &= \frac{49 + 1 + 81 + 0 + 4}{5} - 6.76 \\ &= \frac{135}{5} - 6.76 \\ &= 20.24\end{aligned}$$

$$\begin{aligned}\text{S.S. Column} &= \frac{C_1^2 + C_2^2 + C_3^2 + C_4^2 + C_5^2}{r} - \text{C.F.} \\ &= \frac{4 + 49 + 1 + 4 + 9}{5} - 6.76 \\ &= \frac{67}{5} - 6.76 \\ &= 6.64\end{aligned}$$

C = ผลรวมของแต่ละ Column

R = ผลรวมของแต่ละ Row

c = จำนวน Column

r = จำนวน Row

เมื่อเราทราบรายการทั้งสองทิศทางแล้ว แต่เรายังไม่ทราบรายการของ Varieties ซึ่งเป็นความหมายสำคัญ ก็จำเป็นต้องหาทางออกยอดของแต่ละ Varieties ให้ได้ โดยที่ ตารางใหม่เพื่อจัดพันธุ์แต่ละพันธุ์ไว้ในช่องเดียวกัน

	A	B	C	D	E
	0	-4	4	1	6
	-2	-5	1	-1	6
	1	-2	3	0	7
	-1	-5	4	-2	4
	-3	-7	3	0	5
V	-5	-23	15	-2	28
V ² =	25	529	225	4	784

$$\begin{aligned}
 \text{S.S. Varieties} &= \frac{V_1^2 + V_2^2 + V_3^2 + V_4^2 + V_5^2}{r} - \text{C.F.} \\
 &= \frac{25 + 529 + 225 + 4 + 784}{5} - 6.76 \\
 &= \frac{1567}{5} - 6.76 \\
 &= 306.64
 \end{aligned}$$

ส่วนประกอบในรายการของ Latin Square Design มีดังนี้

Total = Row + Column + Varieties + สิ่งอื่น ๆ (Error)

หรือ Error = Total - (Row + Column + Varieties)

$$\begin{aligned}
 \text{S.S. Error} &= 346.24 - (20.24 + 6.64 + 306.64) \\
 &= 12.72
 \end{aligned}$$

Sources of Difference	D.F. (n-1)	S.S.	Variance S.S./D.F.	F ratio Calculated	F ratio Table
Total	24	346.24			
Column	4	6.64	1.66		
Row	4	20.24	5.06	4.77+	3.26 5 %
Varieties	4	306.64	76.66	72.32++	5.41 1 %
Error	12	12.72	1.06		

Calculated F ratio for Row = $5.06 \div 1.06 = 4.77$

Calculated F ratio for Varieties = $76.66 \div 1.06 = 72.32$

การคำนวณในตาราง Analysis ก็กระทำอย่างเดียวกันกับแบบ Randomized Block ทุกอย่าง ตลอดไปจนถึงการหา F ratio ดังนั้นถ้าหากสงสัยประการใดให้พลิกกลับไปตรวจดูและทำความเข้าใจตาราง Analysis of Variance ของแบบ Randomized Block Design อีกครั้งหนึ่ง

จะเห็นได้ว่านอกจาก F ratio ของพันธุ์แสดงว่าพันธุ์มีความแตกต่างกันจริง เพราะสูงเกินระดับจากตาราง F มากแล้ว ดินทางคานนอน (Row) ก็แสดงว่ามีความสมบูรณ์ ไม่ค่อยสม่ำเสมอเกินไป เพราะ F ratio ของ Row สูงเกินระดับ ๕ % แต่ยังไม่ถึงระดับ ๑ % แต่การที่เรารู้ทิศทางของความแตกต่างในเรื่องความอุดมสมบูรณ์ของดินนั้น แม้ว่าจะไม่เป็นประโยชน์ในการสรุปผลการค้นคว้า แต่ก็ช่วยให้เราได้ศึกษาทราบลักษณะของดินและทิศทางการเปลี่ยนแปลงของดินเพื่อประโยชน์ในการวางแผนการทดลองคราวต่อไป จะได้กระทำได้อย่างถูกต้องตามทิศทางและระดับ เช่นการวาง Block ให้ทุก ๆ Block หรือทุกซ้ำได้มีโอกาสกระทบความอุดมสมบูรณ์ของดินคล้ายคลึงกันโดยทั่วถึง เมื่อปรากฏแล้วว่าพันธุ์มีความแตกต่างกันจริง เราก็ดำเนินการเปรียบเทียบเป็นรายพันธุ์เพื่อทราบรายละเอียดในการเปรียบเทียบต่อไป

วิธีการที่กระทำได้โดยการไ้ระดับ L.S.D. ดังได้กล่าวมาแล้วในเรื่องของ

$$\begin{aligned}
 \text{L.S.D.} &= t \cdot \sqrt{\frac{2 \text{ Variance of Error}}{n - 2_{\text{error}}}} \\
 \text{สำหรับเปรียบเทียบ} &= t \cdot \sqrt{\frac{2 \times 1.06}{5}} \\
 \text{ผลต่างระหว่างพันธุ์} &= t \cdot \sqrt{0.424} \\
 &= 0.6511 t
 \end{aligned}$$

จาก t table เมื่อ D.F. = 12 (จาก D.F. error)

ค่าของ t ระดับ 5 % = 2.179

ค่าของ t ระดับ 1 % = 3.055

L.S.D. ระดับ 5 % = 2.179 X 0.6511 = 1.418

L.S.D. ระดับ 1 % = 3.055 X 0.6511 = 1.989

ลำดับผลเฉลี่ยจากสูงไปต่ำ *4.18 1.989 1.418 0.6 non. 3.6 2.6 3.4*

ผลเฉลี่ยของผลผลิต พันธุ์ E	$= \frac{128}{5} = 25.6$	Difference between E & C 2.6 ++
ผลเฉลี่ยของผลผลิต พันธุ์ C	$= \frac{115}{5} = 23.0$	Difference between C & D 3.4 ++
ผลเฉลี่ยของผลผลิต พันธุ์ D	$= \frac{98}{5} = 19.6$	Difference between D & A 0.6 non.
ผลเฉลี่ยของผลผลิต พันธุ์ A	$= \frac{95}{5} = 19.0$	Difference between A & B 3.6 ++
ผลเฉลี่ยของผลผลิต พันธุ์ B	$= \frac{77}{5} = 15.4$	

จะเห็นได้ว่า เมื่อเอาผลต่างระหว่างผลเฉลี่ยของผลผลิตแต่ละคู่ไปเทียบกับระดับ L.S.D. แล้วปรากฏว่าพันธุ์ E ให้ผลสูงที่สุดอย่างแน่นอน โดยถือเป็นนัยสำคัญยิ่งในทางสถิติ เพราะพันธุ์ E ให้ผลผลิตสูงกว่าพันธุ์ C ซึ่งรองลงมาถึง ๒.๖ สูงกว่าระดับ L.S.D. ๑ % คือ ๑.๔๘๘ และรองลงมาก็เช่นเดียวกัน คือแสดงความแตกต่างอย่างเด่นชัด นอกจากพันธุ์ D และพันธุ์ A ซึ่งต่างกันเพียง ๐.๖ ต่ำกว่า ๑.๔๘๘ ซึ่งเป็นระดับ ๕ % ของ L.S.D. อันเป็นระดับต่ำที่สุดที่จะถือว่าต่างกันได้ ดังนั้นพันธุ์ D และ A ต่างกันเพียง ๐.๖ จึงถือว่าไม่ต่างกัน เราก็อาจสรุปได้ว่าพันธุ์ E ให้ผลสูงที่สุด พันธุ์ C เป็นที่ ๒ พันธุ์ D และพันธุ์ A เป็นที่ ๓ เหมือน ๆ กันเพราะไม่มีความแตกต่างที่จะถือเป็นนัยสำคัญทางสถิติได้ ส่วนพันธุ์ B ให้ผลต่ำที่สุด

ในการค้นคว้าทดลองที่ได้กล่าวมาแล้วทั้งสองแบบนี้ เราจะเห็นได้ว่าหัวใจสำคัญที่สุดที่เราต้องเกี่ยวข้องกันก็คือสถิติตัวเลขหรือ Data ในการที่เราลงทุนลงแรง และเสียเวลา

แรมเดือน เพื่อรอนกว่าพืชหรือสัตว์ที่ใช้ในการค้นคว้าจะเจริญเติบโตแสดงผลให้ปรากฏตามความ
ต้องการ เมื่อได้รวบรวมตัวเลขจากผลที่ปรากฏแล้ว เกิดอุบัติเหตุทำให้ตัวเลขบางจำนวนขาดหาย
หรือเลอะเลือนไป บางทีพืชหรือสัตว์บางส่วนอาจเป็นอันตรายไปก่อนที่จะวัดหรือเก็บ Data
ได้ ก็จะทำให้ส่วนที่ที่เหลือนั้นต้องเสียหายไปเลย และจำเป็นต้องรอเวลาที่เหมาะแก่การค้นคว้าใหม่
ต่อไป อันนี้ว่าต้องเสียทุนเสียแรงเสียเวลาไปอย่างมากมาจากการเสียตัวเลขเพียงไม่กี่จำนวน
แต่นั่นมันเป็นโชคที่วิธีการคำนวณเพื่อหาค่าของตัวเลขที่ขาดหายไปนั้นหาทดแทนได้ แทนที่การ
ทดลองค้นคว้านั้นจะต้องได้รับความเสียหายทั้งหมด แต่ทั้งนี้ขอให้ออกใจแก่นักศึกษาคนกว่าไว้
ประการหนึ่งว่า ไม่ควรใช้วิธีการนี้ไปในทางที่ผิด เช่นทำให้เกียจคร้านหรือตั้งอยู่ในทางประมาท
เช่นเห็นสัตว์เลี้ยงหรือสัตว์อื่น ๆ ทำลายพืชในแปลงทดลองก็เฉยเสียเพราะคิดว่า แม้จะขาดตัวเลข
ไปไม่กี่จำนวนก็คิดคำนวณเลขเอาใหม่ได้ง่าย ๆ ซึ่งนับว่าเป็นการเข้าใจผิด วิธีการนี้ช่วยแก้ไข
อุปสรรคสุดวิสัยเป็นสำคัญ แต่ไม่ควรนำไปใช้ช่วยให้เกิดความเกียจคร้านหรือความประมาท
สูตรที่ใช้คำนวณหาค่าของตัวเลขที่ขาดหายไปมีอยู่ทั้งแบบ Randomized Block และแบบ Latin
Square โดยเฉพาะของแต่ละแบบ

ปัญหา Missing value Randomized Block Design ใช้สูตรดังนี้คือ

$$X = \frac{tT + bB - G}{(t - 1)(b - 1)}$$

X = ตัวเลขที่ขาดหายไปหรือ Missing data

T = ผลรวมของตัวเลขใน Treatment เดียวกันกับ Treatment ที่ตัวเลขขาดหายไป

t = จำนวน Treatment ทั้งหมด

B = ผลรวมของตัวเลขใน Block เดียวกันกับ Block ที่ตัวเลขขาดหายไป

b = จำนวน Blocks ทั้งหมด

G = ผลรวมของตัวเลขทั้งหมดในการค้นคว้า

ตัวอย่าง

Varieties

Block	A	B	C	D	Total
I	18	25	16	23	82
II	16	23	18	23	80
III	19	24	17	25	85
IV	21	(24)	19	24	64 (B)
V	20	25	18	(22)	85
	94	97 (T)	88	117	396

สมมติว่าตัวเลขในของพันธุ์ B และ Block IV คือ ๒๔ นั้นขาดหายไป

$$T = 25 + 23 + 24 + 25 = 97$$

$$t = 4$$

$$B = 21 + 19 + 24 = 64$$

$$b = 5$$

$$G = 396$$

$$\begin{aligned}
 \text{สูตร} \quad X &= \frac{tT + bB - G}{(t-1)(b-1)} \\
 &= \frac{(4 \times 97) + (5 \times 64) - 396}{(4-1)(5-1)} \\
 &= \frac{312}{12} \\
 &= 26
 \end{aligned}$$

เนื่องจากผลรวมเดิม (G) คือ ๓๙๖ นั้นผิดไปมากเพราะขณะนั้นมีตัวเลขที่ขาดหายไปจำนวนหนึ่ง จึงใช้ ๒๖ ซึ่งเป็นผลของการคำนวณได้ในขณะที่ G ต่ำกว่าความจริงมากมายมาก G ให้ใกล้เคียงความจริงเข้า แล้วจึงคำนวณอีกทีหนึ่ง

$$G = 396 + 26 = 422 \quad \text{นอกนั้นคงเดิม}$$

$$\begin{aligned} X &= \frac{(4 \times 97) + (5 \times 64) - 422}{(4 - 1)(5 - 1)} \\ &= \frac{286}{12} \\ &= 23.8 \end{aligned}$$

ปกติเศษจาก ๒๓.๘ ให้เป็นเลขฉนวน ๆ อย่างเลขจำนวนอื่น ๆ ดังนั้นจึงถือ ๒๓.๘ เป็น ๒๔ เท่ากับจำนวนเดิม แม้ว่าถาหากจะเคลื่อนไปจากจำนวนเดิมสัก ๑ หรือ ๒ ก็ไม่เป็นไร เพราะตัวเลขในพันธเดียวกันยังมากน้อยไม่เหมือนกันและอาจนึกเพี้ยนไปได้บางแล้วแต่สภาพของแฟกเตอร์แวดล้อมจึงไม่ถือเป็นสำคัญ แต่อย่างนี้เป็นเพียงการพิสูจน์และยกตัวอย่างให้ดู ถ้าตัวเลขขาดหายไปจริง ๆ เราก็มไม่มีทางทราบค่าเดิมของตัวเลขนั้นได้ เพราะถาทราบก็ไม่ต้องคำนวณหาให้เสียเวลา

ในกรณีที่ตัวเลขขาดหายไปถึง ๒, ๓ จำนวน ค่าของผลรวมยกใหญ่ (G) ก็จะต้องผิดไปมาก ดังนั้นการคำนวณจึงต้องทำกลับไปกลับมาหลาย ๆ ครั้ง โดยการคำนวณค่าที่ละจำนวน แล้วนำค่าที่หาได้ตอนแรก ๆ ถือเป็นค่าชั่วคราวนำไปแก้ผลรวม G ให้ค่อยๆ เดือนเข้าที่ จนกระทั่งเห็นว่า Missing value ที่เรากำหนดในระยะหลังอยู่ตัวคงที่ แม้เปลี่ยนก็เพียงทศนิยมหรือไม่ก็เปอร์เซ็นต์ของค่าตัวเองกันมีว่าใช้ได้ เช่นถาหากมีตัวเลขขาดหายไป ๒ จำนวน ก็นอกจากจำนวนที่กล่าวมาแล้วยังมีอีกจำนวนหนึ่งใน Block ที่ ๕ พันธ D คือ ๒๒ ค่า G ในขณะที่เลขขาดหายไปถึง ๒ จำนวนจะมีค่าเพียง ๓๗๔

เรากำหนดตัวที่อยู่ในพันธ B ก่อน คือ

$$X_1 = \frac{(4 \times 97) + (5 \times 64) - 374}{(4 - 1)(5 - 1)} = 27.8 \quad \text{ถือเป็น ๒๘}$$

เอา ๒๘ ใส่ไว้ในช่อง B ชั่วคราว เพื่อแก้ค่า G ก่อน แล้วคำนวณหาตัวเลขในช่อง G ตอนนั้นค่า G จะเพิ่มขึ้นอีก ๒๘

การหา X ตัวที่ ๒

$$T = 23 + 23 + 25 + 24 = 95$$

$$B = 20 + 25 + 18 = 63$$

$$X_2 = \frac{(4 \times 95) + (5 \times 63) - 402}{(4 - 1)(5 - 1)} = 24$$

นำ ๒๔ ไปใส่ในช่อง D ซ้ำคราว เพื่อแก้ G แล้วหา X_1 ใหม่ ดังนั้น G จะเพิ่มขึ้น
 $G = 402 + 24 = 426$ ส่วนค่า X_1 ที่หาไว้ตอนแรก ลบทิ้งไป คงหาใหม่ ดังนี้

$$X_1 = \frac{(4 \times 97) + (5 \times 64) - 426}{(4 - 1)(5 - 1)} = 24.1 \text{ หรือ } 24..$$

แล้วจึงนำ ๒๔ นี้แทน X_1 ใหม่ หาค่า G ใหม่ แล้วคำนวณหา X_2 อีกครั้งหนึ่ง เมื่อเห็นผลการคำนวณในครั้งหลัง ๆ ไม่ค่อยช่วยแก้ไข Missing value มีค่าเปลี่ยนไปได้อีกเท่าใดนัก แม้แต่ค่าของเลขตัวอื่น ๆ ใน Treatment เดียวกันก็มีการเหลื่อมล้ำมากกว่า ก็หยุดการคำนวณและใช้ผลการคำนวณนั้นแทน Missing data ได้

๑) การใช้สูตร Missing Value for Latin Square Design

k = จำนวน Row หรือ Column หรือ Treatment (เนื่องจากใน Latin Square จำนวนของทั้งสามอย่างนี้เท่ากัน)

R = ผลรวมของตัวเลขใน Row เดียวกันกับที่ตัวเลขขาดหายไป

C = ผลรวมของตัวเลขใน Column เดียวกันกับที่ตัวเลขขาดหายไป

T = ผลรวมของตัวเลขใน Treatment เดียวกันกับที่ตัวเลขขาดหายไป

G = ผลรวมของตัวเลขทั้งหมดใน Experiment

$$X = \frac{k(R + C + T) - 2G}{(k-1)(k-2)}$$

ตัวอย่าง

		<u>Column</u>				
		1	2	3	4	5
<u>Row</u>	1	D/21	C/24	A/20	E/26	B/16
	2	A/18	E/26	C/21	B/15*	D/19
	3	C/23	B/18	E/27	D/20	A/21
	4	B/15	A/19	D/18	C/24	E/24
	5	E/25	D/20	B/13	A/17	C/23

สมมุติ Row ที่ 2 Column ที่ 4 พันธุ์ B ตัวเลขขาดหายไป

$$k = 5$$

$$R = 18 + 26 + 21 + 19 = 84$$

$$C = 26 + 20 + 24 + 17 = 87$$

$$T = 16 + 18 + 15 + 13 = 62$$

$$G = 21 + 18 + 23 + 15 + 25 + \dots + 24 + 23 = 498$$

$$X = \frac{5(84 + 87 + 62) - (2 \times 498)}{(5-1)(5-2)} = 14.1 \text{ หรือ } 14$$

ถ้าได้นำ ๑๔ ที่คำนวณได้นำไปเพิ่มค่าของ G เพื่อแก้ไขค่าจำนวนเต็มเข้าไปอีก แล้วคำนวณซ้ำอีกทีหนึ่งก็จะได้อีกยิ่งขึ้น เช่นเดียวกันในการคำนวณ Missing value ในวิธีการ Randomized Block Design เหมือนกัน และถ้าหากว่าเกิดมีตัวเลขขาดหายไปอีก ๒ - ๓ จำนวนก็ใช้วิธีการคำนวณหาได้ไปที่ละจำนวน แล้วคำนวณกลับเพื่อแก้ค่าของผลรวมให้ใกล้เคียงความจริงยิ่งขึ้น เช่นเดียวกันกับวิธีการใน Randomized Block System ซึ่งได้ยกตัวอย่างมาแล้ว

สหสัมพันธ์

(CORRELATION)

เรื่องที่แล้ว ๆ มาทั้งหมดเป็นเรื่องที่เกี่ยวกับการศึกษาความสัมพันธ์เพียงลักษณะเดียว หรือเป็นการศึกษาชุดตัวเลขที่แสดงผลเพียงลักษณะเดียว เช่นการค้นคว้าเปรียบเทียบผลผลิต หรือการค้นคว้าเกี่ยวกับการเจริญเติบโต เป็นต้น ต่อไปนี้เราจะทำการพิจารณาเพื่อค้นคว้าปรากฏการณ์สองลักษณะควบกันไป ทั้งนี้เนื่องจากว่า การเปลี่ยนแปลงของลักษณะหนึ่งนั้น อาจมีผลทำให้อีกลักษณะหนึ่งเปลี่ยนแปลงไปด้วย ถ้าลักษณะคู่ใดมีความสัมพันธ์กันแล้ว เราก็เรียกสองลักษณะนั้นว่ามี "สหสัมพันธ์" (Correlation) กัน การพิจารณาสองลักษณะควบกันไปในั้น จะขอยกตัวอย่างเช่น ความยาวของฝักกับจำนวนเมล็ดในฝัก ความยาวของช่อดอกกับจำนวนดอก ความหนาของกลีบดอกกับความหนาแน่นในการบาน ความสูงกับน้ำหนัก ความชื้นภายในเมล็ดกับเปอร์เซ็นต์ความงอก ความหนากับความหนาแน่นในการขัดสีของเมล็ดข้าว จำนวนปุ๋ยกับความเจริญเติบโต ฯลฯ แต่เนื่องจากสองลักษณะที่เราทำการตรวจสอบกันไปในั้นมิได้มีสหสัมพันธ์กันเสมอไป จึงจำเป็นต้องทำการตรวจ นอกจากนั้นแต่ละลักษณะต่างก็ได้รับการกระทบกระเทือนจากแฟกเตอร์แวดล้อมนานาประการ ดังได้ศึกษามาแล้ว ดังนั้นเมื่อเอาปรากฏการณ์ทั้งสองลักษณะเข้ามาพิจารณาร่วมกัน ก็ยังทำให้เกิดความผิดเพี้ยนไขว้เขวจนยากที่จะเข้าใจได้ง่าย จึงจำเป็นต้องใช้วิธีการทางสถิติเป็นเครื่องมือเพื่อค้นคว้าเอาปรากฏการณ์นั้นมาตั้งเป็นกฎเกณฑ์ใช้ประโยชน์ต่อไป

ผลของการตรวจสอบปรากฏการณ์ของสองลักษณะควบกันนี้ อาจเป็นไปได้ ๓ สภาวะด้วยกันคือ

๑. ไม่มีสหสัมพันธ์ หมายความว่า การเปลี่ยนแปลงของสองลักษณะที่เรานำมาพิจารณาร่วมกันนั้น ไม่มีความเกี่ยวเนื่องกัน ตัวอย่างเช่น เราจะค้นคว้าทดลองว่า ถ้าช่อดอกยาวจะมีจำนวนดอกมากขึ้นจริงหรือไม่ สมมุติว่าผลการทดลองค้นคว้าแสดงว่าไม่มีสหสัมพันธ์ (Uncorrelation) ก็หมายความว่าแม้ช่อดอกจะยาวขึ้นจำนวนดอกอาจมากขึ้นหรือน้อย หรืออาจคงเดิมก็ได้ ไม่อาจจะถือเป็นกฎเกณฑ์อะไรได้แน่นอน

๒. มีสหสัมพันธ์กันในทางเดียวกันหรือทางบวก (Positive correlation) หรือกล่าวอีกนัยหนึ่งว่า มีการเปลี่ยนแปลงเป็นส่วนตรง คือ เมื่อลักษณะหนึ่งเพิ่มขึ้น อีกลักษณะหนึ่งก็เพิ่มขึ้นด้วย เช่นสมมติว่าจากลิ้มคอกหนาขึ้นก็ยิ่งบานทนทานขึ้นด้วย หรือถ้าฝักยิ่งมีขนาดยาว จำนวนเมล็ดก็ยิ่งมากขึ้น แต่ทั้งนี้จะต้องมีการตรวจสอบและยืนยันทางสถิติด้วยว่า สองลักษณะนั้นมีสหสัมพันธ์กันในทางบวกจริง จึงจะถือเป็นหลักเกณฑ์เพื่อใช้ประโยชน์ได้

๓. มีสหสัมพันธ์กันในทางตรงกันข้ามหรือทางลบ (Negative Correlation) หรืออาจกล่าวได้ว่าการเปลี่ยนแปลงเป็นส่วนกลับซึ่งกันและกัน หมายความว่า เมื่อลักษณะหนึ่งเพิ่มขึ้นอีกลักษณะหนึ่งจะลดลง ตัวอย่างเช่นความชื้นในเมล็ดขณะเก็บรักษากับเปอร์เซ็นต์ความงอก ถ้าผลการค้นคว้าแสดงว่ามีสหสัมพันธ์ในทางลบ ก็แสดงว่าถ้าเก็บเมล็ดที่มีความชื้นสูงขึ้น จะทำให้เปอร์เซ็นต์ความงอกลดลง เป็นต้น

จากการที่ได้ผ่านงานค้นคว้ามาแล้ว ผู้เขียนพบว่ามีสหสัมพันธ์บางเรื่องที่เกิดการวิเคราะห์ทางสถิติอาจแสดงออกทางบวกได้เหมือนกัน หรืออาจจะเป็นทางลบก็ได้ แต่ผลงานก็ลงเอยอย่างเดียวกัน ทั้งนี้ย่อมขึ้นอยู่กับการวางแผนทางของผู้ดำเนินงานเองเป็นสำคัญ ผู้เขียนจึงใคร่จะกล่าวไว้ให้เป็นข้อสังเกตว่า ถ้าลักษณะนั้นมีสภาพเป็นองค์ประกอบสองส่วนซึ่งอาจวัดได้ทั้งสองทาง เช่นชาวสารประกอบด้วยชาวเต็มเมล็ดและชาวหัก ถ้าเราจะหาสหสัมพันธ์ระหว่างความชื้นของเมล็ดข้าวเปลือกที่นำมาสีกับเปอร์เซ็นต์ผลการสี หากเราจะหาสหสัมพันธ์ระหว่างความชื้นกับเปอร์เซ็นต์ชาวหัก ปรากฏผลว่ามีสหสัมพันธ์ทางบวก คือยิ่งเมล็ดข้าวชื้นมาก เมื่อนำไปสียิ่งหักมากขึ้น แต่สหสัมพันธ์อาจเป็นทางลบได้ ถ้าจะเลือกตรวจเปอร์เซ็นต์ชาวเต็มเมล็ด เพราะยิ่งข้าวเปลือกชื้นมาก เมื่อนำไปสีจะได้เปอร์เซ็นต์ชาวหักน้อยลง แต่ผลงานก็สรุปได้อย่างเดียวกัน ทั้งนี้แล้วแต่เราจะใช้วิธีตรวจวัดชาวเต็มเมล็ด หรือชาวหักในการพิจารณาผล

ก่อนที่จะทำการหาสหสัมพันธ์ระหว่างสองลักษณะว่าเป็นลักษณะ X และลักษณะ Y เพื่อนำไปใช้แทนลักษณะใดที่เราต้องการจะค้นคว้าก็ได้ เราทราบมาแล้วว่าในการค้นคว้าลักษณะเดียว แต่ลักษณะก็มีความคลาดเคลื่อนในตัวเองอยู่แล้ว ซึ่งเราวัดได้โดยการหา

Deviation(D) ลักษณะ X ก็มี D_x และลักษณะ Y ก็มี D_y ของตัวเอง เนื่องจากเราจะต้องหาความสัมพันธ์ร่วมกัน เราจึงจำเป็นต้องคำนวณหาว่าอัตราส่วนที่แสดงการเปลี่ยนแปลงคลาดเคลื่อนร่วมกันทั้งหมดนั้น เป็นที่เท่าของความคลาดเคลื่อนของตนเองโดยอิสระ ถ้าหากความคลาดเคลื่อนร่วมกันมีมากมายหลาย ๆ เท่า ก็เป็นที่เชื่อถือได้ว่าทั้งสองลักษณะนั้นมีความสัมพันธ์ร่วมกันจริง แต่ถ้าน้อยเท่า ความสัมพันธ์ร่วมกันก็จะถูกความคลาดเคลื่อนของลักษณะเดียวลบลงไปหมด อัตราส่วนที่กล่าวมานี้เราเรียกว่า Correlation coefficient ซึ่งใช้สัญลักษณ์ r แทน จะเห็นได้จากตอนแรก ๆ ว่า เราหา D แล้วยกกำลังสองเป็น D^2 หรือคือ $D \times D$ นั่นเอง เมื่อรวมผล D^2 ทั้งหมดเราเรียกว่า Sum of square ดังนั้นลักษณะ X เราหา D_x ยกกำลังสอง $D_x \cdot D_x = D_x^2$

ลักษณะ Y เราหา D_y ยกกำลังสอง $D_y \cdot D_y = D_y^2$

ลักษณะ x กับ y เราหา D_x และ D_y แล้วคูณ $D_x \cdot D_y$

คล้ายยกกำลังสอง แต่ D_x คูณ D_y ไม่ใช่ยกกำลังสองจริง ๆ แต่มีความมุ่งหมายคล้ายยกกำลังสอง หากเป็นผลคูณ (product) ดังนั้นเมื่อรวม $D_x \cdot D_y$ ทั้งหมดจึงไม่เรียกว่า Sum of square แต่เรียกว่า Sum of Product x.y หรือย่อว่า $S.P_{xy}$ เป็นเครื่องหมายใช้สำหรับ x หรือ y แยกอย่างเดี่ยว หากแต่เป็นความสัมพันธ์ร่วมกันทั้ง x และ y

ตามที่ได้อธิบายมาแล้วถึงความหมายของ Coefficient "r" จึงทำให้สามารถสรุปได้ว่า คืออัตราส่วนระหว่างความสัมพันธ์ของ X และ Y หาค่ายสิ่งที่เกิดจาก x) อย่างเดี่ยว และ y) อย่างเดี่ยว แต่เนื่องจากเราจำเป็นต้องกล่าวโดยเฉลี่ย จึงจำเป็นต้องหารด้วย

Degree of Freedom ควบ

$$\text{ดังนั้น } r = \frac{D_x D_y}{D.F. \cdot S.D._x \cdot S.D._y}$$

$$\text{หรือ } = \frac{S.P_{xy}}{D.F. \cdot \sqrt{\frac{\sum d_x^2}{D.F._x}} \cdot \sqrt{\frac{\sum d_y^2}{D.F._y}}}$$

คือจำนวน D.F. $X - 1$ คือจำนวน $Y - 1$

และ D.F. ที่อยู่ข้างหน้าคือ (จำนวนคู่ของ x กับ y) $- 1$

ยกตัวอย่างดังต่อไปนี้คือ สมมุติว่า เราต้องการหาความสัมพันธ์ระหว่างจำนวนต้นกับผลผลิตเพื่อต้องการทราบว่า ถ้าจำนวนต้นมากขึ้น ผลผลิตจะเพิ่มขึ้นจริงหรือไม่ เรามีข้าวโพดอยู่ ๑๐ แปลง นับจำนวนต้นของแต่ละแปลงแล้วก็เก็บผลผลิตของแต่ละแปลงมาซึ่ง ควรจำไว้เลยว่า เมื่อนับจำนวนต้นแล้วจะต้องจดน้ำหนักของผลผลิตไว้คู่กันเสมอ จะสลับคู่กันไปได้ ในที่นี้สมมุติว่า กำหนดให้ลักษณะผลผลิตเป็น X และจำนวนต้นเป็น Y

$X = \text{Yield}$

$Y = \text{number of plants}$

X	D $(X - M_x)$	D^2_x	Y N_2	D_y $(Y - M_y)$	D^2_y	$D_x D_y$
28	-1	1	24	-2	4	2
25	-4	16	21	-5	25	20
29	0	0	27	1	1	0
27	-2	4	25	-1	1	2
26	-3	9	23	-3	9	9
32	3	9	28	2	4	6
31	2	4	29	3	9	6
28	-1	1	22	-4	16	4
30	1	1	30	4	16	4
34	5	25	31	5	25	25
290		$\Sigma D_x^2 = 70$	260		$\Sigma D_y^2 = 110$	78

$$M_x = \frac{290}{10} = 29.0$$

$$M_y = \frac{260}{10} = 26.0$$

$$\begin{aligned}
 S.D._x &= \pm \sqrt{\frac{\sum d_x^2}{D.F._x}} \\
 &= \pm \sqrt{\frac{70}{10 - 1}} = \pm 2.789 \\
 S.D._y &= \pm \sqrt{\frac{\sum d_y^2}{D.F._y}} \\
 &= \pm \sqrt{\frac{110}{10 - 1}} = \pm 3.496
 \end{aligned}$$

Sum of D.D. or S.P._{xy} = 78

$$\begin{aligned}
 r &= \frac{S.P._{xy}}{D.F. S.D._x S.D._y} \\
 &= \frac{78}{9 \times 2.789 \times 3.496} \\
 &= 0.889
 \end{aligned}$$

แต่วิธีการนี้เป็นวิธีการที่ยาว เพราะจะต้องหา S.D._x และ S.D._y ที่ละครั้ง ต้องถอดกำลังสองถึงสองครั้ง ถ้าหากเราจะหันกลับไปพิจารณาถึงวิธีการหา S.D. แต่ดั้งเดิมมา วิธีที่ไม่ต้องหา Mean แล้วเอา Mean ไปลบ X ที่ละตัว ๆ ให้เสียเวลา โดยเล็งไปใช้วิธีการยกกำลังสองหรือเอา Mean เป็น 0 ทำให้สั้นและสะดวกกว่า เราอาจนำวิธีนั้นมาใช้ดัดแปลงแก้ไขวิธีการคำนวณเรื่องนี้ได้เช่นเดียวกัน

ก่อนอื่นขอให้พิจารณาเสียก่อนว่า ในสูตรเดิมของการหา S.D. หรือ S.E. นั้นเกี่ยวข้องกับเรื่องยกกำลังสองก็คุณตัวเอง เช่น $d_x^2 = d_{xy} d_x$ เมื่อเปลี่ยนเป็น x^2 ก็คือ x คูณด้วย x นั่นเอง

เมื่อมาถึงเรื่องหนึ่งซึ่งเป็นเรื่อง X กับ Y ดังนั้นแทนที่จะเป็น $\sum d^2 = \sum X^2 - \frac{(\sum X)^2}{n}$

ก็เป็น $\sum d_x \cdot d_y = \sum X \cdot Y - \frac{(\sum X)(\sum Y)}{n}$

๑. เราจึงเปลี่ยนแปลงสูตร

$$r = \frac{\sum d_x \cdot d_y}{D.F. \cdot S.D._x \cdot S.D._y}$$

$$= \frac{\sum X \cdot Y - (\sum X)(\sum Y)/n}{D.F. \cdot S.D._x \cdot S.D._y}$$

$$D.F. \cdot \sqrt{\frac{\sum X^2 - (\sum X)^2/n}{D.F._x}} \cdot \sqrt{\frac{\sum Y^2 - (\sum Y)^2/n}{D.F._y}}$$

๒. เมื่อเราเอา D.F. ให้เขา

$$\text{เครื่องหมายเหมือนกัน} = \frac{\sum X \cdot Y - (\sum X)(\sum Y)/n}{D.F. \cdot S.D._x \cdot S.D._y}$$

$$\sqrt{\frac{\sum X^2 - (\sum X)^2/n}{D.F._x}} \cdot \sqrt{\frac{\sum Y^2 - (\sum Y)^2/n}{D.F._y}}$$

$$D.F. \text{ ข้างบนและข้างล่างตัดกัน} = \frac{\sum X \cdot Y - (\sum X)(\sum Y)/n}{D.F. \cdot S.D._x \cdot S.D._y}$$

$$\text{หมดเพราะค่าของ D.F. เท่ากันทุกตัว} \sqrt{\frac{\sum X^2 - (\sum X)^2/n}{D.F._x}} \cdot \sqrt{\frac{\sum Y^2 - (\sum Y)^2/n}{D.F._y}}$$

อาจเขียนย่อ ๆ ได้ว่า

$$r = \frac{S.P._{xy}}{\sqrt{S.S._x \cdot S.S._y}}$$

ถ้าหากเราจะตั้งตารางคำนวณโดยใช้สูตรนี้ก็สามารถกระทำได้ ดังต่อไปนี้

X	X ²	Y	Y ²	X.Y.
28	784	24	576	672
25	625	21	441	525
29	841	27	729	783
27	729	25	625	675
26	676	23	529	598
32	1024	28	784	896
31	961	29	841	899
28	784	22	484	616
30	900	30	900	900
34	1156	31	961	1054
290	8480	260	6870	7618

$$\begin{aligned}
 r &= \frac{\sum X.Y. - (\sum X)(\sum Y)/n}{\sqrt{\sum X^2 - (\sum X)^2/n} \cdot \sqrt{\sum Y^2 - (\sum Y)^2/n}} \\
 &= \frac{7618 - (290 \times 260)/10}{\sqrt{8480 - \frac{(290)^2}{10}} \times \sqrt{6870 - \frac{(260)^2}{10}}} \\
 &= \frac{78}{\sqrt{70 \times 110}} = 0.889
 \end{aligned}$$

เราอาจใช้วิธี Code เข้าช่วยเพื่อตัดปัญหาเรื่องเลขมากได้ในเมื่อไม่มีเครื่องคิดเลข และเราอาจใช้ตัวเลข Code สำหรับ X คนละตัวกับสำหรับ Y ข้อแต่เพียงว่าในชุด X เกี่ยวกันต้อง Code ด้วยตัวเดียวกัน และในชุด Y ก็เช่นเดียวกัน ทั้งนี้เหตุผลอยู่ว่าการหาสหสัมพันธ์เป็นการหาตำแหน่งความสูงต่ำของค่าตัวเลขแม้จะ Code ด้วยอะไรก็ตาม ถ้าในชุดเดียวกัน Code ด้วยตัวเดียวกัน หรือพูดง่าย ๆ ว่าตัดออกเท่า ๆ กัน ลำดับความสูงต่ำภายในเลขชุดนั้นก็ยังคงเดิม ดังนั้นในกรณีค่าตัวเลขชุด X กับชุด Y ต่างกันมาก ๆ เช่นชุด X มีค่าเรือนสิบ แต่ชุด Y มีค่าเรือนร้อยหรือเรือนพัน ก็จำเป็นต้องหา code ด้วยตัวเลขที่เหมาะสมแก่ชุดนั้น ๆ เพื่อให้แต่ละชุดได้ค่าตัวเลขแต่ละจำนวนน้อยที่สุด จะได้เป็นประโยชน์สวดกแก่การคำนวณมากที่สุด ในที่นี้จะเห็นได้ว่า ค่ากลาง ๆ ของชุด X ประมาณ ๓๐ แต่ค่ากลาง ๆ ของชุด Y ประมาณ ๒๕ จึง code ชุด X ด้วย ๓๐ และ code ชุด Y ด้วย ๒๕

X	X coded by - 30	X ²	Y	Y coded by - 25	Y ²	X.Y.
28	-2	4	24	-1	1	2
25	-5	25	21	-4	16	20
29	-1	1	27	2	4	-2
27	-3	9	25	0	0	0
26	-4	16	23	-2	4	8
32	2	4	28	3	9	6
31	1	1	29	4	16	4
28	2	4	22	-3	9	6
30	0	0	30	5	25	0
34	4	16	31	6	36	24
	-10	80		10	120	68

$$\begin{aligned}
 r &= \frac{\sum X.Y. - (\sum X)(\sum Y)/n}{\sqrt{\sum X^2 - (\sum X)^2/n \cdot \sum Y^2 - (\sum Y)^2/n}} \\
 &= \frac{68 - (-10)(10)/10}{\sqrt{80 - (-10)^2/10 \cdot 120 - (10)^2/10}} \\
 &= \frac{78}{\sqrt{78 \times 110}} = 0.889
 \end{aligned}$$

เราวิธีคำนวณหา "r" อีกวิธีหนึ่งซึ่งวิธีนี้คัดแปลงมาจากวิธีที่โดยกตัวอย่างมาแล้วเพื่อประโยชน์ในการลดค่าของตัวเลขให้น้อยลง สะดวกแก่การคิดเลข โดยเฉพาะในกรณีตัวเลขแต่ละจำนวนมีค่ามากมายถึงเรือนร้อยเรือนพันเรือนหมื่น โดยเหตุที่เราได้ทราบซึ่งแล้วว่าการหาความสัมพันธ์ของการเปลี่ยนแปลงระหว่างสองลักษณะนั้นก็คล้ายกับที่เราหาตำแหน่งหรืออันดับความสูงต่ำของค่าตัวเลขเป็นคู่ ๆ ถ้าหากมีสหสัมพันธ์ต่าง ๆ ก็หมายความว่า ถ้าลักษณะหนึ่งเพิ่มขึ้นอีกลักษณะหนึ่งก็ลดลงเสมอ ดังนั้นแทนที่เราจะเปรียบเทียบความสูงต่ำของค่าที่แท้จริงของตัวเลขซึ่งแต่ละตัวก็มีเรือนร้อยเรือนพันจึงเป็นการเสียเวลา เมื่อเราได้ทราบความมุ่งหมายของเรื่องนี้แล้ว

หากเราจะพิจารณารัangsความสูงต่ำ (Rank) ของค่าตัวเลขแล้วนำไปวิเคราะห์ว่าเมื่อเลขจำนวนหนึ่งมีค่าสูงที่สุด เลขอีกจำนวนหนึ่งในคู่เดียวกันนั้นจะสูงกว่าหรือเปล่าเป็นต้น วิธีการนี้ เราเรียกว่า Ranks Method ดังจะยกตัวอย่างต่อไปนี้

X	Rank of X	Y	Rank of Y	Difference between Rank X & Y	$(X - Y)^2$
28	6.5	24	7	-0.5	0.25
25	10	21	10	0	0.00
29	5	27	5	0	0.00
27	8	25	6	2	4.00
26	9	23	8	1	1.00
32	2	28	4	-2	4.00
31	3	29	3	0	0.00
28	6.5	22	9	-2.5	6.25
30	4	30	2	2	4.00
34	1	31	1	0	0.00
					19.50

การให้อันดับ

ก่อนอื่นสำรวจดูในเลขชุด X จะเห็นได้ว่า ๓๔ มีค่าสูงที่สุดจึงให้เป็นอันดับ ๑ ต่อไป ๓๒ รองลงมาให้เป็นอันดับ ๒ และ ๓๑ ให้เป็นอันดับ ๓ ๓๐ เป็นอันดับ ๔ และ ๒๙ เป็นอันดับ ๕ จะเห็นว่าอันดับที่ ๒ ควรเป็น ๒๘ เพราะมีการรองจาก ๒๙ แต่ ๒๘ มีอยู่ถึง ๒ จำนวน การที่จะให้จำนวนหนึ่งเป็นอันดับที่ ๒ และอีกจำนวนหนึ่งเป็นอันดับที่ ๓ ก็ไม่ยุติธรรม เพราะทั้งสองจำนวนก็มีค่าเท่ากัน ดังนั้นอันดับที่ ๒ และที่ ๓ จึงไม่มี แต่ยอมเอามาแบ่งให้เสมอภาคกัน คือ เอา ๒ รวมกับ ๓ แล้วแบ่งครึ่ง $\frac{2+3}{2}$ ดังนั้น ๒๘ จึงอยู่ในอันดับที่ ๒.๕ ทั้งคู่ ต่อไปก็ถึงอันดับที่ ๔ คือ ๒๗ ๒๖ เป็นอันดับที่ ๕ และ ๒๕ เป็นอันดับที่ ๑๐ ตัวเลขในชุด Y ก็ให้อันดับเดียวกันเมื่อเสร็จแล้วจึงเอาอันดับของ X ลบด้วยอันดับของ Y ได้อีกช่องหนึ่งแล้วเอาผลต่างของแต่ละคู่ยกกำลังสอง $(X - Y)^2$ เสร็จแล้วรวมกันได้ ๑๙.๕ การที่เราวัดสหสัมพันธ์ด้วยผลต่างระหว่างอันดับก็เนื่องจากว่า ถ้าหากสหสัมพันธ์ทางบวกสมบูรณ์เต็มที่แล้วอันดับของ X

และ Y ในคู่เดียวกันจะเหมือนกัน ผลต่างเป็น ๐ ทุกคู่ แต่ย่อมจะเป็นไปไม่ได้ง่ายเสมอไป ยิ่งความเหลื่อมล้ำของค่าตัวเลขมีเพียงเล็กน้อยก็ยิ่งเป็นไปไม่ได้ จึงแสดงว่ามีสาเหตุรบกวนอื่น ๆ อีกมากที่ทำให้ผลต่างมีค่าแสดงออกมาไขว้เขวกันไปทางบวกบ้าง ทางลบบ้าง ดังนั้นผลต่างที่แสดงออกอย่างไม่มีระเบียบนี้ก็คือ สิ่งที่เกิดจากการกระทบกระเทือนจากสาเหตุอื่น ๆ จึงจำเป็นต้องนำมาวิเคราะห์ว่าสหสัมพันธ์นั้นจะเชื่อถือได้หรือไม่ แม้ว่าสหสัมพันธ์ทางลบก็ตามหาสมมุติฐานที่คืออันดับของ X สูงสุดและอันดับของ Y จะต้องต่ำที่สุด ในขณะที่ X ลดอันดับ Y ก็จะเพิ่มอันดับวิ่งสวนทางกัน ดังนั้นผลต่างทั้งทางบวกและทางลบจะต้องมีทั้งสูงสุดและต่ำสุด ทำให้ผลรวมของกำลังสองสูงมาก เมื่อคำนวณหา r จึงมีค่าติดลบเต็มที่ ดังนั้นจึงอาจกล่าวได้ว่าสหสัมพันธ์ทางบวก ผลรวมของผลต่างกำลังสองต้องเป็น ๐ แต่หาสมมุติฐานทางลบผลรวมของผลต่างกำลังสองจะต่ำกว่าผลคูณระหว่างจำนวนคู่กับกำลังสองของจำนวนคู่ลบด้วยหนึ่ง $n(n^2 - 1)$ ถึง ๓ เท่าเสมอ ดังนั้นสหสัมพันธ์สมมุติฐานทางลบ

$$3 \sum (X-Y)^2 = n(n^2 - 1) \text{ or } \frac{3 \sum (X-Y)^2}{n(n^2 - 1)} = 1$$

$$\text{แต่หาสมมุติฐานทางบวกค่า } \sum (X-Y)^2 = 0 \quad \text{และ } r = 1$$

ดังนั้นเมื่อสมมุติฐานทางลบจึงต้องเป็น $r = -1$

คือเปลี่ยนไป ๒ หน่วยจากสูงสุดทางบวกคือ ๑ ลดลงถึง ๐ แล้วลดลงสุดทางลบคือ $\rightarrow$ ห่างกัน ๒ หน่วยจึงต้องเอา ๒ คูณ $\frac{3 \sum (X-Y)^2}{n(n^2 - 1)}$ เป็น $\frac{6 \sum (X-Y)^2}{n(n^2 - 1)} = 2$ เมื่อสหสัมพันธ์สมมุติฐานทางลบ ทำให้เราสามารถเขียนสูตรของ r ได้ว่า $r = 1 - \frac{6 \sum (X-Y)^2}{n(n^2 - 1)}$ สำหรับใช้วิธีคำนวณแบบ Rank ดังนั้นตัวอย่างที่เราแปลค่าในสูตรนี้

$$r = 1 - \frac{6 \sum X 19.50}{10(100 - 1)} = 0.882$$

จะเห็นได้ว่าการหา r จากตัวเลขจริงได้ $r = .889$ นิดกับกับวิธี Rank เพียงเล็กน้อยเท่านั้น แต่อย่างไรก็ตามถ้าหากเรามีตัวเลขน้อยจำนวนและต้องการผลงานละเอียดจริง ๆ ก็ใช้วิธีคำนวณจากตัวเลขจริง ๆ ไม่ได้เพราะวิธี Rank นั้น บางอันดับค่าของตัวเลขก็ต่างกันมาก บางอันดับก็ต่างกันนิดหน่อย แต่เราก็ถือว่าเป็นคนละอันดับและกำหนดอันดับเรียงสม่ำเสมอ

ไปพบวิธี Rank จึงใช้โดยลำดับตัวเลขที่มาก ๆ ซึ่งแม้จะติดกันเล็กน้อยก็ไม่มีผลกระทบกระเทือนใด ๆ และความแตกต่างของค่าตัวเลขแต่ละจำนวนก็ควรไล่เรียงกันหรือสม่ำเสมอขึ้นบ้างจำนวนก็ใกล้เคียงมาก บางจำนวนก็แตกต่างกันมากมายทำให้การหาสหสัมพันธ์แบบ Rank ไขว้ไม่ได้ผล

เมื่อทราบค่าของ Coefficient "r" ได้มาแล้วจำเป็นจะต้องหาทางตัดสินใจว่าควรจะมีสหสัมพันธ์เพียงใดจึงจะถือว่าสองลักษณะที่ปรากฏนั้นมีความสัมพันธ์กันจริง ถ้าเราจะพิจารณาโดยอาศัยหลักเกณฑ์ใดสักอย่างมาไว้ในตอนต้น ๆ เกี่ยวกับความไม่สม่ำเสมอของปรากฏการณ์ตามธรรมชาติ เราก็จะพบว่าแต่ละลักษณะต่างก็ได้รับความกระทบกระเทือนจากแฟกเตอร์แวดล้อมหลายอย่าง ทำให้ค่าของตัวเลขที่ปรากฏขึ้นแต่ละครั้งแปรไปต่าง ๆ ดังนั้นจึงเห็นได้ว่าในสูตรที่ใช้คำนวณหา "r" จึงมี $S.D._y$ และ $S.D._x$ เข้ามาเกี่ยวข้องด้วย

เมื่อเป็นเช่นนี้เราจึงไม่สามารถจะเห็นสหสัมพันธ์ได้ชัดเจน จึงจำเป็นต้องหาขอบเขตมาตรฐานมาใช้เป็นเครื่องวัดปรากฏการณ์นั้นเช่นเดียวกันกับการเปรียบเทียบลักษณะเดียวโดยใช้ระดับมาตรฐาน "t" เป็นเครื่องตัดสินก็มีขอบเขต ๕ % และขอบเขต ๑ % เช่นเดียวกัน ฉะนั้นคำว่า "r" เป็นขอบเขตที่ใช้จัดการกระจายของจุดภายในกราฟ เพราะจุดแต่ละจุดเกิดขึ้นได้จากการปรากฏของสองลักษณะร่วมกันจากแต่ละคู่ ดังนั้นถ้าหากสหสัมพันธ์สมบูรณ์ จุดเหล่านี้จะฟอร์มรูปเป็นแนวตรง แต่ตามธรรมชาติมีแฟกเตอร์มากมาย สหสัมพันธ์ย่อมจะสมบูรณ์ไม่ได้ ดังนั้นจุดเหล่านั้นก็จะกระจัดกระจายออกไป ถ้าหากการกระจายเล็กน้อยยังอยู่ในขอบเขต ๑ % ซึ่ง เป็นขอบเขตแถบก็ถือว่ามีความสัมพันธ์กันอย่างเด่นชัด Highly Significant ถ้าการกระจายมากน้อยก็ยังอยู่ในขอบเขต ๕ % ซึ่งกว้างกว่า ๑ % ก็ถือว่ามีความสัมพันธ์พอเป็นที่เชื่อถือ Significant ถ้าการกระจายมากจนเลยขอบเขต ๕ % ออกไปเราก็ถือว่าการเปลี่ยนแปลงของสองลักษณะร่วมกันนั้นไม่สามารถจะถือว่าเป็นเส้นตรง หรือกล่าวอีกนัยหนึ่งว่าไปก็ว่ามีสหสัมพันธ์กันได้ การหาขอบเขตมาตรฐาน "r" มี r table อยู่ ดูจาก $D.F. = n' - 2$ ($n' =$ จำนวนคู่ที่เราหาสหสัมพันธ์กัน) ตามตัวอย่างมีเลขอยู่ ๑๐ คู่ ฉะนั้น $D.F. = 10 - 2$ เปิด Table r $D.F. = 8$, 5 % = 0.632, 1% = 0.765 ที่คำนวณได้

จริง = .๘๘๘ สูงกว่าขอบเขต ๑ % แสดงว่าลักษณะ X กับ Y มีสหสัมพันธ์กันอย่างเด่นชัด
ถือเป็นนัยสำคัญในทางสถิติ ถ้าค่า r ที่คำนวณได้ต่ำกว่า ๐.๗๖๕ แต่ยังสูงกว่า .๖๗๒ ก็ถือว่า
มีความสัมพันธ์พอเชื่อถือได้ ถ้าต่ำกว่า ๐.๖๗๒ ก็ถือว่าไม่มีสหสัมพันธ์กันเลย

ในกรณีที่เราไม่มี Table r แต่มี Table t อยู่ เราอาจคำนวณหาค่าขอบเขต
มาตรฐาน r ๕ % และ ๑ % ได้โดยวิธีต่อไปนี้

$$r^2 = \frac{t^2 (1 - r^2)}{n' - 2}$$

ค่าของ t ได้จาก Table ที่ D.F. $n' - 2$ เราจะเห็นได้ว่ามี ๒ ค่า คือค่า ๕ % และ
๑ % เช่นเดียวกัน ดังนั้นจึงต้องแทนค่า ๒ ครั้ง

การคำนวณหา "r" ขอบเขต ๕ % เราก็แทนค่าด้วยค่า ๕ % ของ t และการ
คำนวณหา r ขอบเขต ๑ % เราก็แทนค่าด้วย ๑ % ของ t

$n' =$ จำนวนคู่ ตามตัวอย่างคือ ๑๐ คู่

เมื่อดู t Table ที่ D.F. $10 - 2 = 8$ ได้ $t \ 5\% = 2.306$

$$r^2 = \frac{2.306^2 (1 - r^2)}{10 - 2}$$

$$8r^2 = 5.313 - 5.313 r^2$$

$$8r^2 + 5.313 r^2 = 5.313$$

$$r^2 = \frac{5.313}{8 + 5.313}$$

$$= .399084$$

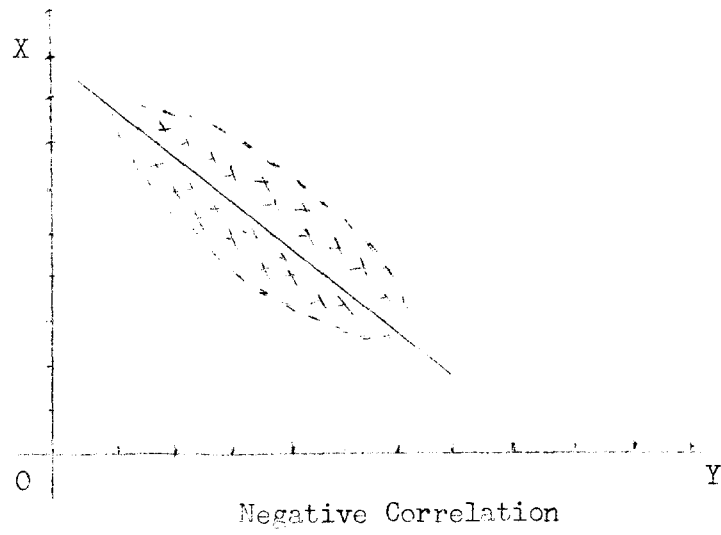
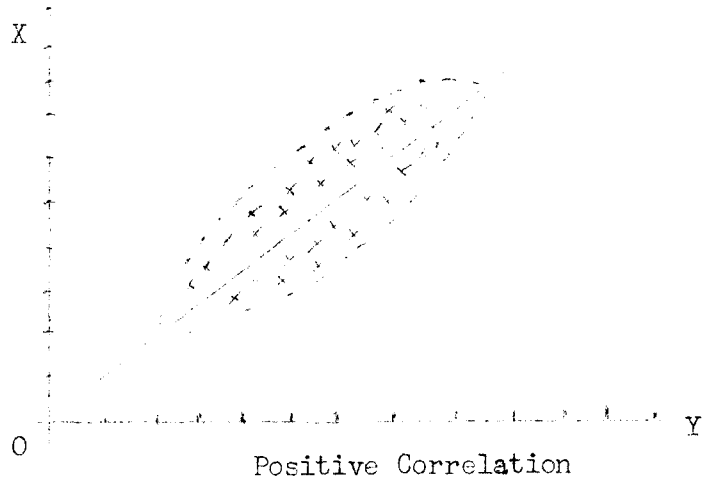
$$r = \sqrt{.399084}$$

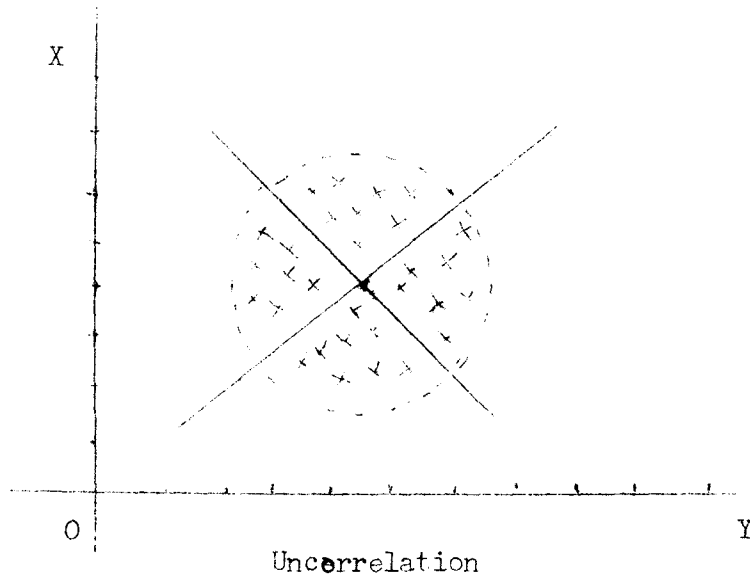
$$= .632 \dots\dots\dots 5\%$$

$$t \ 1\% = 3.355$$

$$r^2 = \frac{3.355^2 (1 - r^2)}{10 - 2}$$

$$r = .765 \dots\dots\dots 1\%$$





REGRESSION

การศึกษาสองลักษณะควบกันนั้น นอกจากจะทราบว่าการเปลี่ยนแปลงของลักษณะหนึ่ง จะสามารถทำให้อีกลักษณะหนึ่งเปลี่ยนไปด้วยหรือไม่ หรือที่เรียกกันตามภาษาสถิติว่ามีสหสัมพันธ์ กันหรือไม่ แล้ว ถ้าหากปรากฏว่าสหสัมพันธ์กันจริง เรายังสามารถทำให้เกิดประโยชน์มากยิ่งขึ้น ไปอีกโดยการหาต่อไปอีกว่า ถ้าลักษณะ X เปลี่ยนไปถึงแค่นั้นแค่นี้แล้ว ลักษณะ Y จะเปลี่ยน ไปถึงไหนเท่าใด การเปลี่ยนแปลงซึ่งสัมพันธ์กันเป็นกฎเกณฑ์เองที่เป็นเรื่องของ Regression ดังจะขอยกตัวอย่างเช่น ความสัมพันธ์ระหว่างจำนวนต้น (Y) กับผลผลิต (X) ซึ่งได้ยก ตัวอย่างไว้ในเรื่องสหสัมพันธ์ (ตารางที่ ๑) จะเห็นได้ว่า เราหาผลเฉลี่ยของ X ได้ ๒๕ ก.ก. และผลเฉลี่ยของ Y ได้ ๒๒ ต้น เราอาจกล่าวได้ว่า เฉลี่ยผลผลิต ๒๒ ต้นต่อ ๒๕ ก.ก. เมื่อการหาสหสัมพันธ์ปรากฏว่ามีสหสัมพันธ์ในทางบวกอย่างแน่นอน คือจำนวนต้น เพิ่มขึ้น ผลผลิตก็เพิ่มขึ้นด้วย เราก็เอาผลเฉลยมากล่าวต่อไปอีกได้ คือเอาผลเฉลี่ยผลผลิตหารด้วยผลเฉลี่ยของจำนวนต้นหรือ $22 \div 25 = 0.88$ แล้วกล่าวด้วยผลเฉลี่ยว่า ถ้าจำนวนต้น

เพิ่มขึ้นทีละหนึ่งตันผลผลิตจะเพิ่มขึ้นทีละ ๑.๑๑๕ ก.ก. หรือถ้าจำนวนต้นลดลง ๑ ต้น ผลผลิตก็จะลดลง ๑.๑๑๕ ก.ก. การที่เรานำเอาผลเฉลยมาตั้งเป็นกฎเกณฑ์ขึ้นมานั้น ถ้าหากเราใช้จุดเฉลยของสองลักษณะรวมกันคือจุดตัดระหว่าง ๒๖ กับ ๒๕ เป็นจุดเริ่มต้น แล้วเขียนเส้นกราฟโดยการเพิ่มจำนวนต้นทีละ ๑ ต้น เพิ่มผลผลิตทีละ ๑.๑๑๕ ก.ก. หรือลดจำนวนต้นทีละ ๑ ต้นลดผลผลิตทีละ ๑.๑๑๕ ก.ก.ก็ตาม เราจะได้กราฟที่เป็นเส้นตรงซึ่งเรียกว่า Linear Regression ดังนั้นจึงเปรียบคล้าย ๆ กับเส้นทฤษฎีหรือเส้นเฉลยของแนวสหสัมพันธ์ เพื่อใช้แสดงทิศทางการกระจายของคู่ลักษณะนั้น การเพิ่มผลผลิตต่อหนึ่งต้นที่แสดงไว้นั้นเป็นเพียงตัวอย่างง่าย ๆ เพื่อใช้อธิบายความหมายของ Regression ให้เข้าใจแจ่มแจ้งขึ้นเท่านั้น ส่วนวิธีการที่แท้จริงนั้นเราจะต้องไม่ลืมว่าลักษณะที่แสดงออกของพืชแต่ละต้น สัตว์แต่ละตัว ย่อมมีความคลาดเคลื่อนประจำ ดังนั้นความสัมพันธ์รวมกันจึงต้องมีความคลาดเคลื่อนรวมกันด้วย จึงจำเป็นต้องคำนึงถึงสิ่งเหล่านี้แล้วนำมาใช้คำนวณทั้งหลักเกณฑ์หา Regression ต่อไป

การหา Regression Coefficient (b_{xy}) คือการหาปฏิภากระหว่างสิ่งที่เกิดจาก X และ Y รวมกันกับสิ่งที่เกิดจาก Y โดยเหตุที่เราใช้คือ Y หรือจำนวนต้นเป็นหลัก

$$\text{Regression Coefficient } b_{xy} = \frac{S.P._{xy}}{S.S._y} = \frac{78}{110} = .709$$

การที่เราใช้คือ Y เป็นหลักก็เพื่อจะต้องการทำนายว่า ถ้าจำนวนต้นเท่านั้นผลผลิตจะต้องเป็นเท่าไร และผลจากการทำนายนั้นเองจะเกิดเป็นเส้น Regression ขึ้น ซึ่งคล้าย ๆ กับเส้นเฉลยที่ใช้แทนลักษณะการกระจายนั้น

ถ้าเรากำหนด Y หรือจำนวนต้นขึ้นแล้วต้องการทำนายผลผลิตตามทฤษฎี

เราก็สามารถใช้สูตรต่อไปนี้ ให้ X = ตัวเลขที่ต้องการทำนาย

Y = ตัวเลขที่กำหนดขึ้นเป็นหลัก

M_x = Mean ของ X

M_y = Mean ของ Y

$$X = M_x + b_{xy} (y - M_y)$$

$$= 29 + .709(y - 26)$$

$$= 29 + .709y - 18.434$$

$$= 10.566 + .709y$$

การทราบว่าค่า Y เป็นเท่าใดแล้ว X จะเป็นเท่าไร

สมมติว่าค่าของการทราบว่า Y = 10 ต้น ผลผลิต X จะเป็นกี่ ก.ก.

$$X = 10.566 + (.709 \times 10) = 17.656$$

ถ้าเราจะทดลองทำนายจากผลที่แท้จริง ๆ ในตารางที่ เราจะได้อะไร

$$Y = 24 \quad X = 10.566 + (.709 \times 24) = 27.582$$

$$Y = 21 \quad X = 10.566 + (.709 \times 21) = 25.455$$

$$Y = 27 \quad X = 10.566 + (.709 \times 27) = 29.709$$

$$Y = 25 \quad X = 10.566 + (.709 \times 25) = 28.291$$

$$Y = 23 \quad X = 10.566 + (.709 \times 23) = 26.873$$

$$Y = 28 \quad X = 10.566 + (.709 \times 28) = 30.418$$

$$Y = 29 \quad X = 10.566 + (.709 \times 29) = 31.127$$

$$Y = 22 \quad X = 10.566 + (.709 \times 22) = 26.164$$

$$Y = 30 \quad X = 10.566 + (.709 \times 30) = 31.836$$

$$Y = 31 \quad X = 10.566 + (.709 \times 31) = 32.545$$

จะเห็นได้ว่าค่าจริง ๆ ของ X ใกล้กับผลทำนายมาก ทั้งนี้เนื่องจากการหาสหสัมพันธ์เราได้ค่าของ r ถึง 0.๘๘๘ ซึ่งแสดงว่ามีความสัมพันธ์กันมาก ถ้าหากสหสัมพันธ์ค่าที่แท้จริงก็ยิ่งกระจัดกระจายห่างไปจากผลการทำนายหรือเส้นทฤษฎีมากขึ้น ระยะห่างที่ X ซึ่งปรากฏจริง ๆ (observed "X") ผิดไปจาก X ทำนาย (Predicted "X") นี้เราเรียกว่า "deviation from Linear Regression"

ต่อไปเราก็เอาค่าของ X ที่คำนวณได้กับค่าของ X ที่ observed ได้จากการค้นคว้าทดลองจริงๆ มาคำนวณหามาตรฐานความผิดเพี้ยน โดยการเอาของทุกคู่ที่ได้อ้างอิง มาเทียบกัน เพื่อหา deviation มาแสดงว่าสิ่งที่ได้จริง ๆ นั้นคลาดเคลื่อนไปจากแนวทฤษฎีหรือเคลื่อนไปจากเส้นตรงมากน้อยเพียงไร

Observed X	Calculated X'	$X - X'$	$(X - X')^2$
28.0	27.582	0.418	0.174724
25.0	25.455	-0.455	0.207025
29.0	29.709	-0.709	0.502681
27.0	28.291	-1.291	1.666681
26.0	26.873	-0.873	0.762129
32.0	30.418	1.582	2.502724
31.0	31.127	-0.127	0.016129
28.0	26.164	1.836	3.370896
30.0	31.836	-1.836	3.370896
34.0	32.545	1.455	2.111025
$\Sigma X = 290.00$	$\Sigma X' = 290.000$	0.000	14.690910

หมายเหตุ ถ้าหากการปกคลุมใกล้เคียงกันดี ผลรวมของ Observed X และ Calculated X' นั้นทำหน้าที่คล้ายผลเฉลี่ยหรือแนวเฉลี่ยซึ่งอยู่ระดับกลางระหว่างของจริงซึ่งเพียงทั้งทางบวกและทางลบ แต่บางที่ปกคลุมไม่เหมาะก็อาจผิดเพี้ยนกันไปบ้างเล็กน้อย

ต่อไปเราก็หามาตรฐานความคลาดเคลื่อนของ Regression Coefficient

$$\begin{aligned} \text{คือ } S.E._{b_{xy}} &= \pm \sqrt{\frac{(X - X')^2}{(n' - 2) S.S._y}} \quad (n' = \text{จำนวนคู่}) \\ &= \pm \sqrt{\frac{14.690910}{8 \times 110}} = 0.129 \end{aligned}$$

$$\begin{aligned} \text{Calculated "t" value} &= \frac{\text{Regression Coefficient}}{S.E. \text{ of Regression Coefficient}} = \frac{b_{xy}}{S.E._{b_{xy}}} \\ &= \frac{0.709}{0.129} = 5.496 \end{aligned}$$

เปิด "t" Table D.F. คูณที่ $n' - 2 = 10 - 2 = 8$ ได้ค่า $t_{5\%} = 2.306$
 และ "t" $1\% = 3.355$, Calculated "t" ที่คำนวณได้คือ ๕.๔๘ สูงกว่าระดับ
 ๑ % มาก ดังนั้นจึงเป็น Highly Significant แสดงว่าลักษณะ X กับ X' ไปด้วยกัน
 ได้ X ที่ปรากฏจึงมีความสัมพันธ์กับ Y จริงตลอดรายการของ Experiment คือ
 Calculated X' ได้ก็มีความคลาดเคลื่อนน้อย ดังนั้นความสัมพันธ์ระหว่าง X กับ Y
 ก็ถือเป็นสิ่งแน่นอนจริงจึงได้

ถ้าหากเราจะหาความแตกต่างของผลผลิต X เองโดยไม่คิดถึง Y เราก็กระทำได้
 โดยการหา S.E. of Estimate (S.E. Est.) ดังนั้น S.E. Est. จึงเป็นเครื่องแสดง
 ความแตกต่างที่เกิดจากตัวของ X เอง ทั้งนี้หาได้โดยคำนวณจากความคลาดเคลื่อนที่ Observed X
 เคลื่อนไปจาก Calculated X' นั้นเอง

$$\begin{aligned} \text{S.E. Est.} &= \pm \sqrt{\frac{(X - X')^2}{n' - 2}} \\ \text{S.E. Est.} &= \pm \sqrt{\frac{14.690910}{8}} \\ &= \pm 1.355 \end{aligned}$$

เมื่อเราพลิกกลับไปดูการหา Correlation ในตอนต้นเราจะพบว่าเราได้โดย
 หา S.D._x ไว้ครั้งหนึ่งแล้ว = ๒.๗๘๘ เป็นความคลาดเคลื่อน ของ X ทั้งหมด ยิ่งเราได้
 ทราบว่า X กับ Y สัมพันธ์กันอย่างแท้จริง ดังนั้น S.D._x จึงรวมเอาความคลาดเคลื่อนที่
 เกิดจาก X เอง และความคลาดเคลื่อน ของ X ที่เกิดจากอำนาจของ Y เข้าไว้ด้วย
 ถ้าเราเขียนแผนผังก็จะได้นั่นคือ

ความคลาดเคลื่อนที่ X แสดงออกทั้งหมด = สิ่งที่เกิดจาก X เอง + สิ่งที่เกิดจาก Y

หรือเขียนเป็นสูตร $\text{S.D.}_x = \text{S.E. Est.} + \text{สิ่งที่เกิดจาก Y}$

แทนค่า..... 2.789 = ๑.๓๕๕ + สิ่งที่เกิดจาก Y

หรือสิ่งที่เกิดจาก Y = ๒.๗๘๘ - ๑.๓๕๕ = ๑.๔๓๓

ส่วน S.E. Est. นั้นเป็นความแตกต่างที่เกิดจาก X เอง ซึ่งนอกเหนือไปจากอำนาจของ Y อาจเป็นลักษณะและแฟกเตอร์อื่น ๆ รวมทั้งโชค (Chance) ด้วย

ในการวางแผนคนความแบบ Randomized Block Design สมมุติว่าเราต้องการเปรียบเทียบการใส่ปุ๋ยผสม Combination หนึ่ง โดยการตั้งอัตราของจำนวนปุ๋ยที่ใช้ต่าง ๆ กัน เพื่อบ่งจากน้อยไปหามาก สมมุติว่า

A_0	=	Control คือไม่ใส่ปุ๋ยเลย
A_1	=	ปุ๋ยผสมจำนวน ๑๐๐ ก.ก.ต่อไร่
A_2	=	ปุ๋ยผสมจำนวน ๒๐๐ ก.ก.ต่อไร่
A_3	=	ปุ๋ยผสมจำนวน ๓๐๐ ก.ก.ต่อไร่
A_4	=	ปุ๋ยผสมจำนวน ๔๐๐ ก.ก.ต่อไร่

เรามีความมุ่งหมายในการค้นคว้าทั้งอาจแบ่งแยกออกได้เป็น สาม ข้อดังนี้

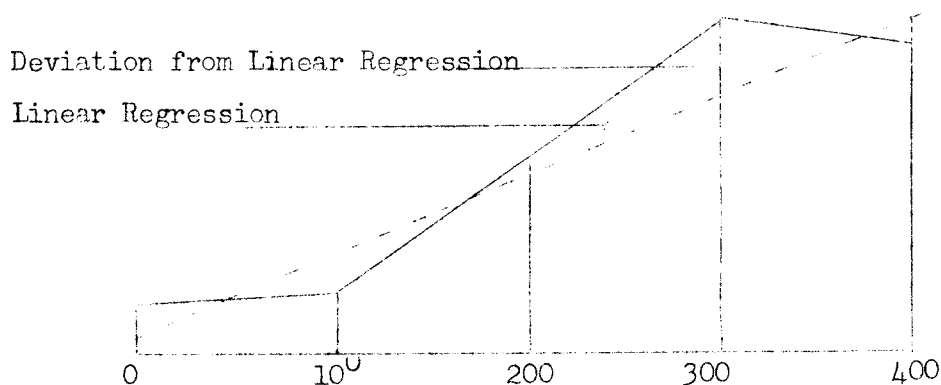
(๑) ต้องการจะทราบว่า การใส่ปุ๋ยกับไม่ใส่ปุ๋ยนั้นอย่างไหนจะให้ผลดีกว่ากัน ทั้งนี้เนื่องจากบางทีในดินหรือในเครื่องปลูก (Medium) อาจมีปุ๋ยหรืออาหารอยู่เพียงพอแล้วก็ได้ มิใช่ว่าการใส่ปุ๋ยจะให้ผลดีเสมอไป

(๒) ถ้าปรากฏผลว่าการใส่ปุ๋ยให้ผลดีกว่าไม่ใส่ปุ๋ย เช่นทำให้เจริญเติบโตเร็ว ช่วยเพิ่มปริมาณผลผลิต หรือช่วยปรับปรุงคุณลักษณะต่าง ๆ เราก็จำเป็นต้องทราบต่อไปอีกว่า ถ้าเพิ่มจำนวนปุ๋ยให้มากขึ้นเป็นลำดับ ผลผลิตจะเพิ่มขึ้นเรื่อย ๆ หรือไม่ การวางจำนวนปุ๋ยหรือกำหนดอัตราปุ๋ยที่ใช้ให้เพิ่มขึ้นเป็นลำดับนี้ เราจะต้องมีความรู้เกี่ยวกับเรื่องดิน ตลอดจนความต้องการปุ๋ยของพืชที่เราทำการทดลองเป็นอย่างดี มิฉะนั้นอาจทำให้เราวางระยะระหว่างอัตราปุ๋ยของแต่ละ Treatment ไม่ได้เหมาะสม กล่าวคือถ้าวางระยะแคบเกินไป ผลงานที่ได้ก็จะไม่เด่นชัด ยิ่งกว่านั้นก็อาจต้องเสียเวลาปฏิบัติซ้ำ ๆ กันต่อไปอีกหลายปี กว่าจะถึงจุดที่เหมาะสมซึ่งเป็นจุดที่ให้ผลดีที่สุด แต่ถ้าวางระยะห่างเกินไปเราก็จะไม่สามารถจะทราบจุดเหมาะสมที่แน่นอน เพราะอาจครอบคลุมไว้อย่างกว้าง ๆ ดังนั้นการจัดระยะจึงเปรียบเสมือนเครื่องมือวัดความยาวของวัตถุ ถ้าใช้เครื่องมือเล็ก ๆ แต่ละเอียงเพื่อวัดของที่มีขนาดยาว ซึ่งเปรียบคล้ายกับการวางระยะแคบๆ

จนเกินไปก็จะเสียเวลาโดยไม่จำเป็น แต่ถ้าใช้เครื่องมือหยาบเกินไปสำหรับวัดของเด็ก หากความแตกต่างมีน้อยเราก็ไม่สามารถจะเห็นได้ ดังนั้นจึงควรได้ศึกษาลักษณะสิ่งที่เราจะทำการวัดให้ละเอียดเสียก่อน ในเรื่องของปุ๋ยก็ควรจะได้อบรมคุณสมบัติของดิน ความต้องการปุ๋ยของพืชที่เราเกี่ยวข้องด้วย และสภาพสิ่งแวดล้อมอื่น ๆ บาง

(๓) สมมุติว่าเราได้อบรมแล้วว่าการใส่ปุ๋ยเพิ่มขึ้น จะทำให้ผลผลิตเพิ่มขึ้นด้วย ซึ่งก็เป็นเรื่องสหสัมพันธ์ระหว่างจำนวนปุ๋ยกับผลผลิต อันให้กำเนิดแนวของ Regression Line หลักเกณฑ์ของเส้นตรง (Linear Regression) เสนอไป ถ้าใส่ปุ๋ยเพิ่มขึ้นผลผลิตก็จะเพิ่มขึ้น เราถือว่าเป็น Linear Regression คือเป็นเส้นตรงขึ้นไป แต่ถ้าวางปุ๋ยเพิ่มมากขึ้นจนกระทั่งแรงเกินไปสามารถทำอันตรายหรือกระทบกระเทือนทางเจริญเติบโต ต้นไม้ก็ให้ผลต่ำลง ถ้าปุ๋ยมากเกินไปมาก ๆ ก็อาจเป็นอันตรายถึงตายได้ นี่ก็แสดงให้เห็นแล้วว่าเส้น Linear Regression จะต้องมีความตกลงมาบ้างจะสูงขึ้นตลอดไป ถ้าหากเราวางอัตราปุ๋ยชิดกันมากเกินไป เส้นก็จะค้อม ๆ ตกเป็นรูปโค้งเห็นได้ชัด แต่ถ้าอัตราปุ๋ยห่างมากเกินไป เราจะรู้การตกของเส้นได้ก็ต่อเมื่อต้นไม้ตายเสียก็ได้ ถ้าเราวางระยะพอเหมาะ ก็จะสังเกตเห็นการตกของผลผลิตได้ชัดเจนพอสมควร ความจริงแล้วเราไม่จำเป็นต้องรอให้ผลผลิตตกต่ำเสียก่อน เพียงแต่เห็นว่าเส้นนั้นเริ่มจะโค้ง แสดงว่าเมื่อเพิ่มปุ๋ยขึ้นไปแต่ ผลผลิตที่เคยเพิ่มขึ้นมากกลับเพิ่มเพียงเล็กน้อย แสดงว่าไม่คุ้มกับการลงทุน เราก็อาจสรุปได้ว่าไม่ควรเพิ่มปุ๋ยขึ้นไปอีก เพราะถ้าเพิ่มปุ๋ยขึ้นไปอีก ผลผลิตจะตกลงแทน

เราทราบอยู่แล้วว่า มีสาเหตุทางธรรมชาติหลายประการที่คอยรบกวนปรากฏการณ์ให้เกิดความคลาดเคลื่อนอยู่เสมอ ดังนั้นการตั้งอัตราปุ๋ยเพิ่มขึ้นเป็นระยะ ๆ เท่า ๆ กัน ถ้าผลผลิตเพิ่มขึ้นก็จะไม่เพิ่มขึ้นระยะละเท่า ๆ กันพอดีเสมอไป อาจเพิ่มขึ้นมากไปบ้างน้อยไปบ้างเล็กน้อย ๆ แต่เราก็ถือว่ามีความเป็นเส้นตรงสูงขึ้นไป ส่วนความเพี้ยนที่ทำให้เส้นตกไปบ้างเล็กน้อย ๆ นั้นเราก็คือว่าเป็นสาเหตุอื่น ๆ ทางธรรมชาติ ไม่ใช่เรื่องปุ๋ย จึงไม่สามารถทำให้เส้นนั้นโค้งออกนอกแนวขอบเขตของเส้นตรงไปได้



ถ้าเราจะกล่าวเป็นทางสถิติแล้ว ก็อาจกล่าวได้ว่า ระยะที่เพี้ยนไปจากเส้นตรงหรือ Deviation from Linearity ไม่ถือเป็นนัยสำคัญในทางสถิติ (unsignificant) แต่ถ้าในการทดลองค้นคว้าได้ Deviation from Linear Regression แสดงความแตกต่างซึ่งถือเป็นนัยสำคัญทางสถิติ หรือกล่าวได้ว่าเส้นนั้นเริ่มจะหนีจากความเป็นเส้นตรง เช่นจำนวนปุ๋ยเพิ่มขึ้นถึงจุดสูงสุด ถ้าใส่ปุ๋ยเพิ่มขึ้นอีกผลผลิตก็จะลดลงแน่นอน เมื่อการหนีจากเส้นตรงมี Significant ก็แสดงว่ามีสาเหตุธรรมชาติของธรรมชาติ แต่เนื่องมาจาก Treatment ของเรา

ต่อไปนี้เป็นตัวเลขผลผลิตของพืชที่ได้จากการค้นคว้าเปรียบเทียบจำนวนปุ๋ยตั้งแต่ A_0 ถึง A_4 ดังใ้กล่าวมาแล้วจากสองหน้าที่ผ่านมา

(Coded by - 50)

Block	Treatment										Block Total
	A_0		A_1		A_2		A_3		A_4		
	X	X^2	X	X^2	X	X^2	X	X^2	X	X^2	
I	48		56		61		62		75		302
	-2	4	6	36	11	121	12	144	25	625	52
II	53		62		65		68		66		314
	3	9	12	144	15	225	18	374	16	256	64
III	50		56		66		77		71		320
	0	0	6	36	16	256	27	729	21	441	70
IV	42		58		60		67		72		299
	-8	64	8	64	10	100	17	289	22	484	49
V	42		59		61		65		67		294
	-8	64	9	81	11	121	15	225	17	289	44
Treatment Total	235		291		313		339		351		1529
	-15		41		63		89		101		279

$$\text{Grand Total} = 1529$$

$$\text{Grand Total (Coded value)} = 279$$

เนื่องจากเราพิจารณาทั้งในด้านผลผลิตและจำนวนปุ๋ยที่เพิ่มขึ้น จึงมีส่วนที่ผิดแปลกไปจากการวิเคราะห์เปรียบเทียบผลผลิตแต่ด้านน้อย ในเรื่องของ Randomized Block ขรรคชาติได้เรียนมาแล้วเล็กน้อย เมื่อเกิดมีการพิจารณาสองด้านขึ้น จึงจำเป็นต้องมีการสมมุติทั้ง X และ Y คือให้ผลผลิตเป็น X ดังนั้นตัวเลขที่แสดงผลผลิตจึงเป็นเรื่องของ X และตัวเลขที่แสดงจำนวนปุ๋ยจึงเป็นเรื่องของ Y แต่ตัวเลขที่แสดงจำนวนปุ๋ยคือ ๐ - ๑๐๐ - ๒๐๐ - ๓๐๐ - ๔๐๐ ปีระยะเท่ากันหมด เราก็อาจใช้หน่วยแทนโดยคิดเป็นหน่วยละ ๑๐๐ ดังนั้นจึงมีค่าเป็น ๐ - ๑ - ๒ - ๓ - ๔ ได้ ขั้นแรกการคำนวณ X ก็กระทำตามแบบ Randomized Block ขรรคชาติได้เรียนมาแล้ว คือ

$$C.F.X = \frac{(\text{Grand total } X)^2}{n} = \frac{279^2}{25} = 3113.64$$

$$\begin{aligned} S.S. \text{ Total} &= X_1^2 + X_2^2 + X_3^2 + X_4^2 + X_5^2 + X_6^2 + \dots + X_{25}^2 - C.F.X \\ &= 4 + 9 + 0 + \dots + 289 - 3113.64 \\ &= 2017.36 \end{aligned}$$

$$\begin{aligned} S.S. \text{ Block} &= \frac{B_1^2 + B_2^2 + B_3^2 + B_4^2 + B_5^2}{t} - C.F.X \\ &= \frac{52^2 + 64^2 + 70^2 + 49^2 + 44^2}{5} - 3113.64 \\ &= 93.76 \end{aligned}$$

$$\begin{aligned} S.S. \text{ Treatments} &= \frac{T_1^2 + T_2^2 + T_3^2 + T_4^2 + T_5^2}{b} - C.F.X \\ &= \frac{-15^2 + 41^2 + 63^2 + 89^2 + 101^2}{5} - 3113.64 \\ &= 1685.76 \end{aligned}$$

$$\begin{aligned}
 \text{S.S. Error} &= \text{S.S. Total} - (\text{S.S. Blocks} + \text{S.S. Treatments}) \\
 &= 2017.36 - (93.76 + 1685.76) \\
 &= 237.84
 \end{aligned}$$

แล้วจึงหา Sum of Square ของ Y เพื่อจะหาความแตกต่างของ Y ในแต่ละ Treatment

$$\text{S.S. Treatment for Y} = \frac{Y_1^2 + Y_2^2 + Y_3^2 + Y_4^2 + Y_5^2}{b} - \text{C.F.}_y$$

$$\text{C.F.}_y = \frac{(\text{Grand total Y})^2}{n} = \frac{(0 + 1 + 2 + 3 + 4)^2}{25} = 4.0$$

$$\begin{aligned}
 \text{S.S. Treatment for Y} &= \frac{0^2 + 1^2 + 2^2 + 3^2 + 4^2}{5} - 4 \\
 &= 6 - 4 = 2
 \end{aligned}$$

แต่เนื่องจากเราต้องการความสัมพันธ์ระหว่าง X กับ Y ด้วย จึงจำเป็นต้องหา S.P._{xy} เราจะเห็นได้ว่าในเรื่องของ Treatments เวลาเราหาของ X เราก็เอา X คูณตัวเอง เวลาหาของ Y ก็เอา Y คูณตัวเอง (คือยกกำลังสอง) แต่เวลาเราหา X กับ Y รวมกันจึงต้องเอา X คูณกับ Y ซึ่งเราเรียกว่า Sum of Product XY หรือ S.P._{xy} ซึ่งแท้จริงมีความหมายคล้ายคลึงกันกับ S.S. แต่เป็นของสองลักษณะรวมกัน

$$\text{S.P.}_{xy} = \frac{\left(\frac{T}{y} \times 1\right) + \left(\frac{T}{y} \times 2\right) + \left(\frac{T}{y} \times 3\right) + \left(\frac{T}{y} \times 4\right) + \left(\frac{T}{y} \times 5\right)}{b} - \text{C.F.}_{xy}$$

C.F._{xy} ก็เช่นเดียวกัน เราหา C.F._x เราก็เอา X คูณตัวเอง และหา C.F._y ก็เอา Y คูณตัวเอง

$$\text{ดังนั้น } \text{C.F.}_{xy} = \frac{(\text{Grand total X}) (\text{Grand total Y})}{n}$$

$$= \frac{279 \times 10}{25}$$

$$= 111.6$$

$$\text{S.P.}_{xy} = \frac{(0 \times -15) + (1 \times 4) + (2 \times 63) + (3 \times 89) + (4 \times 101)}{5} - 111.6$$

$$= 56$$

ต่อไปเราก้หา S.S. ของ Linear Regression เพื่อต้องการจะทราบว่า การเพิ่มจำนวนปุ๋ยขึ้นไปนั้นจะยังคงทำให้ผลผลิตเพิ่มขึ้นเป็นเส้นตรงขึ้นไปอีกหรือไม่

$$\begin{aligned} \text{S.S. for Linear Regression} &= \frac{(S.P. \cdot xy)^2}{S.S. \cdot y} \\ &= \frac{56^2}{2} = 1568.00 \end{aligned}$$

ส่วน D.F. ของ Linear Regression ต้องเป็นหนึ่งเสมอ ทั้งนี้เพราะเหตุว่าจำนวนมีอยู่ ๒ เสมอ คือ ๑ เป็นเส้นตรง และอีก ๑ ก็คือ ถ้าไม่เป็นเส้นตรงก็ต้องโค้งแน่ ๆ ดังนั้น D.F. จึง = $n - 1$ หรือ $2 - 1 = 1$ ส่วน S.S. ก็เหมือนกัน ใน S.S. ของ Treatments นั้นเป็นเครื่องแสดงความแตกต่างของผลผลิตของแต่ละ Treatments ดังนั้นความแตกต่างของผลผลิตระหว่างแต่ละ Treatments ถ้ามีการใส่ปุ๋ยเพิ่มขึ้น ผลผลิตอาจเพิ่มขึ้นเป็นเส้นตรงหรือโค้งตกลงก็ได้ ดังนั้น $S.S. \text{ Treatments} = S.S. \text{ Linear Regression} + S.S.$

$$\begin{aligned} \text{Deviation from Linear Regression} \quad S.S. \text{ Deviation from Linear Regression} &= S.S. \text{ Treatments} - S.S. \text{ Linear Regression} = 1685.76 - 1568.00 = 117.76 \end{aligned}$$

การหา D.F. ของ Deviation ก็เช่นเดียวกัน

$$\begin{aligned} &= D.F. \text{ Treatments} - D.F. \text{ Linear Regression} \\ &= 4 - 1 = D.F. \quad 3 \end{aligned}$$

ต่อไปเราก้หาว่าการใส่ปุ๋ยกับไม่ใส่ปุ๋ยจะให้ผลดีต่างกันหรือไม่ โดยการหา S.S. ของ Fertilizer &

$$\begin{aligned} \text{No fertilizer} &= \frac{(\text{No fertilizer})^2}{b} + \frac{(\text{Fertilizer})^2}{n - b} - C.F. \\ &= \frac{(-15)^2}{5} + \frac{(41 + 63 + 89 + 101)^2}{25 - 5} - \frac{279^2}{25} \\ &= \frac{225}{5} + \frac{294^2}{20} - 3113.64 \\ &= (45.0 + 4321.8) - 3113.64 \\ &= 1253.16 \end{aligned}$$

D.F. = 1 (เพราะมีการเทียบ ๒ อย่างคือปุ๋ยกับไม่ใส่ปุ๋ย) ต่อไปก็หาว่าจำนวนปุ๋ยเท่าไรจึงจะให้ผลดี คือ

หา S.S. Amount of Fertilizer เราตัดเอาพวกที่ไม่ได้ใส่ปุ๋ยออกไป เพราะ
ต้องการเปรียบเทียบเฉพาะที่ใส่ปุ๋ยเท่านั้น C.F. ก็ใช้คำนวณจากค่าที่ใส่ปุ๋ยทั้งหมด ไม่รวม
พวกที่ไม่ได้ใส่ปุ๋ย

$$\begin{aligned} \text{S.S. Amount of Fertilizer} &= \frac{T_2^2 + T_3^2 + T_4^2 + T_5^2}{b} - \text{C.F. fertilizer} \\ &= \frac{41^2 + 63^2 + 89^2 + 101^2}{5} - \frac{(41+63+89+101)^2}{20} \\ &= 432.60 \text{ with } 3 \text{ D.F.} \end{aligned}$$

Analysis of Variance

Sources of Difference	D.F.	S.S.	Variance (S.S./D.F.)	F ratio	
				Calculated	Table
Total	24	2017.36			
Blocks	4	93.76	23.44		
Treatments	4	1685.76	421.44	28.36 **	
Linear Reg.	1	1568.00	1568.00	105.51 **	
Deviation	3	117.76	39.25	2.64	
Fertilizer & No fertilizer	1	1253.16	1253.16	84.33 **	
Amount of fertilizer	3	432.60	144.20	9.70 **	
ERROR	16	273.84	14.86		

จากการวิเคราะห์จะเห็นว่า Treatments สามารถแตกต่างออกเป็นรายการย่อยได้ดังนี้ $\text{S.S. Treatments} = \text{S.S. Linear} + \text{S.S. Deviation}$

ความหมายว่าความแตกต่างระหว่าง Treatment อาจเป็นไปได้สองลักษณะ คือ อาจทำให้เกิดเป็นเส้นขึ้นไป ถ้าวลยผลิตเพิ่มขึ้นเรื่อย ๆ หรือสหสัมพันธ์ทางบวกแต่ก็อาจเป็น

เส้นตรงเรื่อยลงมา ถ้าสหสัมพันธ์เป็นไปทางลบ แต่ต้องเป็นไปเพียงทางเดียวคือทางหนึ่งทางใด จึงจะเป็นเส้นตรงได้ ถ้าไม่เป็นเส้นตรงก็อาจเป็นเส้นโค้ง (คือมี Deviation จากเส้นตรง) การจะแสดงให้ถือเป็นเส้นตรงหรือเส้นโค้งก็ยอมแล้วแต่ลักษณะ Y จะเพิ่มขึ้นถึงจุดอิ่มตัวหรือยัง คือทางสถิติ ลักษณะนั้นจะต้องแสดง Significant ดังนั้นตัวอย่างนี้จะเห็นได้ว่า Deviation ไม่แสดง Significant หมายความว่ายังไม่มีโค้งหรือแตกแยกจากเส้นตรงที่จะถือเป็นสำคัญได้ แต่ Linear Regression หรือการเป็นเส้นตรงแสดง Highly Significant แสดงว่าเส้นยังคงตรงต่อไปอีกหรืออาจกล่าวได้ว่าถ้าใส่ปุ๋ยเพิ่มขึ้น ผลผลิตก็ยังเพิ่มขึ้นต่อไปอีกได้ ถ้าจะพิจารณาการใส่ปุ๋ยต่าง ๆ $S.S. \text{ Treatments} = S.S. \text{ Fertilizer} + S.S. \text{ Amount of Fertilizer}$ หมายความว่าความแตกต่างในเรื่องปุ๋ยประกอบด้วย ความแตกต่างระหว่างพวกที่ใส่ปุ๋ยกับไม่ใส่ปุ๋ยอย่างหนึ่ง อีกอย่างหนึ่งในระหว่างพวกที่ใส่ปุ๋ยด้วยกันแต่ใส่ต่างจำนวนกัน ในตัวอย่าง F ratio ของใส่ปุ๋ยกับไม่ใส่ปุ๋ยแสดง Highly Significant หมายความว่าพวกที่ใส่ปุ๋ยให้ผลต่างกับพวกไม่ใส่ปุ๋ยจริง เมื่อเราทราบแล้วก็หา F ratio ของจำนวนปุ๋ยซึ่งแสดง Highly Significant อีก แสดงว่าจำนวนปุ๋ยที่เพิ่มขึ้นแสดงผลผลิตแตกต่างกันอย่างแน่นอน และสหสัมพันธ์เป็นทางบวกจึงแสดงว่ายิ่งใส่ปุ๋ยเพิ่มขึ้นผลผลิตก็ยิ่งเพิ่มขึ้นด้วย

เมื่อเห็นประจักษ์แล้วว่าพวกที่ใส่ปุ๋ยด้วยกัน แสดงมีความแตกต่างกันโดยถือเป็นนัยสำคัญยิ่งทางสถิติ จึงจำเป็นต้องหาต่อไปอีกว่าปุ๋ยอัตราไหนให้ผลมากน้อยอย่างไร และแตกต่างกันอย่างเด่นชัดตลอดทุกคู่ไปหรือไม่

การเปรียบเทียบต่อไปนี้ D.F. = 1 เสมอเพราะเป็นการเปรียบเทียบภายใน ๑ คู่ หรือสองชนิด

การเปรียบเทียบผลผลิตระหว่างจำนวนปุ๋ยที่ใช้

๑. เปรียบเทียบอัตราปุ๋ย ๑๐๐ ก.ก. กับ ๒๐๐, ๓๐๐, ๔๐๐ ก.ก.

S.S. for A_{100} & A_{200} , A_{300} , A_{400}

$$= \frac{A_{100}^2}{b} + \frac{(A_{200} + A_{300} + A_{400})^2}{3b} - C.F.$$

$$= \frac{41^2}{5} + \frac{(63 + 89 + 101)^2}{15} - \frac{(41 + 63 + 101)^2}{20}$$

= 281.67 with 1 D.F. Variance also = 281.67

Calculated F ratio = $\frac{281.67}{14.86} = 18.95 **$

2. เปรียบเทียบอัตรา A_{200} A_{300} A_{400}

S.S. A_{200} กับ A_{300} , $A_{400} = \frac{A_{200}^2}{b} + \frac{(A_{300} + A_{400})^2}{2b} - C.F.$

$$= \frac{63^2}{5} + \frac{(89 + 101)^2}{10} - \frac{(63 + 89 + 101)^2}{15}$$

= 136.53 with 1 D.F., Variance = 136.53

Calculated F ratio = $\frac{136.53}{14.86}$

= 9.19 **

3. เปรียบเทียบ A_{300} กับ A_{400}

S.S. A_{300} กับ $A_{400} = \frac{A_{300}^2}{b} + \frac{A_{400}^2}{b} - C.F.$

$$= \frac{89^2}{5} + \frac{101^2}{5} - \frac{(89 + 101)^2}{10}$$

= 14.40 with 1 D.F., Variance = 14.40

Calculated F ratio = $\frac{14.40}{14.86}$ less than 1 (unsignificant)

หลักเกณฑ์ในการคำนวณตัวอย่างนี้ ขอเตือนไว้ว่าจงพยายามเข้าใจความมุ่งหมายของ Correction Factor (C.F.) ให้มาก เราจะเห็นได้ว่า C.F. นั้นได้จากการเอาผลรวมทั้งหมดยกกำลังสองแล้วหารด้วยจำนวนตัวเลขทั้งหมด ในการเปรียบเทียบที่แลมาแล้ว เรายกกันออกมาหา เพื่อต้องการวัดความแตกต่างทุก ๆ ระยะไป จากปุ๋ยอัตราต่ำมาหาอัตราสูง จาก ๑ ถึง ๓ จึงเห็นได้ว่าปุ๋ยอัตราต่ำที่วัดแล้วถูกตัดทิ้งเรื่อยมา ดังนั้น C.F. จึงต้องเปลี่ยนเรื่อยมา โดยใช้ตัวเลขที่อยู่ในขอบเขตของการเปรียบเทียบโดยเฉพาะเท่านั้น

ตามปกติ หลังจากหา S.S. แล้ว เราต้องเอา S.S. หารด้วย D.F. จึงจะเป็น Variance และ เอา Variance ของสิ่งที่ต้องการเปรียบเทียบหารด้วย Variance ของ Error จึงจะได้ Calculated F ratio แต่ในที่นี้ D.F. = 1 ดังนั้น S.S. จึงเท่ากับ Variance เพราะแม้จะหาร S.S. ด้วย ๑ ก็ได้ลงเดิม ดังนั้นจึงเอา S.S. หารด้วย Variance ของ Error ได้เลย

จากผลการเปรียบเทียบปรากฏว่า การใส่ปุ๋ยเพิ่มขึ้นจาก ๑๐๐ ก.ก. ให้ผลผลิตสูงขึ้นจริงคือ F ratio = 18.95 แสดง Highly Significant และเมื่อเพิ่มเป็น ๒๐๐ ก.ก. แล้วเทียบกับอัตราสูงขึ้นไปอีกก็ได้ F ratio = 9.19 แสดง Highly Significant แสดงว่าผลผลิตยังสูงขึ้นได้อีก แต่เมื่ออัตราจาก ๓๐๐ ก.ก. เป็น ๔๐๐ ก.ก. F ratio ไม่ถึงหนึ่ง แสดงว่าผลผลิตไม่มีความแตกต่างในเมื่อเพิ่มปุ๋ยจาก ๓๐๐ เป็น ๔๐๐ หมายความว่าถึงขีดสุดแล้ว ถ้าเพิ่มปุ๋ยขึ้นไปอีกจะแรงเกินไปและเป็นโทษ นอกจากจะเสียปุ๋ยไปเปล่า ๆ แล้วจะทำให้ผลผลิตตกต่ำลงอีกด้วย แต่ใน Regression นั้นจะเห็นได้ว่าในตาราง Analysis of Variance Deviation from Linear Regression ไม่ได้แสดงว่าจะหนีไปจากเส้นตรง (ไม่มี Significant) ทำให้ชักกันจึงขออธิบายไว้ด้วยว่า นั้นเป็นผลเฉลี่ยทั้งหมดของทุก ๆ Treatments ใน Experiment เพราะทั้ง Experiment ส่วนมากแสดงเส้นตรง แต่เมื่อปลายของเส้นเริ่มจะโค้ง Regression จึงไม่สามารถแสดงได้ละเอียด แต่เราก็มาพบในตอนเทียบกันว่าสุดท้ายไม่แสดงความแตกต่าง เพราะผลผลิตคงตัวและเริ่มจะลดลงในที่สุดท้ายนั่นเอง การวางระยะห่างระหว่างอัตราปุ๋ยกันนี้ว่าเป็นสิ่งสำคัญอันหนึ่ง ถ้าเราวาง

ระยะที่มากเกินไป การโค้งของเส้นตอนปลาย ๆ ก็อาจไม่แสดงให้เห็นปรากฏใน Regression ได้ แต่ก็สามารถทราบได้จากการเทียบคู่ในตอนท้าย เพราะอัตราชดกันมาขึ้นเอง แม้เส้นโค้งเพียงเล็กน้อยก็จะทำให้เกิดความแตกต่าง ถ้าช่วงห่างกันมาก Regression จะแสดงออกได้อย่างเด่นชัด แต่ก็ทำให้งานหายากมาก เพราะไม่รู้จักของส่วนโค้งที่ใกล้เคียงความจริง เพียงแต่ตรวจข้ามไปเป็นตอน ๆ จำเป็นต้องวางแผนการค้นคว้าใหม่ให้ระยะถี่ยิ่งขึ้น เพื่อตรวจรายละเอียดระหว่างส่วนที่ห่างอยู่อีกครั้งหนึ่ง

การวิเคราะห์ทางสถิติบอกได้แต่เพียงในการว่าปุ๋ยชนิดนั้นจำนวนเท่านี้ให้ผลดีที่สุดหรือเหมาะที่สุด แต่บางทีในตอนปลาย ๆ ของเส้นตรงที่เริ่มจะโค้ง ปุ๋ยมีราใส่ลงไปก็ยิ่งทำให้ผลผลิตเพิ่มขึ้นได้ แต่บางทีการเพิ่มของผลผลิตอาจน้อยลง ไม่ได้สัดส่วนกับปุ๋ยและแรงงานในการใส่ปุ๋ยที่เพิ่มขึ้นเป็นอัตราตายตัว จึงควรนำเอาตัวเลขเกี่ยวกับการเงินมาวิเคราะห์ด้วย เพื่อจะได้ทราบว่าเมื่อผลผลิตจะเพิ่มขึ้น แต่ต้นทุนสูงจนเกินส่วน ก็อาจสู้การใส่ปุ๋ยอัตราต่ำกว่าไม่ได้ จึงควรจะได้พิจารณาถึงผลกำไรที่แท้จริงของแต่ละ Treatment เป็นสำคัญ แต่อย่างไรก็ดีจำเป็นต้องคิดเป็นผลผลิตอย่างที่ทำกันมาแล้วเสียก่อน จึงเทียบเป็นเงินที่หลัง มิใช่คิดออกมาเป็นตัวเงินเลย เพราะค่าแรงและค่าใช้จ่ายต่าง ๆ ย่อมไม่เหมือนกันทุกแห่งเสมอไป จึงควรรายงานผลผลิตเป็นหลัก

ANALYSIS OF COVARIANCE

เราเคยศึกษากันมาแล้วถึงเรื่อง Analysis of Variance และเรื่อง Correlation ซึ่งต่างก็ได้ให้ความรู้ขึ้นเป็นรากฐานสำคัญในการวิเคราะห์ผลการค้นคว้า แต่ในบางโอกาสอาจจะต้องมีความจำเป็นในการที่จะนำทั้งสองวิธีการดังกล่าวแล้วมาใช้ร่วมกัน เพื่อวิเคราะห์ผลการค้นคว้าให้ได้ผลแน่นอนและกระจ่างแจ้งยิ่งขึ้น เราจะได้เห็นว่าในการวางแผนการค้นคว้าแบบ Randomized Block Design นั้น เรามีการควบคุมสภาพสิ่งแวดล้อมต่าง ๆ ให้ใกล้เคียงมากที่สุด เพื่อผลของการเปรียบเทียบจะได้แน่นอนเที่ยงตรง เช่น ในการเปรียบเทียบพันธุ์หรือเปรียบเทียบ Treatments เรามีการควบคุมความอุดมสมบูรณ์

ของคืนให้ใกล้เคียงกันโดยทำเป็น Block และเป็น Replication นอกจากนั้นยังควบคุมการปลูกปฏิบัติทั่ว ๆ ไปให้เหมือน ๆ กัน สิ่งเหล่านี้เราเรียกว่า Local Control และในการวิเคราะห์ผลการค้นคว้าด้วยวิธีการทางสถิติเราก็สามารถแยกเอาความแปรปรวนที่เกิดขึ้นเช่นที่ปรากฏจากความไม่สม่ำเสมอของคืนหรือ Block ออกจากได้จนกระทั่งเหลือ error สุดท้ายซึ่งเกิดจากสาเหตุอื่นๆ ที่เราไม่รู้

เรายังมีวิธีที่สามารถจะลด error ลงได้อีก ซึ่งเป็นการลด error ที่จะมากระทบกระเทือนทำให้ผลการเปรียบเทียบคลาดเคลื่อนไปจากความเป็นจริง โดยเหตุที่มี variation บางอย่างซึ่งเกิดจากการสุ่มที่ไม่ได้มีการควบคุมหรือไม่สามารถจะควบคุมแก้ไขได้และ variation นี้ก็รวมเข้าไปอยู่ใน error ของการทดลอง เราจึงจำเป็นจะต้องวัดสิ่งนี้แล้วนำมาใช้คำนวณแก้ไขผลผลิตให้ถูกต้อง วิธีนี้เราเรียกว่าควบคุมทางสถิติ "Statistical control of error" คือเป็นวิธีการควบคุมโดยการคำนวณทางสถิติ ดังนั้นนอกจากจะเก็บตัวเลขแสดงผลผลิตของแต่ละแปลงแล้วยังจะต้องเก็บสถิติตัวเลขของลักษณะที่เห็นว่ามีผลกระทบกระเทือนผลผลิตของแปลงนั้น ๆ เพื่อจะนำมาวิเคราะห์ว่าลักษณะหรือปรากฏการณ์นั้นจะมีผลกระทบกระเทือนผลผลิตหรือไม่ ถ้ามีจริงก็จะใช้คำนวณแก้ไขผลผลิตให้ถูกต้องเพื่อจะได้เปรียบเทียบกันด้วยความยุติธรรมและได้แน่นอน

ตัวอย่างเช่น เราตรวจสอบความอุดมสมบูรณ์ของดินแต่ละแปลงเสียก่อนว่ามีความสม่ำเสมอเพียงไร (Uniformity Trial) แล้วจึงทำการทดลองกันจริงๆ ทั่วแปลงในแปลงที่ตรวจสอบแล้ว เมื่อเสร็จการทดลองกันแล้วทุก ๆ แปลงเราก็มีตัวเลขอยู่ ๒ จำนวน จำนวนหนึ่งแสดงความอุดมสมบูรณ์ของดิน อีกจำนวนหนึ่งได้จากการค้นคว้าจริง ๆ เสร็จแล้วเราก็ใช้วิธีวิเคราะห์ทางสถิติตรวจสอบดูว่า ความไม่สม่ำเสมอของความอุดมสมบูรณ์ของดินนั้นจะกระทบกระเทือนผลผลิตหรือไม่ ถ้ากระทบกระเทือนเราก็ดำเนินการแก้ไขผลผลิตด้วยวิธีการคำนวณ การตรวจสอบความสม่ำเสมอของดินก็กระทำได้โดยการปลูกพืชชนิดเดียวกันที่ทำการทดลองไปแล้วเก็บผลผลิตเป็นแปลง ๆ ในเมื่อเราทำพืชอย่างเดียวกันหมด ปลูกปฏิบัติทดลองจนได้ปุ๋ยและบำรุงรักษาเหมือนกัน ถ้าแต่ละแปลงให้ผลผลิตกันไปก็หมายความว่าดินมีความอุดมสมบูรณ์

ไม่เหมือนกัน บางครั้งทำเลและสภาพของพื้นที่อาจทำให้เกิดความไม่สม่ำเสมอขึ้นได้ แม้ความอุดมสมบูรณ์จะใกล้เคียงกัน เช่นที่บางตอนลุ่มน้ำวังและ พืชก็ไม่งาม หรือบางตอนอยู่ใกล้บ่อน้ำพีช อาจได้รับความชุ่มชื้นดีกว่าก็จะงอกงามและให้ผลดีกว่า ฯลฯ เหล่านี้เป็นต้น ซึ่งถ้าได้อาศัยวิธีการตรวจสอบด้วย Uniformity trial แล้วคำนวณผลการทดลองจริง ๆ ด้วยวิธีการ Covariance ก็จะทำให้ผลการทดลองได้ผลแน่นอนยิ่งขึ้น ตัวอย่างอีกอันหนึ่งก็คือ สมมุติว่าเราทำการทดลองเปรียบเทียบพันธุ์หรือเปรียบเทียบปุ๋ยก็ดี ทุก ๆ แปลงจะต้องปฏิบัติเหมือนกัน เช่นเนื้อที่แปลงเท่ากัน ระยะปลูกเท่ากัน จำนวนต้นในแต่ละแปลงเท่ากัน ฯลฯ นอกจากสิ่งที่เราต้องการเปรียบเทียบ บางทีเราปลูกแปลงละเท่า ๆ กัน แต่ปลูกเสร็จแล้วบางต้นตายไป บางต้นก็ถูกศัตรูทำลาย และปลูกซ่อมไม่โดยผล เช่นอาจล่าช้าเกินไปทำให้เติบโตไม่ทันกัน ทำให้ผลผลิตแต่ละแปลงคลาดเคลื่อนไป ดังนั้นเมื่อเวลาเก็บผลผลิตก็ใช้วิธีเก็บตัวเลขสองจำนวนในแต่ละแปลง จำนวนหนึ่งคือจำนวนต้นที่ไถ่ลงในแปลงนั้นและอีกจำนวนหนึ่ง คือผลผลิต เสร็จแล้วเราก็นำมาวิเคราะห์ว่าจำนวนต้นทำให้ผลผลิตคลาดเคลื่อนไปจากความจริงหรือเปล่า ถ้าเป็นเช่นนั้นจริงก็นำเอาจำนวนต้นมาทำการปรับแก้ผลผลิตให้ถูกต้องจึงจะทำการเปรียบเทียบผลผลิตกันได้ ด้วยวิธีการคำนวณเพื่อให้ผลผลิตอยู่ในสภาพที่ไถ่มาจากจำนวนต้นเท่า ๆ กันแล้ว

ในการทดลองเกี่ยวกับสัตว์ โดยเฉพาะอย่างยิ่งการทดลองเปรียบเทียบความเจริญเติบโต เช่นการทดลองเรื่องอาหารสัตว์ก็ดี หรือการทดลองเปรียบเทียบพันธุ์สัตว์ ฯลฯ ก่อนทำการทดลอง สัตว์ย่อมจะต้องมีขนาดหรือน้ำหนักเฉพาะตัวของสัตว์เอง เราไม่สามารถจะจัดหาสัตว์ที่มีขนาดหรือน้ำหนักตัวเหมือนกัน ๆ กันทุกตัวได้ ย่อมจะต้องมีเล็กลงบ้างใหญ่บ้างตามสมควร น้ำหนักตัวก่อนทำการทดลองนี้เราเรียกว่า initial weight เมื่อเสร็จการทดลองแล้วจึงยังไม่สามารถเปรียบเทียบผลได้ เพราะน้ำหนักก่อนทำการทดลองไม่เท่ากัน จำเป็นจะต้องใช้น้ำหนักก่อนทดลองกับน้ำหนักเมื่อเสร็จการทดลองแล้วมาพิจารณาร่วมกัน โดยการวิเคราะห์ว่าการที่น้ำหนักก่อนทดลองไม่เท่ากันนั้นจะกระทบกระเทือนผลการทดลองหรือไม่ ถ้ากระทบกระเทือนก็ดำเนินการแก้ผลการทดลองให้ถูกต้องแล้วจึงเปรียบเทียบกัน โดยเฉพาะสัตว์มีผลกระทบกระเทือนเรื่องน้อยกว่าพืช เพราะถาขนาดและน้ำหนักดึกกัน ตัวใหญ่กว่าก็โตเร็วกว่าเพราะมี

กำลังในการแก่งแย่ง (competition) มาก ตัวเล็กกว่าก็จะยิ่งเล็กกว่ามากเข้า เพราะอาจถูกรบกวนถูกรังแกและแย่งอาหารสู้ตัวใหญ่ไม่ได้ ความสัมพันธ์ระหว่างน้ำหนักก่อนการทดลองกับหลังการทดลองจึงเด่นชัดกว่าในการทดลองทางพืช ดังนั้นวิธีการ Covariance จึงจำเป็นมากสำหรับการอธิบายวิธีการคำนวณเพื่อวิเคราะห์ผลการทดลองค้นคว้าแบบ Covariance นี้ใครจะขอยกตัวอย่างเพื่อประกอบการอธิบายซึ่งจะดำเนินไปเป็นตอน ๆ จึงจะเข้าใจแจ่มแจ้งยิ่งขึ้น

สมมติว่าในการเปรียบเทียบผลผลิตของพืช ๔ พันธุ์ คือพันธุ์ A, B, C และพันธุ์ D วางแผนการค้นคว้าแบบ Randomized Block Design ซึ่งแต่ละพันธุ์กระทำ 5 Replication ดังนั้นการทดลองนี้จึงแบ่งออกเป็น 5 Blocks เมื่อเก็บผลผลิตของแต่ละแปลงมาชั่งน้ำหนักเพื่อดำเนินการเปรียบเทียบพันธุ์ ปรากฏว่าแต่ละแปลงของแต่ละพันธุ์ที่ให้ผลผลิตนั้นมีจำนวนต้นไม่เท่ากัน เนื่องจากบางแปลงก็ตายมาก บางแปลงก็ตายน้อย จึงเกิดเป็นปัญหาขึ้นว่า ถ้าจำนวนต้นของแต่ละพันธุ์ไม่เท่ากันแล้วจะเป็นที่ยุติธรรมที่จะนำผลผลิตของพันธุ์เหล่านั้นมาเปรียบเทียบกัน จึงจำเป็นต้องใช้วิธีการ Covariance ตรวจสอบว่าการที่จำนวนต้นไม่เท่ากันนั้นจะกระทบกระเทือนผลผลิตหรือไม่ ถ้ากระทบกระเทือนก็ต้องคำนวณแก้ไขผลผลิตให้อยู่ในสภาพที่มีจำนวนต้นเท่า ๆ กันเสียก่อน ส่วนวิธีการก็กระทำได้โดยนอกจากการจัดบันทึกน้ำหนักผลผลิตของแต่ละแปลงแล้วจะต้องนับจำนวนต้นแล้วจดสถิติตัวเลขแสดงจำนวนต้นของแต่ละแปลงควบไปกับผลผลิตของแปลงนั้น ๆ ด้วย ฉะนั้นในแปลงทุกแปลงจึงมีตัวเลขประจำ ๒ จำนวน จำนวนหนึ่งแสดงน้ำหนักของผลผลิต อีกจำนวนหนึ่งแสดงจำนวนต้น ทั้งสองจำนวนนี้ต้องคู่กันและอยู่ประจำแปลงโดยเฉพาะ จะสลับกับแปลงอื่นหรือคู่อื่นใดไม่ได้ทั้งสิ้น เพราะทั้งจำนวนต้นและผลผลิตได้จากแปลงนั้นโดยเฉพาะ และต่างฝ่ายต่างควบคุมซึ่งกันและกัน

ดังนั้นในตัวอย่างที่จะกล่าวต่อไปนี้ แต่ละแปลงมีตัวเลขอยู่ ๒ จำนวนซึ่งแสดงสองลักษณะ คือ $X = \text{Yield}$ หรือผลผลิตและอีกจำนวนหนึ่งคือ $Y = \text{Stand}$ หรือจำนวนต้น

Block	X=Yield	Varieties				Block Total	Square of Block Total
	Y=Stand	A	B	C	D		
I	X	28	23	23	27	101	10201
	Y	24	21	17	28	90	8100
II	X	25	26	16	26	93	8649
	Y	21	28	13	26	38	7744
III	X	29	26	19	26	100	10000
	Y	28	27	13	27	95	9025
IV	X	27	29	19	30	105	11025
	Y	27	26	17	24	94	8836
V	X	33	30	30	30	123	15129
	Y	25	26	21	23	95	9025
Varieties	X	142	134	107	139	522 (Grand total X)	
Total	Y	125	128	81	128	462 (Grand total Y)	
Square of	X	20164	17956	11449	19321		
Varieties	Y	15625	16384	6561	16384		
Total							

โดยเหตุที่การวางแผนทดลองครั้งนี้เป็นแบบ Randomized Block Design ซึ่งกล่าวมาแล้วในตอนต้น จากความที่ได้เรียนมาแล้วในเรื่อง Randomized Block Design นั้นเอง ถ้าเราพิจารณากันโดยถ่วงน้ำหนักจะเห็นว่าตารางที่แสดงอยู่ข้างบนนี้เป็นตาราง Randomized Block ๒ ตารางทับกันอยู่คือ X ตารางหนึ่ง กับ Y อีกตารางหนึ่ง ถ้าเราจะแยกออกเป็นตารางของ X อันหนึ่งกับตาราง Y อีกอันหนึ่งก็กระทำได้ เพื่อผู้ศึกษาใหม่จะได้เข้าใจง่ายขึ้น แต่จำเป็นต้องเินดและยกเยื่อไปโดยใช้เหตุ เพราะตารางที่มีอักษร X และ Y กำกับไว้แล้ว เมื่อจะวิเคราะห์แต่เรื่องใดก็ดูแต่เรื่องนั้น การที่สามารถแยกออกได้เป็น ๒ ตารางนี้ได้หมายความว่าสิ่งที่จะพิจารณาเพียง ๒ ปัญหา โดยเหตุที่ X กับ Y

จะต้องคู่กันไปเสมอ จะสลับคู่มิได้ จึงเกิดเป็น ๓ ปัญหาขึ้น ดังจะแยกให้ดูเป็น ๓ ข้อต่อไป

(๑) ปัญหาเรื่อง X คือต้องวิเคราะห์ว่าผลผลิตของแต่ละพันธุ์แตกต่างกันจริงหรือไม่ การคำนวณก็ใช้ Analysis of Variance แบบ Randomized Block ซึ่งรวมค่าทั้งเรียนมาแล้วในตอนต้น แยกเฉพาะตัวเลขในตารางของ X เท่านั้น

(๒) ปัญหาเรื่อง Y เราก็วิเคราะห์เช่นเดียวกับ X แต่จะเฉพาะตัวเลขในตารางของ Y เพื่อต้องการวิเคราะห์ว่าแต่ละพันธุ์มีจำนวนต้นต่างกันหรือไม่ เนื่องจากเป็นการเปรียบเทียบในเรื่องของ Y โดยเฉพาะการวิเคราะห์จึงไม่เกี่ยวข้องกับ X คงใช้ตัวเลขในตาราง Y ดำเนินการวิเคราะห์ Analysis of Variance แบบ Randomized Block System อย่างที่เรียนมาแล้วในตอนต้น

(๓) ปัญหาเรื่อง X กับ Y ร่วมกัน ในตอนนี้เราจำเป็นจะต้องวิเคราะห์ทั้ง X และ Y ร่วมกัน โดยใช้วิธีการทำนองเดียวกันกับ Analysis of Variance หากแต่เป็นการวิเคราะห์ความแตกต่างอันเกิดจากทั้ง X และทั้ง Y ร่วมกัน ดังนั้นตอนใดที่เป็นการหา Sum of square (S.S.) เราก็หา Sum of Product of X and Y (S.P._{XY}) แทน คือแทนที่จะใช้รายการนั้น ๆ ของตัวเอง แต่ใช้ X ของรายการนั้นคูณกับ Y ในรายการเดียวกันต่อไปนี้จะได้แสดงการคำนวณเป็นขั้น ๆ ตามลำดับหัวข้อทั้ง ๓ หัวข้อที่ได้กล่าวมาแล้ว

(๑) การวิเคราะห์ความแตกต่างระหว่างผลผลิต (X) ของแต่ละพันธุ์

$$C.F._X = \frac{(\text{Grand total of } X)^2}{n} = \frac{522^2}{20} = 13624.2$$

$$S.S.\text{total}_X = X_1^2 + X_2^2 + X_3^2 + X_4^2 + \dots + X_{20}^2 - C.F._X$$

$$= 28^2 + 25^2 + 29^2 + \dots + 29^2 + 30^2 + 30^2 - 13624.2 = 353.8$$

(with 20-1=19 D.F.)

$$S.S.\text{Block}_X = \frac{B_X^2 + B_{X_2}^2 + B_{X_3}^2 + B_{X_4}^2 + B_{X_5}^2}{V} - C.F._X$$

$$= \frac{101^2 + 93^2 + 100^2 + 105^2 + 123^2}{4} - 13624.2$$

$$= 126.8 \text{ with } 5-1 = 4 \text{ D.F.}$$

๑(๒๕)

$$\begin{aligned} \text{S.S. Varieties}_X &= \frac{v_{XA}^2 + v_{XB}^2 + v_{XC}^2 + v_{XD}^2}{b} - \text{C.F.}_X \\ &= \frac{142^2 + 134^2 + 107^2 + 139^2}{5} - 13624.2 \\ &= 153.8 \text{ with } 4-1 = 3 \text{ D.F.} \end{aligned}$$

$$\begin{aligned} \text{S.S. Error}_X &= \text{S.S. Total}_X - (\text{S.S. Blocks}_X + \text{S.S. Varieties}_X) \\ &= 353.8 - (126.8 + 153.8) \\ &= 73.2 \\ &\text{with } 19 - (4 + 3) = 12 \text{ D.F.} \end{aligned}$$

Analysis of Variance (for X)

Sources of Difference	D.F.	S.S.	Variance	F ratio	
				Calculated	Table
Total	19	353.8			
Block	4	126.8			
Varieties	3	153.8	51.3	8.41**	
Error	12	73.2	6.1		

**

ปรากฏว่าผลผลิตของแต่ละพันธุ์แตกต่างกันจริง

(๒) การวิเคราะห์ความแตกต่างระหว่างจำนวนต้น (Y) ของแต่ละพันธุ์

$$C.F._Y = \frac{(\text{Grand total } Y)^2}{n}$$

$$= \frac{462^2}{20} = 10672.2$$

$$\begin{aligned} S.S.Total_Y &= Y_1^2 + Y_2^2 + Y_3^2 + Y_4^2 + \dots + Y_{20}^2 - C.F._Y \\ &= 24^2 + 21^2 + 28^2 + \dots + 27^2 + 24^2 + 23^2 - 10672.2 \\ &= 439.8 \text{ with } 20-1 = 19 \text{ D.F.} \end{aligned}$$

$$\begin{aligned} S.S.Block_Y &= \frac{B_{Y1}^2 + B_{Y2}^2 + B_{Y3}^2 + B_{Y4}^2 + B_{Y5}^2}{v} - C.F._Y \\ &= \frac{90^2 + 88^2 + 95^2 + 94^2 + 95^2}{4} - 10672.2 \\ &= 10.3 \text{ with } 5-1 = 4 \text{ D.F.} \end{aligned}$$

$$\begin{aligned} S.S.Varieties_Y &= \frac{V_{YA}^2 + V_{YB}^2 + V_{YC}^2 + V_{YD}^2}{b} - C.F._Y \\ &= \frac{125^2 + 128^2 + 81^2 + 128^2}{5} - 10672.2 \\ &= 318.6 \text{ with } 4 - 1 = 3 \text{ D.F.} \end{aligned}$$

$$\begin{aligned} S.S.Error_Y &= S.S. Total_Y - (S.S Block_Y + S.S.Varieties_Y) \\ &= 439.8 - (10.3 + 318.6) \\ &= 110.9 \text{ with } 19 - (4 + 3) = 12 \text{ D.F.} \end{aligned}$$

ตาราง Analysis of Variance (for Y)

เพื่อวิเคราะห์ความแตกต่างของจำนวนต้น

Sources of Difference	D.F.	S.S.	Variance	F. ratio	
				Calculated	Table
Total	19	439.8			
Block	4	10.3			
Varieties	3	318.6	106.2	11.49 **	
Error	12	110.9	9.24		

ตารางวิเคราะห์ความแตกต่างของจำนวนต้น แสดงว่าแต่ละพันธุ์จำนวนต้นแตกต่างกันจริง และจากตารางวิเคราะห์ความแตกต่างของผลผลิตในข้อหนึ่ง ก็แสดงว่าแต่ละพันธุ์ให้ผลผลิตไม่เท่ากัน ดังนั้นจึงทำให้เกิดปัญหาขึ้นว่า การที่ผลผลิตของแต่ละพันธุ์ไม่เท่ากันนั้น อาจมีสาเหตุอันเนื่องมาจากความแตกต่างในเรื่องพันธุ์ก็ได้ จึงอยากถามหาว่าพันธุ์ที่ให้ผลผลิตสูงจะดีกว่าพันธุ์ที่ให้ผลผลิตต่ำ เพราะจำนวนต้นของแต่ละพันธุ์ไม่เท่ากัน ดังนั้นการที่ผลผลิตต่างกันอาจเนื่องมาจากจำนวนต้นไม่เท่ากันก็ได้ จึงจำเป็นต้องทำการวิเคราะห์ต่อไปอีก

(๓) การวิเคราะห์ความสัมพันธ์ระหว่างจำนวนต้นกับผลผลิต เพื่อตรวจสอบว่าการที่ผลผลิตแตกต่างกันนั้นเนื่องมาจากจำนวนต้นไม่เท่ากันจริงหรือไม่

วิธีการคำนวณ กระทำเช่นเดียวกันกับที่แล้วมา หากแต่เรื่องนี้เป็นเรื่องของ X กับ Y ร่วมกัน ดังนั้นแทนที่จะยกกำลังสอง (หรือคูณตัวเอง) ใช้ X กับ Y ในช่องเดียวกันคูณกัน เช่นการหา C.F. โดยปกติเราเอาผลรวมทั้งหมดมายกกำลังสอง แล้วหารด้วยจำนวนตัวเลข แต่ในเรื่องนี้ C.F. ใช้ผลรวมของ X ทั้งหมดคูณกับผลรวมของ Y ทั้งหมดแล้วหารด้วยจำนวนตัวเลขรายการอื่น ๆ ก็เช่นเดียวกัน ดังจะได้แสดงให้เห็นต่อไปนี้

$$C.F._{XY} = \frac{(\text{Grand total } X)(\text{Grand total } Y)}{n}$$

$$= \frac{522 \times 462}{20} = 12058.2$$

$$S.P._{XY} \text{ Total} = X_1Y_1 + X_2Y_2 + X_3Y_3 + X_4Y_4 + \dots + X_{20}Y_{20} - C.F._{XY}$$

$$= (28 \times 24) + (25 \times 21) + \dots + (30 \times 24) + (30 \times 23) - 12058.2$$

$$= 294.8 \text{ with } 20-1 = 19 \text{ D.F.}$$

$$S.P._{XY} \text{ Block} = \frac{B_{X_1}B_{Y_1} + B_{X_2}B_{Y_2} + B_{X_3}B_{Y_3} + B_{X_4}B_{Y_4} + B_{X_5}B_{Y_5}}{v} - C.F._{XY}$$

$$= \frac{(101 \times 90) + (93 \times 88) + (100 \times 95) + (105 \times 94) + (123 \times 95)}{4} - 12058.2$$

$$= 24.0 \text{ with } 5-1 = 4 \text{ D.F.}$$

$$S.P._{XY} \text{ Varieties} = \frac{V_{X_A}V_{Y_A} + V_{X_B}V_{Y_B} + V_{X_C}V_{Y_C} + V_{X_D}V_{Y_D}}{b} - C.F._{XY}$$

$$= \frac{(142 \times 125) + (134 \times 128) + (107 \times 81) + (139 \times 128)}{5} - 12058.2$$

$$= 214.0 \text{ with } 4-1 = 3 \text{ D.F.}$$

$$S.P._{XY} \text{ Error} = S.P._{XY} \text{ Total} - (S.P._{XY} \text{ Blocks} + S.P._{XY} \text{ Varieties})$$

$$= 294.8 - (24.0 + 214.0)$$

$$= 56.8 \text{ with } 19 - (4 + 3) = 12 \text{ D.F.}$$

เสร็จแล้วจึงทำตาราง Analysis of Variance รวมทั้ง ๓ เรื่อง คือ
เรื่อง X เรื่อง Y และเรื่องความสัมพันธ์ระหว่าง X กับ Y ดังต่อไปนี้

Analysis of Variance

Sources of Difference	D.F.	S.S. _X	S.S. _Y	S.P. _{XY}
Total	19	353.8	439.8	294.8
Blocks	4	126.8	10.3	24.0
Varieties	3	153.8	318.6	214.0
Error	12	73.2	110.9	56.8

ต่อไปเราก็คำนวณการหา Correlation Coefficient "r" ดังต่อไปนี้

$$\text{สูตร } r = \frac{\text{S.P.}_{xy}}{\sqrt{\text{S.S.}_X \cdot \text{S.S.}_Y}}$$

การแทนค่าในสูตร ใช้ค่าของ Error ทั้งหมด
ดังนั้น

$$r = \frac{56.8}{\sqrt{73.2 \times 110.9}} = 0.63 \star (\text{Significant})$$

เปิด "r" Table โดย D.F. ที่ n' - 2

(n' = D.F. of Error = 12) ฉะนั้น 12 - 2 = 10 r 0.576

เพราะฉะนั้น ค่าของ Calculated r = 0.63 จึงแสดงว่า X กับ Y มีความสัมพันธ์กันจริง คือการที่จำนวนต้นไม้มากหรือน้อยทำให้ผลผลิตไม้ไปด้วยจริง เราจำเป็นต้องตรวจสอบทาง Regression อีกทีหนึ่ง แล้วจึงไปดำเนินการแก้ผลผลิตที่ผิดไปจากความจริงใหญ่ถูกต้อง

$$\text{Regression Coefficient} = \frac{\text{S.P.}_{xy}}{\text{S.S.}_y} = \frac{56.8}{110.9} = 0.51 = b_{xy}$$

เมื่อเราหา b_{xy} ได้แล้ว ก็ต้องหา S.E. ของ b_{xy} ด้วย เพื่อจะได้นำไปหา t ratio สำหรับใช้ตรวจสอบ regression

$$\begin{aligned} \text{S.E. } b_{xy} &= \sqrt{\frac{\text{S.S.}_x - (b_{xy}^2 \times \text{S.S.}_y)}{(n' - 2) (\text{S.S.}_y)}} \\ &= \sqrt{\frac{73.2 - (0.51^2 \times 110.9)}{(12-2) (110.9)}} \\ &= 0.19 \end{aligned}$$

$$\begin{aligned} \text{"t" ratio} &= \frac{b_{xy}}{\text{S.E. } b_{xy}} = \frac{0.51}{0.19} \\ &= 2.684 \end{aligned}$$

เมื่อเปิด "t" Table ได้ 2.201 แสดงว่า Significant หรือ 3.106
กล่าวคือว่าจำนวนต้นทำให้ผลผลิตผิดไปจริง จึงดำเนินการแก้ผลผลิตเพื่อให้อยู่ในสภาพที่มีจำนวนต้นเท่า ๆ กัน จะได้ทำการเปรียบเทียบผลผลิตกันได้แน่นอน
สูตรที่ใช้ในการแก้ผลผลิตมีดังนี้

$$\text{Corrected Yield} = M_x - b_{xy} (M_y - Y)$$

M_x = Observed mean X of each variety

เป็นผลเฉลี่ยของผลผลิตของแต่ละพันธุ์ ที่ได้จากการทดลองจริง ๆ

M_y = Observed mean Y of each variety

เป็นผลเฉลี่ยของผลผลิตของแต่ละพันธุ์ ที่ได้จากการทดลองจริง ๆ

Y = ผลเฉลี่ยของจำนวนต้นทั้งหมดในการทดลองนั้น

Varieties	M_y	$M_y - Y$	$b_{xy}(M_y - Y)$	Observed M_x	Corrected Yields
A	25.0	1.9	0.969	28.4	27.431
B	25.6	2.5	1.275	26.8	25.525
C	16.2	-6.9	-3.519	21.4	24.919
D	25.6	2.5	1.275	27.8	26.525

Mean ของ Variety A ที่ได้จาก Experiment = 28.4 จาก Experiment นี้แสดงว่า M_x ของ Variety A สูงไป ๐.๙๖๙ จึงเอา ๒๘.๔ - ๐.๙๖๙ = ๒๗.๔๓๑ ส่วน Variety B & Variety D ก็เช่นเดียวกัน แต่ M_x ของ Variety C แสดงว่าต่ำไป ๓.๕๑๙ เพราะมีเครื่องหมายลบอยู่หน้า จึงต้องเอาเพิ่ม M_x เข้าไปอีก คือ ๒๑.๔ + ๓.๕๑๙ = ๒๔.๙๑๙ เสร็จแล้วจึงเอา Corrected Yield นี้ไปเปรียบเทียบกันแบบ Analysis of Variance แต่เนื่องจากผลผลิตเดิมจึงทำให้ Variance ที่หาได้ผิดควยจึงจำต้องทำตารางแก้ Variance อีกทีหนึ่ง ดังต่อไปนี้

Sources of Difference	D.F.	S.S. _x	S.P. _{xy}	S.S. _y	Values of Reduced Variance		
					Deviation Reduced S.S.	D.F.	Reduced $\frac{S.S.}{D.F.}$
Varieties	3	153.8	214.0	318.6			(Variance)
Error	12	73.2	56.8	110.9	44.11	11	4.01
Var. + Error	15	227.0	270.8	429.5	56.26	14	
					12.15	3	4.05

Reduced S.S. for varieties =

$$\begin{aligned} & (\text{Reduced S.S.} + \text{var. Error}) - (\text{Reduced S.S. Error}) \\ & = 56.26 - 44.11 = 12.15 \end{aligned}$$

$$\begin{aligned} \text{Reduced S.S. Error} &= \text{S.S.}_x \text{ for error} - \frac{\text{S.P.}_{xy}^2 \text{ for error}}{\text{S.S.}_y \text{ for error}} \\ &= 73.2 - \frac{56.8^2}{110.9} \\ &= 44.11 \end{aligned}$$

สำหรับ Reduced D.F. นั้นให้เอา D.F. เก็บคณด้วยหนึ่งเสมอ

$$\begin{aligned} \text{Reduced S.S. for Varieties} &= \text{S.S.}_x \text{ var.} + \text{error} - \frac{\text{S.P.}_{xy}^2 \text{ var.} + \text{error}}{\text{S.S.}_y \text{ var.} + \text{error}} \\ &= 227.0 - \frac{271.8^2}{429.5} \\ &= 56.26 \end{aligned}$$

$$\text{Reduced Variance} = \frac{\text{Reduced S.S.}}{\text{Reduced D.F.}}$$

$$\text{for Variance} = \frac{12.15}{11} = 4.01$$

$$\text{for Error} = \frac{44.11}{11} = 4.01$$

เสร็จแล้วก็หา Calculated F ratio เปรียบเทียบ หากแต่ใช้ตัวเลขในชุดที่แก้ไข
(Reduced value) แล้วทั้งหมด

$$\begin{aligned} \text{F ratio} &= \frac{\text{Reduced Variance of Varieties}}{\text{Reduced Variance of Error}} \\ &= \frac{4.05}{4.01} = 1.00 \text{ insignificant} \end{aligned}$$

(เปิด Table "t" คงดูที่ Reduced D.F. ของ Varieties และ Reduced D.F. ของ Error)

ไม่มีความแตกต่างที่เชื่อถือเป็นสำคัญได้ แสดงว่า เบื่อไก่แก๊สชนิดนี้ให้อยู่ในสภาพที่มีจำนวนตนเท่า ๆ กันแล้ว ปรากฏว่าผลผลิตของแต่ละพันธุ์ได้แสดงความแตกต่างกันเลย การที่

ผลผลิตของแต่ละพันธุ์แสดงความแตกต่างในตอนแรกนั้น เนื่องมาจากจำนวนต้นไม่เท่ากัน มิได้
เนื่องมาจากความแตกต่างระหว่างพันธุ์

ถ้าหากเราต้องการจะตรวจสอบว่า เมื่อเราแก้จำนวนต้นให้เท่ากันแล้ว Error จะลดลง
อีกเท่าไร เราก็กระทำดังนี้

$$\begin{aligned} \text{Precision ratio} &= \frac{\text{Variance of Error}}{\text{Reduced variance of error}} \\ &= \frac{6.1}{4.01} = 1.52 \end{aligned}$$

ในตัวอย่างนี้ยังเห็นพันธุ์ไม่แสดงความแตกต่าง จึงสรุปได้เพียงแค่นี้ แต่สมมุติว่าถ้าหาก
กระทำมาถึงขั้นนี้ ตรวจพบว่า "F" ratio มี significant ซึ่งแสดงว่าพันธุ์มีความแตกต่าง
กันจริง เราก็จะต้องกำหนดค่าต่อไปอีกว่าพันธุ์ไหนต่างกับพันธุ์ไหน โดยใช้ผลเฉลี่ยที่แก้ไขจำนวน
ต้นเท่า ๆ กันแล้ว (Corrected Mean Yield)

สมมุติว่าเราต้องการจะเปรียบเทียบพันธุ์ A กับพันธุ์ C ก่อนอื่นเราต้องหา S.E.D
ซึ่งเป็น S.E. ของความแตกต่างระหว่างผลผลิตที่แก้ไขจำนวนต้นเท่ากันแล้ว ของพันธุ์ A
กับพันธุ์ C ให้ E เป็น Variance of Error ที่แก้ไขแล้ว ซึ่งในที่นี้ = ๔.๐๑ n เป็น
จำนวน plot ของแต่ละ Mean = 5 D เป็นผลต่างระหว่างผลเฉลี่ยของจำนวนต้น ของพันธุ์
ที่นำมาเปรียบเทียบกัน ซึ่งในที่นี้เราเทียบพันธุ์ A & C ผลเฉลี่ยจำนวนต้นของพันธุ์ A =
25.0 ผลเฉลี่ยจำนวนต้นของพันธุ์ C = 16.2

$$\begin{aligned} D_y \text{ จึง} &= 25.0 - 16.2 = 8.8 \\ \text{S.E.}_D &= \sqrt{E \left(\frac{2}{n} + \frac{D_y^2}{\text{S.S.}_y \text{ of Error}} \right)} \\ &= \sqrt{4.01 \left(\frac{2}{5} + \frac{8.8^2}{110.9} \right)} \\ &= 2.1 \end{aligned}$$

$S.E._D = 2.1$ เป็น $S.E._D$ ของผลต่างระหว่างพันธุ์ A และ C เราเอาผลผลิตที่แก้แล้วมาเทียบกัน เพื่อหาผลต่างหรือ "D"

$$\text{Corrected Mean of Variety A} = 27.4$$

$$\text{Corrected Mean of Variety C} = 24.9$$

$$D = 27.4 - 24.9 = 2.5$$

$$\begin{aligned} \text{Calculated "t"} &= D / S.E._D = 2.5 / 2.1 \\ &= 1.2 \text{ insignificant} \end{aligned}$$

แล้วไปเปิด "t" Table เพื่อตรวจสอบ โดยใช้ D.F. ของพันธุ์ A บวกกับ D.F. ของพันธุ์ C เช่นเดียวกันกับการเปรียบเทียบธรรมดา ถ้าต้องการจะเทียบพันธุ์อื่นก็กระทำเช่นเดียวกัน

Technique เรื่อง Covariance เราเอาไปใช้ได้หลายอย่าง แต่ไม่ควรใช้เมื่ออยู่กรณีหนึ่ง คือ เมื่อเราสุ่มสาเหตุที่เราจะเอามาหา Covariance นั้นมันสืบสายพันธุ์ เช่น ถ้าเรา variety C นั้นเพาะงอกน้อยเสมอและความงอกเวลานี้สืบทอดลักษณะไปถึงชั่วลูกชั่วหลานเรื่อยไป แสดงว่าเกี่ยวกับคุณลักษณะประจำพันธุ์ของมันเอง เราแก้ด้วย Covariance ไม่ได้เพราะมันเป็นพันธุ์ที่แท้จริง ๆ มิใช่เกิดจากสาเหตุภายนอก เราจะใช้ Covariance ก็ได้ก็แต่เฉพาะที่มันตายโดยได้รับอันตราย หรือการปลูกไม่ดี หรือเมล็ดสูญหายถูกศัตรูทำลาย เป็นต้น ถ้าเป็นการทดลองเกี่ยวกับสัตว์ เช่นการทดลองเกี่ยวกับการให้อาหาร (Feeding) เราไม่สามารถหาสัตว์ที่มีน้ำหนักเท่า ๆ กันได้ และเราก็ไม่สามารถหา Replications ได้ตั้งร้อยตัวพัน ฉะนั้นเราต้องหาโดยชั่งน้ำหนักเดิมของแต่ละตัวก่อนเพื่อมาคิดหา Covariance กับอาหาร แต่บางทีเกิดมีบางตัวกินเก่ง บางตัวกินไม่เก่ง ต้องมี Multiple Covariance หรือ Covariance แบบซับซ้อนขึ้นไปอีก

SPLIT PLOT DESIGN

ในการทดลองกันกว่า เพื่อเปรียบเทียบนั้น ที่ใกล้เข้ามาแล้ว เป็นเพียงการกันกว่า เรื่องเดียวหรือปัญหาเดียวโดยเฉพาะ เช่นการทดลองเปรียบเทียบพันธุ์กุ่มแก่ป่าหาเรื่องพันธุ์อย่างเดียว หรือการทดลองเรื่องปุ๋ยกุ่มที่จะแก่ป่าหาเรื่องปุ๋ยอย่างเดียวนั้น มีได้คำนึงถึงเรื่องพันธุ์ด้วย ในบางโอกาสเราจำเป็นต้องแก่ป่าหาสองปัญหาพร้อม ๆ กันไป เพื่อตัดทอนความหมดเปลืองทั้งในด้านเวลาและทุนรอน ซึ่งก็เท่ากับยอมเอาการกันกว่าสองเรื่อง เข้ามารวมกันเป็นเรื่องเดียว แต่ถ้ามองจะพิจารณาจริงๆ แล้ว การที่นำป่าหาสองปัญหาเข้ามารวมกันในการทดลองกันกว่านั้น มิใช่จะทำให้การทดลองกันกว่าเรื่องนั้นขยายตัวออกเป็นสองปัญหาเท่านั้น หากแต่ทำให้เกิดเป็นสามปัญหานั้นได้ ดังจะยกตัวอย่างต่อไปนี้

สมมุติว่าเราทำการทดลองเปรียบเทียบพันธุ์พืชหลาย ๆ พันธุ์ และในขณะเดียวกันก็ทำการทดลองเปรียบเทียบปุ๋ยในอัตราส่วนผสมต่าง ๆ กันด้วย ถ้าหากเรามองแต่เพียงผิวเผินก็จะถึงความเห็นว่าการทดลองมีความมุ่งหมายที่จะเฉลยป่าหาเรื่องพันธุ์เรื่องหนึ่ง และเกี่ยวกับปุ๋ยอีกเรื่องหนึ่ง แต่ผลงานอันแท้จริงนั้นมีโอกาสที่จะเป็นไปได้ถึงสามกรณีด้วยกัน ดังจะจำแนกให้เห็นเป็นข้อ ๆ ต่อไปนี้

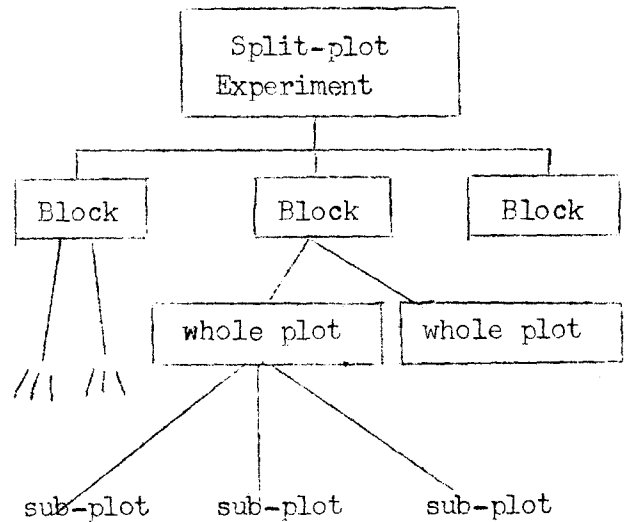
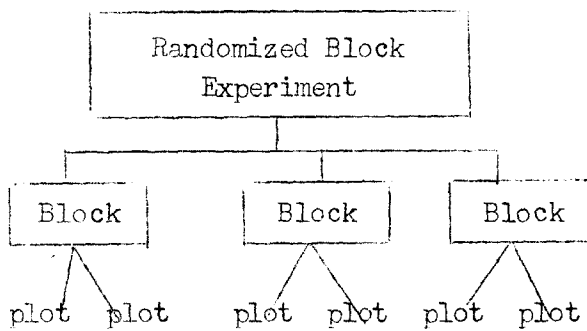
ผลงานแสดงความแตกต่างเรื่องปุ๋ยโดยเฉพาะ ก็โดยเหตุที่มีการทดลองเปรียบเทียบพันธุ์อยู่ด้วย ดังนั้นจึงต้องมีการเฉลยป่าหาในเรื่องพันธุ์เป็นธรรมดา

เนื่องจากการทดลองปุ๋ยและทดลองเปรียบเทียบพันธุ์กระทำร่วมกันไปในการทดลองเดียวกัน ดังนั้นทุก ๆ พันธุ์จึงมีโอกาสที่จะร่วมกับปุ๋ยทุกชนิดซึ่งโดยปกติสมมุติว่าถ้าเราจะทำการทดลองเปรียบเทียบปุ๋ยสี่อัตรา เราก็ใช้พืชชนิดเดียวพันธุ์เดียว เมื่อสรุปผลการทดลองแล้วถ้าปุ๋ยอัตราใดให้ผลดี เราจะกล่าวว่าปุ๋ยอัตรานั้นให้ผลดีเฉย ๆ ไม่ได้ เราจำเป็นต้องกล่าวเน้นลงไปอีกว่า ในการทดลองเปรียบเทียบปุ๋ยสี่อัตรา สำหรับปลูกพืชพันธุ์นั้น ปรากฏว่าปุ๋ยอัตรานั้นให้ผลดีที่สุด ซึ่งแสดงว่าผลของการใช้ปุ๋ยนั้นขึ้นอยู่กับพันธุ์พืชด้วย ยิ่งใน Split-plot Design ซึ่งเป็นการทดลองที่ทั้งปัญหาเรื่องปุ๋ยและปัญหาเรื่องพันธุ์อยู่ด้วยกัน จึงเป็นการไม่แน่นอนเสมอไปว่า ถ้าหากปุ๋ยชนิดนี้ให้ผลดีสำหรับพืชพันธุ์นั้นแล้ว เมื่อนำไปใช้กับพืชอีกพันธุ์หนึ่งจะได้รับผลดีอยู่ใน

อันคับเกินเรื่อยไป เช่นเมื่อเปรียบเทียบปุ๋ยสูตร ๖ กับปุ๋ยสูตร ๙ อาจให้ผลดีในเมื่อใช้ปุ๋ยเบอร์ ๒ แต่หาพืชพันธุ์ ๖ ปุ๋ยเบอร์ ๔ อาจให้ผลดีที่สุดก็ได้ การที่พืชพันธุ์หนึ่ง ๆ เหมาะกับปุ๋ยอย่างหนึ่ง โดยเฉพาะ หรือพืชคนละพันธุ์เหมาะกับปุ๋ยคนละอย่าง ตลอดจนพืชคนละพันธุ์เหมาะกับวิธีการปฏิบัติคนละอย่าง เหล่านี้เราเรียกว่า Interaction ระหว่างพันธุ์พืชกับปุ๋ย หรือระหว่างพันธุ์พืชกับวิธีการปฏิบัติ เช่นพืชพันธุ์ชอบดินหรือสิ่งปลูกอย่างหนึ่ง พืชอีกอย่างหนึ่งก็ชอบดินหรือสิ่งปลูกอีกอย่างหนึ่ง เป็นต้น ในกรณีที่เกี่ยวข้องกับสัตว์ก็มี Interaction มากเหมือนกัน เช่นเกี่ยวกับดินฟ้าอากาศและแฟกเตอร์แวดล้อมต่าง ๆ ตลอดจนกระทั่งเกี่ยวกับส่วนผสมของอาหาร รสชาติของอาหาร กลิ่นของอาหาร ตัวอย่างง่าย ๆ เช่น บางคนชอบอาหารรสเค็ม บางคนชอบจืด บางคนชอบอากาศหนาว บางคนชอบอากาศร้อน ฯลฯ สรุปแล้วพืชหลาย ๆ พันธุ์ที่นำมาใช้ปุ๋ยหลาย ๆ อย่างก็ย่อมจะมี Interaction ได้ มีใจว่าปุ๋ยชนิดนี้ใช้ได้ผลดีกับพืชพันธุ์นั้นแล้วจะใช้ได้ผลดีกับพืชพันธุ์อื่น ๆ ด้วย

ในการวางแผนวิจัยสองปัญหาคงกันไปโดยใช้แบบ Split-plot นั้น เราจะต้อง random ปัญหาหนึ่งลงในอีกปัญหาหนึ่ง ดังนั้นปัญหาที่ถูก random ลงไปที่หลังจึงมีขนาดแปลงเล็กกว่าและมีจำนวนซ้ำมากกว่าปัญหาแรก ตามหลักของ Plot Technique แสดงว่าปัญหาที่มีขนาดแปลงเล็กแต่ทำซ้ำบ่อยจะปรากฏผลแน่นอนกว่าปัญหาที่มีขนาดแปลงใหญ่กว่าแต่จำเป็นต้องทำน้อยซ้ำ ดังนั้นก่อนที่จะวางแผนเรื่องนี้ เราจำเป็นต้องพิจารณาเสียก่อนว่าในระหว่างสองปัญหาที่เราประสงค์จะทำการทดลองกันนั้น ปัญหาใดมีความสำคัญกว่ากัน หรือปัญหาใดที่เราหวังประโยชน์มากกว่า ให้เอาปัญหานั้นไว้ในแปลงย่อยที่สุด ดังจะได้แสดงให้เห็นต่อไปนี้

วิธีการวางแผนแบบ Split-plot ก็คล้าย ๆ กับการวางแผนแบบ Randomized Block นั่นเอง หากแต่เป็น Randomized Block Design หลาย ๆ อันซ้อนกันอยู่ ซึ่งจะได้แยกออกให้ดูในภายหลัง แต่หลักเกณฑ์ในการวางแผน เมื่อเปรียบเทียบกับ Randomized Block Design แล้วเป็นดังนี้คือ



จะเห็นว่า Randomized Block นั้น จาก Block ก็แบ่งเป็น plot
แต่ใน Split-plot นั้น จาก plot หรือ whole-plot ยังแบ่งเป็น plot ย่อย
หรือ sub-plot อีกทีหนึ่ง

สมมุติว่าเราประสงค์จะทำการทดลองเปรียบเทียบปุ๋ย ๔ อัตรา เพื่อต้องการจะ
ทราบว่าปุ๋ยอัตราไหนให้ผลดีที่สุด แต่เรามีพื้นที่ที่จะใช้ทดลองเปรียบเทียบปุ๋ยอยู่ถึง ๓ พันธุ์
ก็เลยอยากจะทำการศึกษาเปรียบเทียบพันธุ์ไปในตัวด้วยเลย ในที่นี้เราเห็นได้ว่าปัญหาเรื่องปุ๋ยสำคัญ
กว่าเรื่องพันธุ์ ดังนั้นในการวางแผนการทดลองจำเป็นจะต้องเอาปัญหาเรื่องปุ๋ยไว้ใน sub-plot
และเอาปัญหาเรื่องพันธุ์ไว้ใน whole-plot สมมุติว่าในการทดลองนี้เรากระทำ 4 Replications

(๑) ในการวางแผนขั้นแรก เราวางเช่นเดียวกับ Randomized Block Design
แบบธรรมดา โดยการสร้าง Block แล้วก็เอาปัญหาเรื่องพันธุ์วางลงใน whole-plot
เพื่อประกอบเป็น Block เสียก่อน เรามีพันธุ์อยู่สามพันธุ์ ดังนั้นใน 1 Block จึงมี 3
whole-plots ดังแผนต่อไปนี้

ให้ V_1 V_2 V_3 เป็นพันธุ์ทั้งสามพันธุ์ random ลงใน Block

แผนผังที่หนึ่ง

	V_2	V_1	V_3	V_1	V_3	V_2	
I							III
II							IV
	V_1	V_2	V_3	V_2	V_1	V_3	

(๒) การวางแผนผังชั้นที่สอง ให้พิจารณาว่าเรา มีปุ๋ยอยู่สี่อัตรา ดังนั้นจึงแบ่ง whole-plot แต่ละ whole-plot ออกเป็นสี่ส่วนเท่า ๆ กันก็จะเป็นสี่ sub-plot แล้วจึง random เอาปุ๋ยทั้งสี่อัตราลงไปใน sub - plot ภายใน whole - plot แต่ละ whole - plot ดังนี้.-

แผนผังที่สอง

	V_2	V_1	V_3	V_1	V_3	V_2	
Block I	T 3	T 2	T 4	T 2	T 1	T 3	Block III
	T 1	T 4	T 2	T 4	T 3	T 1	
	T 4	T 1	T 3	T 3	T 2	T 4	
	T 2	T 3	T 1	T 1	T 4	T 2	
Block II	T 1	T 4	T 2	T 4	T 3	T 1	Block IV
	T 3	T 1	T 3	T 2	T 4	T 3	
	T 2	T 3	T 1	T 3	T 1	T 2	
	T 4	T 2	T 4	T 1	T 2	T 4	
	V_1	V_2	V_3	V_2	V_1	V_3	

สรุปแล้ว Plot ในเรื่องของพันธุ์ (ตามแผนผังที่หนึ่ง) ก็คือ Block ในเรื่องของปุ๋ย (ตามแผนผังที่สอง) นั่นเอง จึงเห็นได้ว่าการทดลองเรื่องนี้ ถ้าเปรียบเทียบพันธุ์ก็ประกอบด้วย ๔ Block เท่านั้น แต่ถ้าเปรียบเทียบปรากฏว่ามี Block อยู่ทั้ง

๑๒ Block แสดงว่าปัญหาเรื่องปัญหามีความแน่นอนสูงกว่าปัญหาเรื่องพันธุ์ ทั้งที่ได้กล่าวไว้แล้ว จึงเป็นการเข้าหลักเกณฑ์อยู่แล้วที่เอาปัญหานี้ถือว่าเป็นเรื่องสำคัญกว่าไว้ใน sub - plot ซึ่งเป็นแปลงย่อยที่สุด

ต่อไปนี้จะขออธิบายแผนผังแบบ Split-plot เพื่อให้เข้าใจแจ่มแจ้งและลึกซึ้งลงไปอีก โดยการนำเอาแผนผังแบบนี้มาแยกออกเป็นแบบ Randomized Block ใหญ่ เพราะวิธีการคำนวณก็ลงใช้แบบ Randomized Block นั้นเอง หากแต่แผนผังแบบ Split-plot รวมเอาแผนผังแบบ Randomized Block ไว้ถึงสามอันซ้อนกัน ถ้าผู้คำนวณแยกไม่ได้ว่าจะไรเป็นของอันไหน ก็อาจทำให้การคำนวณไขว้เขวผิดพลาดได้ง่ายและยากแก่การเข้าใจ แต่ถ้ามองเห็นกระจ่างแจ่มว่าจะไรเป็นอะไรแล้ว ก็ทำให้้ง่ายมากขึ้นเพราะการคำนวณแบบ Randomized Block นั้นเรารู้เคยมาดีแล้วตั้งแต่ตอนต้น ๆ

สมมติว่าเราต้องการเปรียบเทียบวิธีการปลูกสองวิธี คือปลูกลึกกับปลูกตื้น (ให้เป็น T 1 และ T 2) และมีพันธุ์พืชอยู่สี่พันธุ์ด้วยกันคือ V₁ V₂ V₃ V₄ การวางแผนกระทำเป็นสามซ้ำ (3 Replications) เราถือว่าปัญหาเรื่องวิธีปลูกสำคัญกว่าจึงเอาไว้ใน sub-plot ต่อไปนี้จะแสดงให้เห็นว่าแผนผังรวมนั้นสามารถแยกออกเป็นแผนผังแบบ Randomized Block ธรรมดาได้ถึงสามอันดังนี้

แผนผังรวม

V ₂				V ₁				V ₃				V ₄			
T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T
2	1	1	2	1	2	1	2	1	2	1	2	1	2	1	2
Block I				Block II				Block III							

(๑) แยกเป็นแผนผังเปรียบเทียบพันธุ์

V_2	V_1	V_3	V_4	V_1	V_4	V_3	V_2	V_3	V_2	V_4	V_1
I				II				III			

(๒) แยกเป็นแผนผังเปรียบเทียบวิธีปฏิบัติ

T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	
2	1	1	2	1	2	2	1	2	1	1	2	1	2	1	2	1	2	1	2	1	2	1	
I		II		III		IV		V		VI		VII		VIII		IX		X		XI		XII	
Block																							

(๓) แยกเป็นแผนผังเปรียบเทียบความสัมพันธ์ระหว่างพันธุ์กับปี (Interaction)

T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T	T
2	1	1	2	1	2	2	1	2	1	2	1	1	2	1	2	1	2	2	1	2	1	2	1
V	V	V	V	V	V	V	V	V	V	V	V	V	V	V	V	V	V	V	V	V	V	V	V
2	2	1	1	3	3	4	4	1	1	4	4	3	3	2	2	3	3	2	2	4	4	1	1
I				II				III															

เราจะเห็นได้ว่าแผนผังแต่ละแผนผังเป็นแบบ Randomized Block Design
 ธรรมดาแน่นอน หากแต่ซ้อน ๆ กันอยู่ เมื่อเราต้องการวิเคราะห์เรื่องอะไรก็ยกหรือถอดเอา
 แผนผังนั้นออกมาได้ แต่ถ้าหากมีความจำแนกเพียงพอแล้วก็สามารถทำในแผนผังรวมได้เลย

เราจะเห็นได้ว่าในแผนผังเปรียบเทียบพันธุ์นั้น plot แต่ละ plot ของพันธุ์มีขนาด
ใหญ่ แต่พันธุ์แต่ละพันธุ์ก็มีเพียงสามซ้ำเท่านั้น เมื่อดูในแผนผังเปรียบเทียบวิธีปฏิบัติปรากฏว่า plot
เล็กลง แต่มีถึง ๑๒ ซ้ำ เพราะ plot ในเรื่องพันธุ์ได้กลายเป็น Block ในเรื่องวิธีปฏิบัติไป
แล้วที่กล่าวมาแล้วนเป็นการพิจารณาเพียงปัญหาเดียวโดยไม่คำนึงถึงปัญหาหนึ่ง

เมื่อเราพิจารณาสองปัญหาคงกันไปแล้วเราก็แบ่ง Interaction ออกได้เป็น
แปดอย่าง ฉะนั้นแม้แต่อย่างหนึ่งเหมือนกันแต่ก็อย่างหนึ่งผิดกันเช่น $V_1 T_1$ กับ $V_1 T_2$
เราก็ต้องถือเป็นคนละอย่าง เพราะจะต้องดูทั้งคู่ควบกันไป ตามแผนผังข้อสาม ดังนั้นในเรื่องนี้
จึงมีขนาด plot เล็กเท่า ๆ กับปัญหาเรื่องวิธีปฏิบัติในข้อสอง และยังมีน้อยซ้ำคือมีสามซ้ำเท่า ๆ
กับปัญหาเรื่องพันธุ์ในข้อหนึ่ง นอกจากนั้นเราจะเห็นได้ว่าใน Interaction นั้นวิธีการ Random
ยังไม่สมบูรณ์ เพราะพันธุ์แต่ละพันธุ์ติดกันไปเป็นคู่ ไม่ได้แยกออกเป็นอิสระ ด้วยเหตุผลข้อนี้เอง
จึงมีบางครั้งที่มีผู้ใช้วิธี random โดยตลอด คือเอาคู่ทั้ง ๘ คู่มา random โดยอิสระ ไม่มี
การแบ่งเป็น whole-plot หรือ sub-plot ทำแบบ Randomized Block สมบูรณ์
หากแต่ Interaction เข้ามาเกี่ยวข้อง ซึ่งเราเรียกวิธีการนี้ว่า Factorial Design
ซึ่งเรากล่าวในตอนหลัง เพราะถ้าเข้าใจ Split-plot แล้วก็จะทำให้เข้าใจ Factorial
Design ได้ง่ายเข้า

ต่อไปนี้เป็นตัวอย่างการคำนวณ Analysis of Variance แบบ Split-plot
สมมุติว่าจากการทดลองตามแผนผังที่แยกให้ดูนั้น เราเก็บตัวเลขผลผลิตได้ และนำมากรอก
ตารางดังต่อไปนี้

Coded by - 25

Block	Var.	Treatment				whole plot Total	Block Total
			X	Coded (X- $\bar{x}$)	X ²		
I	V ₂	T 2.	27	2	4	-4	1
		T 1.	19	-6	36		
	V ₁	T 1.	24	-1	1	-5	
		T 2.	21	-4	16		
	V ₃	T 1.	29	4	16	4	
		T 2.	25	0	0		
	V ₄	T 2.	30	5	25	6	
		T 1.	26	1	1		
II	V ₁	T 2.	21	-4	16	-3	-1
		T 1.	26	1	1		
	V ₄	T 2.	31	6	36	4	
		T 1.	23	-2	4		
	V ₃	T 1.	28	3	9	5	
		T 2.	27	2	4		
	V ₂	T 2.	25	0	0	-7	
		T 1.	18	-7	49		
III	V ₃	T 1.	27	2	4	4	-19
		T 2.	27	2	4		
	V ₂	T 2.	23	-2	4	-12	
		T 1.	15	-10	100		
	V ₄	T 2.	30	5	25	1	
		T 1.	21	-4	16		
	V ₁	T 2.	18	-7	49	-12	
		T 1.	20	-5	25		

$$\text{Grand total} = (1) + (-1) + (-19) = -19$$

$$\begin{aligned} \text{C.F.} &= \frac{G^2}{n} \\ &= \frac{-19^2}{24} \\ &= 15.04 \end{aligned}$$

$$\begin{aligned} \text{S.S. Total} &= X_1^2 + X_2^2 + X_3^2 + \dots + X_{24}^2 - \text{C.F.} \\ &= (4 + 36 + 1 + 16 + 0 + 25 + 1 + \dots + 49 + 25) - 15.04 \\ &= 445.00 - 15.04 = \underline{\underline{429.96}} \end{aligned}$$

$$\begin{aligned} \text{S.S. Whole-plot} &= \frac{(-4)^2 + (-5)^2 + (4)^2 + (6)^2 + (-3)^2 + \dots + (-12)^2}{n \text{ of sub-plot in each whole-plot}} - \text{C.F.} \\ &= \frac{497}{2} - 15.04 \\ &= 233.46 \end{aligned}$$

$$\begin{aligned} \text{S.S. varieties} &= \frac{v_1^2 + v_2^2 + v_3^2 + v_4^2}{\text{Block} \times n \text{ of sub-plot in each whole-plot}} - \text{C.F.} \\ &= \frac{(-5 + -3 + -12)^2 + (\dots)^2}{3 \times 2} - 15.04 \\ &= \frac{-20^2 + -23^2 + 13^2 + 11^2}{6} - 15.04 \\ &= \frac{1219}{6} - 15.04 \\ &= 188.13 \end{aligned}$$

$$\begin{aligned} \text{S.S. Block} &= \frac{B_1^2 + B_2^2 + B_3^2}{v \times \text{tr.}} - \text{C.F.} \\ &= \frac{1^2 + -1^2 + -19^2}{4 \times 2} - 15.04 \\ &= \frac{363}{8} - 15.04 \\ &= 30.34 \end{aligned}$$

$$\begin{aligned} \text{S.S. Error (a)} &= \text{S.S. Whole plot} - (\text{S.S. Block} + \text{S.S. Varieties}) \\ &= 233.46 - (188.13 + 30.34) \\ &= 14.99 \end{aligned}$$

$$\begin{aligned} \text{S.S. Treatments} &= \frac{\text{Tr}_1^2 + \text{Tr}_2^2}{b \times v} - \text{C.F.} \\ &= \frac{(-6 + -1 + \dots + -5)^2 (2 + -4 + \dots + -7)^2}{3 \times 4} - \text{C.F.} \\ &= \frac{(-24)^2 + (5)^2}{12} - 15.04 \\ &= 50.08 - 15.04 \\ &= 35.04 \end{aligned}$$

ในการหา S.S. ของความสัมพันธ์รวมกัน (Interaction) ระหว่างพันธุ์พืชพันธุ์ต่าง ๆ กับวิธีการปฏิบัติซึ่งเหมาะสมสำหรับแต่ละพันธุ์โดยเฉพาะนั้น เราจำเป็นต้องหาเสียก่อนว่าความสัมพันธ์รวมกันนั้นมีอยู่ชนิดใดกัน ในการนี้เราจะต้องดูทั้ง V และ T ประกอบกันไปคือ TV จะดูเฉพาะอย่างหนึ่งอย่างใดไม่ได้ เพื่อให้วิธีการรักษุณและสะดวกป้องกันการผิดพลาดจึงจำเป็นต้องทำการวางสองทางหรือสองด้าน ด้านหนึ่งเป็นพันธุ์และอีกด้านหนึ่งเป็นวิธีปฏิบัติ เมื่อนำมาพิจารณาสองด้านรวมกันคือคู่ตักกัน ลงไปตักกันตรงของใดก็เป็นเรื่องของช่องนั้นโดยเฉพาะ ตารางสองด้านนี้สะดวกมากในการป้องกันการผิดพลาด เราเรียกตารางนี้ว่า Two-way Table ดังจะได้แสดงให้เห็นต่อไปนี้

Two-way Table

Treatments				Varieties Total
Varieties		T_1	T_2	
V_1	$V_1 T_1$	-5	$V_1 T_2$ -15	-20
V_2	$V_2 T_1$	-23	$V_2 T_2$ 0	-23
V_3	$V_3 T_1$	9	$V_3 T_2$ 4	13
V_4	$V_4 T_1$	-5	$V_4 T_2$ 16	11
Treatment Total		-24	5	-19

ถ้าเราพิจารณาคูตารางข้างบนนี้โดยถือว่าเป็นตารางสำหรับวิเคราะห์ Analysis of Variance แบบ Randomized Block ธรรมดาแล้วจะเห็นได้ว่ามียอดหรือผลรวมอยู่สองด้านคือด้านตั้งและด้านนอน ผิดกันแต่เพียงว่าถ้าเป็นแบบธรรมดาแล้ว ด้านหนึ่งจะต้องเป็นผลรวมในเรื่องของ Block และอีกด้านหนึ่งก็เป็นผลรวมในเรื่องพันธุ์หรือ Treatments

ดังนั้นในการหาวิธีคำนวณของตารางก็กระทำแบบเดียวกันกับแบบธรรมดา คือขั้นแรกก็หา S.S.Total โดยเอาตัวเลขแต่ละจำนวนมากำลังสองแล้วรวมกันหมด แต่จะต้องหารด้วยจำนวน Block คือ ๓ เพราะที่แท้จริงนั้นเลขแต่ละจำนวนในตารางไม่ใช่เลขจำนวนเดียว แต่ละจำนวนมาจากเลขย่อย จำนวนในเรื่องใหญ่แล้วจึงหักลบด้วย C.F. ส่วนของพันธุ์และ Treatment เราได้หาไว้แล้ว

S.S.Total for Interaction

$$\begin{aligned}
 &= \frac{(v_1 T_1)^2 + (v_1 T_2)^2 + (v_2 T_1)^2 + \dots + (v_4 T_2)^2}{b} - C.F. \\
 &= \frac{(-5)^2 + (-15)^2 + (-23)^2 + \dots + (16)^2}{3} - 15.04 \\
 &= 385.67 - 15.04 \\
 &= 370.63
 \end{aligned}$$

$$S.S.Treatments = 35.04$$

$$S.S.Varieties = 188.13$$

$$\begin{aligned}
 &\underline{S.S. \text{ Interaction between treatment and varieties} = S.S. \text{ Total for} \\
 &\quad \text{Interaction} - (S.S.Treatment + S.S.Var.)} \\
 &= 370.63 - (35.04 + 188.13) \\
 &= \underline{147.46}
 \end{aligned}$$

$$\begin{aligned}
 D.F. \text{ Interaction} &= D.F.Total \text{ for Interaction} - (D.F.Treatment + D.F.Var.) \\
 &= 7 - (1 + 3) = 3
 \end{aligned}$$

ต่อไปเราก็กู้หา Error ของ Experiment ใหญ่ทั้งหมด เราก็บอก Error นี้ว่า Error (b) ทั้งนี้เนื่องจากถ้าเราสังเกตดูแล้วจะเห็นได้ว่า Experiment ใหญ่นี้ประกอบด้วย Experiment ย่อย ๆ อยู่ดังได้แยกแผนผังไว้ให้ดูโดยกระจ่างแล้วในหน้า สำหรับ Error (a) จึงเป็น Error ของ Experiment ย่อยเพื่อตรวจสอบความแตกต่างใน

เรื่องพันธุ์เท่านั้น แต่ Error (b) นั้นกลุ่มได้หาให้ของ Experiment ไว้ทั้งหมด ดังนั้นใน
การคำนวณหา Error (b) จึงจำเป็นต้องเอา S.S.Total ให้ตั้งแต่เอาออกของ
Experiment ย่อยคือ Total whole - plot หักออก แล้วจึงเอาเรื่องของ Treatment
และ Interaction หักออกอีกทีหนึ่ง

$$\begin{aligned} \text{S.S.Error(b)} &= \text{S.S.Total} - (\text{S.S.Total whole-plot} + \text{S.S.Treatment} + \text{S.S. Interaction}) \\ &= 429.96 - (233.46 + 35.04 + 147.46) \\ &= 14.00 \end{aligned}$$

$$\begin{aligned} \text{D.F.Error(b)} &= 23 - (11 + 1 + 3) \\ &= 8 \end{aligned}$$

Analysis of Variance

Sources of Difference	D.F.	S.S.	Variance	F ratio	
				Calculated	Table
<u>Total</u>	23	429.96			
<u>Whole-plot</u>					
Total	11	233.46			
Block	2	30.34			
Varieties	3	188.13	62.71	25.18**	
Error(a)	6	14.99	2.49		
<u>Sub-plot</u>					
Treatments	1	35.04	35.04	20.02**	
<u>Interaction</u>					
V.T.	3	147.46	49.15	28.08**	
<u>Error (b)</u>	8	14.00	1.75		

ในการคำนวณทุก ๆ ระยะเรากระทำเฉพาะภายในขอบเขตของแต่ละแผนผังการทดลอง ไม่เอาไปปะปนกัน ส่วนที่ขีดขอบเขตไว้ภายในขอบเขตนั้นก็เป็นการทดลองที่แบ่งย่อยออกไปโดยเฉพาะเรื่องพันธุ์เท่านั้น ดังนั้นการหาสำหรับเรื่องพันธุ์จึงใช้ Error (a) โดยตลอดนับตั้งแต่การหา Calculated F ratio และเมื่อปรากฏว่าพันธุ์แสดงความแตกต่างอย่างสำคัญยิ่ง เมื่อเรานำไปหา L.S.D. ของพันธุ์เพื่อต้องการทราบว่าพันธุ์ไหนให้ผลดีเร็วกว่ากันประการใด เราก็ต้องใช้ Variance ของ Error (a) แทนค่าในสูตร L.S.D. และเวลาเปิด t Table ก็ต้องใช้ D.F. ของ Error (a) แต่สำหรับสิ่งที่อยู่ใน Experiment 'ใหญ่' เช่น เรื่อง

Treatment Interaction

L.S.D. for Varieties

$$\text{Mean } V_1 = \frac{24 + 21 + 21 + 26 + 18 + 20}{3} = 43.33 / 1 \text{ whole plot}$$

$$\text{" } V_2 = \frac{27 + 19 + 25 + 18 + 23 + 15}{3} = 42.33 / 1 \text{ "}$$

$$\text{" } V_3 = \frac{29 + 25 + 28 + 27 + 27 + 27}{3} = 54.33 / 1 \text{ "}$$

$$\text{" } V_4 = \frac{30 + 26 + 31 + 23 + 30 + 21}{3} = 53.66 / 1 \text{ "}$$

$$\text{L.S.D.} = t \times \sqrt{\frac{2 \text{ Variance of Error (a)}}{n}}$$

$$= t \times \sqrt{\frac{2 \times 2.49}{3}}$$

$$= t \times 1.288$$

$$t \text{ value at D.F. } 6 = \begin{matrix} 2.477 & 5 \% \\ 3.707 & 1 \% \end{matrix}$$

$$\text{L.S.D. } 5 \% = 2.477 \times 1.288 = 3.19$$

$$\text{L.S.D. } 1 \% = 3.707 \times 1.288 = 4.775$$

$$V_3 = 54.33 \quad \text{difference} = 0.67 \text{ insignificant}$$

$$V_4 = 53.66$$

$$V_1 = 43.33 \quad \text{difference} = 10.33 \quad \text{** Highly significant}$$

$$V_2 = 42.33 \quad \text{"} = 1.00 \text{ insignificant}$$

L.S.D. for Treatments

$$\text{Mean of } T_1 = \frac{19+21+29+26+26+23+18+27+15+21+20+28}{12} = 23.00$$

$$\text{Mean of } T_2 = \frac{27+21+25+30+21+31+27+25+27+23+30+18}{12} = 25.417$$

$$\text{L.S.D.} = t \times \sqrt{\frac{2 \times \text{Variance of Error (b)}}{n}}$$

$$= t \times \sqrt{\frac{2 \times 1.75}{12}}$$

$$= t \times 0.54 \quad \begin{array}{l} t \text{ value at D.F. } 8 = 5 \% = 2.306 \\ \phantom{t \text{ value at D.F. } 8} = 1 \% = 3.355 \end{array}$$

$$\text{L.S.D. } 5 \% = 2.306 \times 0.54 = 1.245$$

$$\text{L.S.D. } 1 \% = 3.355 \times 0.54 = 1.812$$

$$\text{Difference between treatment} = 25.417 - 23.00$$

$$= 2.417^{**} \text{ Highly significant}$$

L.S.D. for Interaction

FACTORIAL DESIGN

เราจะเห็นได้ว่าในวิธีการวางแผนแบบ Split-plot Design นั้นปัญหาที่อยู่ใน Whole-plot ไม่สามารถจะ Random ให้กระจายออกไปทั่วทั้ง Block ได้ ปัญหาที่อยู่ใน sub-plot เท่านั้นที่สามารถกระจายออกไปได้ ดังนั้นใน Interaction คือ V.T. จึงทำให้ปัญหาทั้งสองนั้นได้รับผลงานแน่นอนไม่เท่ากัน ตัวอย่างเช่น V_1 นั้นจะต้องมี V_1T_1 และ V_1T_2 ซึ่งไม่สามารถแยกออกจากกันได้ แต่วิธีการ Split-plot นั้นมีข้อได้เปรียบอยู่ที่แค่เพียงช่วยให้การปฏิบัติสะดวกง่ายดายขึ้น เช่น V_1 ก็อยู่ร่วมกับ V_1 เมื่อแบ่งเป็นสอง Treatment ก็ทำให้สะดวก ไม่ต้องกังวลเรื่อง V เพราะ V แบ่ง

เป็นสัดส่วนอยู่แล้ว แต่ถ้าหากเราจะให้ทั้งสองปัญหามีความเสมอภาคกันก็กระทำได้ และจะยิ่งทำให้การคำนวณง่ายยิ่งขึ้น คือทำให้ V.T. เป็นเสมือน Treatment เดียวร่วมกัน ไม่มี Whole-plot และไม่มี Sub-plot และ V ทุก ๆ Combination จะถูก Random ลงไปใน Plot แบบ Randomized Block ธรรมดาเหมือนกันหมด แต่จำจะต้องระวังการสับสนปนเปในวิธีการปฏิบัติให้มาก เพราะทั้ง V และ T มีสิทธิได้รับวิธีการ Random เท่า ๆ กัน จึงต้องกังวลทั้งสองทาง วิธีการวางแผนแบบนี้เราเรียกว่า Factorial Design และที่ยกตัวอย่างมาแล้วนี้เป็นสองปัญหาระดับ 2 Factors เราอาจทำ 3, 4, 5.... Factors ได้ แต่ก็ยังต้องระวังความสับสนมากยิ่งขึ้น เช่นทดลองเปรียบเทียบพันธุ์และในขณะเดียวกันเปรียบเทียบชนิดของสิ่งปลูก และเปรียบเทียบปุ๋ยไปด้วยจึงเป็น 3 Factors

จากตัวอย่างเดิมในเรื่อง Split-plot Design ถ้าหากเราจะทำแบบ Factorial Design เราอาจกระทำได้ดังนี้

เรามีพันธุ์อยู่ 4 พันธุ์ คือ V_1, V_2, V_3 และ V_4

เรามีวิธีปฏิบัติอยู่ 2 แบบ คือ T_1, T_2

เราสามารถกระทำเป็น Combination ได้ 8 ชนิดด้วยกันคือ

1. V_1T_1
2. V_1T_2
3. V_2T_1
4. V_2T_2
5. V_3T_1
6. V_3T_2
7. V_4T_1
8. V_4T_2

ถ้าเราทำ 3 ขั้ว เราก็สร้างเป็น 3 Blocks แต่ละ Block แบ่งออกเป็น 8 plots

เสร็จแล้วก็ random เอาทั้ง 8 combination ลงไปที่ละ plot จนครบ 3 Blocks

แบบ Randomized Block ขรรรคกั กั้แฉฉฉัฉฉัฉฉั

2 V_{1T2} 21	7 V_{4T1} 26	1 V_{1T1} 24	4 V_{2T2} 27	6 V_{3T2} 25	3 V_{2T1} 19	8 V_{4T2} 30	5 V_{3T1} 29	Block I
6 V_{3T2} 27	5 V_{3T1} 28	2 V_{1T2} 21	1 V_{1T1} 26	7 V_{4T1} 23	4 V_{2T2} 25	3 V_{2T1} 18	8 V_{4T2} 31	Block II
7 V_{4T1} 21	8 V_{4T2} 30	3 V_{2T1} 15	2 V_{1T2} 18	4 V_{2T2} 23	1 V_{1T1} 20	5 V_{3T1} 27	6 V_{3T2} 27	Block III

เมื่อเราเก็บผลผลิตของแต่ละแปลงได้และนำมาชั่งน้ำหนัก ดังตัวเลขที่ปรากฏอยู่
 ด้านล่างของแปลงแต่ละแปลงนั้น เมื่อนำมาวิเคราะห์ความแตกต่างด้วยวิธีการทางสถิติ
 ชั้นแรกก็กระทำเช่นเดียวกันกับแบบ Randomized Block ขรรรคกั กั้แฉฉฉัฉฉัฉฉั

Coded by - 25

Combination	Block						Combination Total	Total V
	I		II		III			
	X	X ²	X	X ²	X	X ²		
V ₁ T ₁	24 -1	1	26 1	1	20 -5	25	-5	-20
V ₁ T ₂	21 -4	16	21 -4	16	18 - 7	49	-15	
V ₂ T ₁	19 -6	36	18 -7	49	15 -10	100	-23	-23
V ₂ T ₂	27 2	4	25 0	0	23 -2	4	0	
V ₃ T ₁	29 4	16	28 3	9	27 2	4	9	13
V ₃ T ₂	25 0	0	27 2	4	27 2	4	4	
V ₄ T ₁	26 1	1	23 -2	4	21 -4	16	-5	11
V ₄ T ₂	30 5	25	31 6	36	30 5	25	16	
Block Total	1		-1		-19		-19	

$$C.F. = \frac{(-19)^2}{24} = 15.04$$

$$S.S.Total = (-1)^2 + (1)^2 + (-5)^2 + (-4)^2 + \dots + (6)^2 + (5)^2 - C.F.$$

$$= 1 + 1 + 25 + 16 + 16 + 49 + \dots + 25 + 36 + 25 - 15.04$$

$$= 429.96 \text{ with } 24 - 1 = 23 \text{ D.F.}$$

$$\begin{aligned} \text{S.S. Block} &= \frac{(1)^2 + (-1)^2 + (-19)^2}{8} - \text{C.F.} \\ &= \frac{1 + 1 + 361}{8} - 15.04 \\ &= 30.34 \quad \text{with } 3-1 = 2 \text{ D.F.} \end{aligned}$$

$$\begin{aligned} \text{S.S. Combination} &= \frac{(-5)^2 + (-15)^2 + (-23)^2 + (0)^2 + (9)^2 + (4)^2 + (-5)^2 + (16)^2}{3} - \text{C.F.} \\ &= \frac{25 + 225 + 529 + 0 + 81 + 16 + 25 + 256}{3} - 15.04 \\ &= 370.63 \quad \text{with } 8-1 = 7 \text{ D.F.} \end{aligned}$$

$$\begin{aligned} \text{S.S. Error} &= \text{S.S. Total} - (\text{S.S. Block} + \text{S.S. Combinations}) \\ &= 429.96 - (30.34 + 370.63) \\ &= 28.99 \end{aligned}$$

$$\begin{aligned} \text{D.F. Error} &= \text{D.F. Total} - (\text{D.F. Blocks} + \text{D.F. Combinations}) \\ &= 23 - (2 + 7) \\ &= 14 \end{aligned}$$

เราจะเห็นได้ว่าขณะนี้เรามีตัวเลขที่จะกรอกลงในตาราง Analysis of Variance ได้โดยสมบูรณ์แล้ว

แต่ถ้าจะพิจารณากันให้ละเอียดจริงๆ แล้ว S.S. ของ Combination นั้น เป็นเรื่องของพันธุ์และวิธีปฏิบัติซึ่งเป็นสองประการรวมกันอยู่ แต่เราจะสามารถเปรียบเทียบได้ ว่าทั้ง ๘ Combination นั้น อันไหนจะให้ผลดีที่สุดได้ แต่ถ้าวัดการทราบเพียง ปัญหาหนึ่งปัญหาใดโดยเฉพาะเราก็กระทำไม่ได้ จึงจำเป็นต้องหา S.S. ของแต่ละเรื่องด้วย ดังนั้นถ้าจะให้สะดวกจึงควรทำ Two-way Table

ขั้นแรกเราหา S.S. ของ Varieties ก่อน โดยไม่คำนึงถึง Treatments ดังนั้นจึงเอายอดของแต่ละชนิดบวกกำลังสองแล้วรวมกันเสร็จแล้วหารเฉลี่ย แล้วจึงลบด้วย C.F.

$$\begin{aligned}
 \text{S.S. Varieties} &= \frac{V_1^2 + V_2^2 + V_3^2 + V_4^2}{b \times t} - \text{C.F.} \\
 &= \frac{(-20)^2 + (-23)^2 + (13)^2 + (11)^2}{3 \times 2} - 15.04 \\
 &= 188.13 \quad \text{with } 4-1 = 3 \text{ D.F.}
 \end{aligned}$$

ขั้นตอนต่อไปหา S.S. ของ Treatments โดยเฉพาะ โดยไม่คำนึงถึงพันธุ์ เพื่อต้องการเปรียบเทียบว่า Treatments ทั้งสองนั้นมีความแตกต่างกันจริงหรือไม่

$$\text{S.S. Treatments} = \frac{T_1^2}{v \times b} + \frac{T_2^2}{v \times b} - \text{C.F.}$$

หมายเหตุ ในการหา S.S. V และ S.S. T นี้ การที่เอาสามคูณตัวหารก็เพราะเหตุว่า ตัวเลขแต่ละจำนวนที่ปรากฏใน Two-way Table นั้น ไม่ใช่จำนวนย่อยที่สุดของ Experiment หากแต่ละตัวนั้นมาจากตัวเลขสามจำนวนใน Experiment

$$\begin{aligned}
 \text{S.S. Treatments} &= \frac{(-24)^2 + (5)^2}{4 \times 3} - 15.04 \\
 &= \frac{576 + 25}{12} - 15.04 \\
 &= 35.04
 \end{aligned}$$

S.S. Combinations ประกอบด้วย

- S.S. Varieties (V)
- S.S. Treatments (T)
- S.S. Interaction between V and T

$$\begin{aligned}
 \text{S.S. Interaction} &= \text{S.S. Combinations} - (\text{S.S. T.} + \text{S.S. V}) \\
 &= 370.63 - (35.04 + 188.13) \\
 &= 147.46
 \end{aligned}$$

Analysis of Variance

Sources of Difference	D.F.	S.S.	Variance	F ratio	
				Calculated	Table
Total	23	429.96			
Combination	Block	30.34	15.17	7.328	
	Treatments	35.04	35.04	16.927 **	4.60
	Varieties	188.13	62.71	30.294 **	8.86
	Interaction	147.46	49.15	23.744 **	3.34
Error	14	28.99	2.07		5.56

$$\begin{aligned} \text{D.F. Interaction} &= \text{D.F. Combination} - (\text{D.F. Treatments} + \text{D.F. Varieties}) \\ &= 7 - (1 + 3) = 3 \end{aligned}$$

จะเห็นได้ว่าเป็นเรื่องของ Randomized Block ขรรดกา ไม่มี Whole - plot หรือ Sub - plot ทุก ๆ รายการมีโอกาสที่จะ random โดยเสมอภาคกันหมด สำหรับ S.S. Combinations นั้นเราไม่นำมาใส่ไว้เพราะเหตุว่าเราแยกออกเป็นรายการย่อยสาม รายการคือ Treatments, Varieties และ Interaction แล้ว ในตารางวิเคราะห์ แบบ Split-plot เราจะเห็นได้ว่า S.S. Error (a) นี้ค่าต่ำทำให้มีความแน่นอนต่ำ (เพราะตัวหารต่ำลง) ปัญหาที่อยู่ใน whole-plot (คือเรื่องพันธุ์จึงมีความแน่นอนน้อยกว่าปัญหาที่อยู่ใน sub-plot แต่ใน Factorial Design นี้มี error อันเดียว ทุก รายการจึงให้ผลที่มีความแน่นอนเท่า ๆ กัน เนื่องจากค่าของ F ratio ของแต่ละรายการนั้น คำนวณได้จากการใช้ Variance ของแต่ละรายการนั้นหารด้วย Variance ของ Error อันเดียวกัน สำหรับการสรุปผลการค้นคว้าโดยใช้ L.S.D. ก็มีวิธีปฏิบัติในการคำนวณโดยใช้หลัก เกณฑ์เดียวกันกับที่ปฏิบัติมาแล้วในแบบ Split-plot

CHI - SQUARE (χ^2)

สิ่งใดก็ตาม ที่ใคร่รูไว้เป็นทฤษฎีแล้วอย่างชัดเจน เช่นทฤษฎีแห่งการผสมพันธุ์ สมมติว่าเราเอาคนไม่ที่มีคอกสีแดงผสมกับคนไม่ที่มีคอกสีขาว เราอาจใช้ทฤษฎีเป็นเครื่องทำนายได้ว่าลักษณะภายในของพ่อพันธุ์และแม่พันธุ์เป็นอย่างนั้นเป็นอย่างนั้นแล้ว ลูกผสมที่เกิดขึ้นจะมีลักษณะเป็นอย่างนั้นบางส่วน เป็นอย่างนี้บางส่วน และเป็นอย่างโน้นก็ส่วน ฯลฯ ถ้าได้ลูกผสมเท่านั้นตนเท่านั้น จะมีสีแดงก็คน สีขาวก็คน สีชมพูก็คน กลับไปก็คน กลับไปก็คน ฯลฯ ซึ่งการทำนายเหล่านี้ได้จากวิธีการศึกษานวทางทฤษฎี เมื่อปฏิบัติจริง ๆ แล้วผลที่ได้รับจะไม่ลงตัวตรงตามทฤษฎีพอดีทีเดียว จำเป็นจะต้องมีความคลาดเคลื่อนมาบ้างน้อยบ้าง ซึ่งอาจเกิดจากสิ่งแวดล้อมตามธรรมชาติได้น่า ๆ ประการ เช่นเดียวกันกับการโยนสสารกหนึ่งอัน ตามทฤษฎีกล่าวไว้ว่าจะออกหัวและก้อยอย่างละครึ่ง คือถ้าโยนร้อยครั้งจะต้องออกหัวห้าสิบครั้ง และก้อยห้าสิบครั้ง แต่ในทางปฏิบัติการโยนสสารกร้อยครั้งอาจออกหัวสี่สิบเจ็ดครั้ง ออกก้อยห้าสิบสามครั้ง หรือโยนอีกร้อยครั้งอาจออกหัวห้าสิบเก้าครั้ง ออกก้อยสี่สิบเอ็ดครั้ง และถ้าโยนซ้ำ ๆ กันชุดละร้อยครั้งหลาย ๆ ชุด แต่ละชุดแม้จะโยนร้อยครั้งเท่า ๆ กัน ก็จะมีออกหัวและก้อยไม่เท่ากัน การกระทบกระเทือนต่อสิ่งแวดล้อมตามธรรมชาตินั้นบางครั้งก็มาก บางครั้งก็น้อย เราจึงไม่สามารถตัดสินเอาตามข้อบ่งชี้ได้ว่าผลการปฏิบัติของเราเข้าใกล้หลักเกณฑ์หรือทฤษฎีใดหรือไม่ เพราะเราไม่สามารถจะทราบได้ว่าการปฏิบัตินั้น ๆ ได้รับความกระทบกระเทือนจากสิ่งแวดล้อมมากน้อยเพียงไร ดังนั้นจึงจำเป็นต้องใช้วิธีการ Chi-square นี่เป็นเครื่องพิสูจน์ว่าผลงานของเราเข้าใกล้ทฤษฎีไหนได้ และเมื่อทราบแล้วก็สามารถทำให้เราทราบลักษณะภายในของพ่อพันธุ์แม่พันธุ์ได้จากการพิสูจน์ย้อนกลับขึ้นไปโดยอาศัยทฤษฎีที่งานชิ้นนั้นขึ้นอยู่กับ

ดังนั้นวิธีการ Chi-square จึงนับว่าเป็นวิธีการแบบ Deduction คือมีทฤษฎีวางไว้แล้ว เพียงแต่พิสูจน์ผลงานให้เข้ากับทฤษฎีเท่านั้นเอง ผิดกับวิธีการก่อน ๆ ที่ได้ศึกษามา เพราะวิธีการเหล่านั้นเป็นแบบ Induction คือคนคิดของใหม่ขึ้นมาสร้างเป็นกฎเกณฑ์ เช่นการเปรียบเทียบพันธุ์ เปรียบเทียบบุุ่ เปรียบเทียบวิธีปลูกปฏิบัติเป็นต้น

วิธีการ Chi- square นี้ใช้มากในทางวิชาการผสมพันธุ์ เช่น ทฤษฎีการผสมพันธุ์กล่าวถึงผลของการผสมระหว่างพ่อแม่ที่มีลักษณะอย่างนั้นอย่างนั้นว่า เมื่อผสมกันแล้วลูกที่ออกมาจะมีลักษณะอย่างนั้นกี่ส่วนอย่างนั้นกี่ส่วน เราจึงจำเป็นต้องใช้วิธีการนี้ช่วยตรวจสอบผลงานของเราว่าเป็นไปตามทฤษฎีหรือกฎเกณฑ์ข้อใด โดยตรวจสอบว่า deviation ที่เกิดขึ้นระหว่างปรากฏการณ์จริง ๆ กับทฤษฎีนั้นจะสามารถถือเป็นความแตกต่างสำคัญได้หรือไม่ ถ้าปรากฏผลจากการวิเคราะห์ว่าไม่มี ความแตกต่างก็แสดงว่า ปรากฏการณ์นั้นเข้ากับทฤษฎีนั้นได้

ตัวอย่าง

ในการผสมพันธุ์พริก ลูกชั่วที่สอง (F_2) เราหวังว่าจะได้ลูกที่มีลักษณะต่าง ๆ แยกออกเป็นอัตราส่วน คือ ๑:๓:๘:๔ เมื่อเราได้ลูก ๑๑๐๐ ต้น เราจะต้องใช้วิธีการ Chi-square ตรวจสอบดังนี้

Phenotype	Expected		Observed	O-E	$\frac{(O-E)^2}{E}$
	ratio	number of plants(E)			
Deep purple	1	68.75	65	-3.75	0.2045
Medium purple	3	206.25	203	-3.25	0.0512
Light purple	8	550.00	563	13.00	0.3073
Green	4	275.00	269	-6.00	0.1309
		1100	1100		0.6939

$$\chi^2 = \sum \frac{(O-E)^2}{E}$$

$$\chi^2 = 0.6939$$

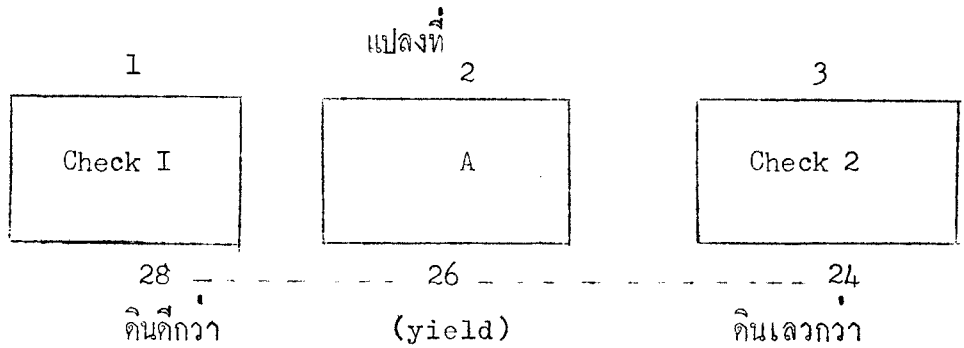
เราจะเห็นได้ว่า O-E นั้นคือตัวเลขที่แสดงความคลาดเคลื่อนซึ่ง ผลการปฏิบัติ ผิดไปจากทฤษฎีในที่นี้ D.F. = 3 ไปเปิด คูตาราง χ^2 ในหน้าหลังปรากฏว่าค่า χ^2 ที่เราคำนวณได้คือ ๐.๖๙๓๙ นี้อยู่ในระหว่างระดับ ๐.๔ ถึง ๐.๕ แสดงว่าอัตราส่วนที่ได้จากปฏิบัตินี้เข้ากับอัตราส่วน ๑:๓:๘:๔ นี้ได้

CHECK METHOD

วิธีการ Check Method นี้ ท่านศาสตราจารย์ H.H.Love นักสถิติผู้
เลื่องคนหนึ่งของโลก และเป็นอาจารย์สอนวิชาสถิติประจำมหาวิทยาลัย Cornell แห่ง
สหรัฐอเมริกาเป็นผู้คิดขึ้น ดังนั้นวิธีการนี้จึงนิยมเรียกกันว่า Love's Method หรือ
Cornell Method วิธีการนี้ในปัจจุบันไม่ค่อยมีผู้นิยมใช้กันนัก เพราะถือว่าเป็นวิธีการ
เก่าล้าสมัย แต่ข้าพเจ้ามีความเห็นว่าเราจะทิ้งวิธีการนี้เสียเสียไม่ได้ ในบางกรณียังมีความ
จำเป็นต้องใช้วิธีการนี้อยู่ โดยเฉพาะอย่างยิ่งในการคัดเลือก (Selection) และเปรียบเทียบ
พันธุ์พืชจำพวก Grain crop เช่นข้าวต่าง ๆ ถั่วต่าง ๆ ซึ่งมีจำนวนพันธุ์เป็นจำนวน
มาก ๆ

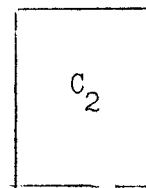
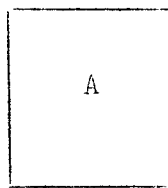
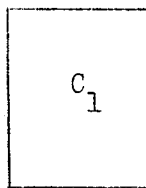
ในการวางแผนการคัดเลือกโดยใช้วิธีนี้ เราจำเป็นต้องมีพันธุ์พืชซึ่งใช้เป็นพันธุ์หลัก
(Check) อยู่พันธุ์หนึ่ง พันธุ์หลักหรือพันธุ์ Check นี้จะต้องเป็นพันธุ์ซึ่งได้ศึกษาและรู้จัก
ลักษณะอุปนิสัยดีมาแล้ว ยิ่งกว่านั้นพันธุ์ Check จะต้องเป็นพันธุ์ซึ่งในแถบท้องถิ่นนั้น ๆ นิยม
ปลูกพืชพันธุ์นั้นกันมาเป็นเวลานาน อาจมีพันธุ์อื่น ๆ ที่ได้รับการปรับปรุงขึ้นใหม่ ๆ และอาจมีพันธุ์
หนึ่งพันธุ์ใดหรือหลายพันธุ์ใหม่ดีกว่าพันธุ์ที่นิยมปลูกกันอยู่แต่ดั้งเดิม จึงนำเอาพันธุ์พืชหลาย ๆ พันธุ์
ที่ได้รับการปรับปรุงขึ้นใหม่นั้นมาทดลองปลูกเปรียบเทียบเทียบกับพันธุ์ Check ด้

หลักเกณฑ์ของวิธีการนี้คือ ใช้วิธีการปรับความอุดมสมบูรณ์ของดินให้พันธุ์ที่จะนำมา
เปรียบเทียบกับพันธุ์ Check นั้นได้รับความอุดมสมบูรณ์ของดินเท่า ๆ กับพันธุ์ Check ซึ่ง
คล้าย ๆ กับปลูกลงในแปลงเดียวกันและซ้ำที่กันพร้อม ๆ กัน ซึ่งในทางปฏิบัติย่อมเป็นไปได้
แต่ถ้าใช้วิธีการคำนวณช่วยปรับผลผลิตซึ่งเราเรียกวิธีการปรับผลผลิตนี้ว่า "Graded Method"
เมื่อปรับแล้วจึงนำเอาผลผลิตมาเปรียบเทียบกัน ดังจะขอยกตัวอย่างอธิบายต่อไปนี้ สมมุติว่าเรา
มีพันธุ์พืชอยู่พันธุ์หนึ่ง ให้เป็นพันธุ์ A เมื่อนำมาปลูกเปรียบเทียบกับพันธุ์ Check โดยเหตุที่ดิน
แต่ละแห่งย่อมมีความอุดมสมบูรณ์ไม่เท่ากัน ดังนั้นเราจึงทำแปลงปลูกพันธุ์ Check กระหนาบข้าง
เพื่อวัดความ สมบูรณ์ของดินว่าด้านไหนจะดีและเลวกว่ากัน โดยดูจากผลผลิตพันธุ์ Check
แล้วเราจึงมาคำนวณหาค่าปลูก Check ลงระหว่างกลางหรือตรงที่เดียวกับพันธุ์
พันธุ์ Check จะให้ผลเท่าใด



แสดงว่าคืนคานขายมือดีกว่าคานขามือ จึงคำนวณดูว่าถ้าปลูกพันธุ์ Check ในแปลงเดียวกับพันธุ์ A พันธุ์ Check จะให้ผลเท่าใด เรากำหนดให้ Check I มีอำนาจหนึ่งหน่วยเมื่ออยู่ในแปลงของตนเองคือแปลงที่หนึ่ง ดังนั้นถ้า Check 1 อยู่นั้นที่แปลงที่สอง อำนาจก็เหลือเพียงครึ่งเดียว และยิ่ง Check 1 ไปอยู่ในแปลงที่สามก็หมดอำนาจเลยคือเป็นศูนย์ ส่วน Check 2 ก็เช่นเดียวกัน เมื่ออยู่ในแปลงที่สามคือแปลงของตนเองก็มีอำนาจเป็นหนึ่ง เมื่อมาอยู่ในแปลงที่สองอำนาจของ Check 2 ก็เหลือเพียงครึ่งเดียว แต่ถ้าไปอยู่ในแปลงที่หนึ่ง Check 2 ก็หมดอำนาจคือเป็นศูนย์ เพราะแปลงที่หนึ่งนั้น Check 1 มีอำนาจเต็ม ถ้าเขียนเป็นแผนผังก็อาจเขียนได้ดังนี้ คือ

$$\begin{array}{ccccc}
 C_1 = 1 & \longrightarrow & C_1 = \frac{1}{2} & \longrightarrow & C_1 = 0 \\
 C_2 = 0 & \longleftarrow & C_2 = \frac{1}{2} & \longleftarrow & C_2 = 1
 \end{array}$$

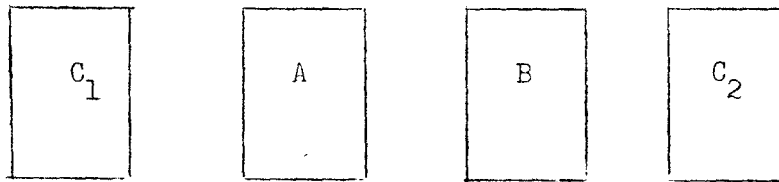


ดังนั้น Check ที่นำไปเทียบกับพันธุ์ A จะต้องเป็น

$$\frac{1}{2} C_1 + \frac{1}{2} C_2 = \frac{1}{2} \times 28 + \frac{1}{2} \times 24 = 26$$

ดังนั้นผลผลิตของ Check ที่จะเทียบกับผลผลิตของพันธุ์ A ได้ก็คือ ๒๖

ถ้าหากเราจะวางแผนเอาพันธุ์พืชสองพันธุ์คือ A, B สอดเข้าระหว่างพันธุ์ Check ก็กระทำได้



สูตรที่ใช้คำนวณ Check ก็คือเราจะต้องแบ่งแต่ละ Check ออกเป็นสามส่วน

$$\begin{array}{ccccccc}
 C_1 = 1 & \longrightarrow & C_1 = \frac{2}{3} & \longrightarrow & C_1 = \frac{1}{3} & \longrightarrow & C_1 = 0 \\
 C_2 = 0 & \longleftarrow & C_2 = \frac{1}{3} & \longleftarrow & C_2 = \frac{2}{3} & \longleftarrow & C_2 = 1
 \end{array}$$

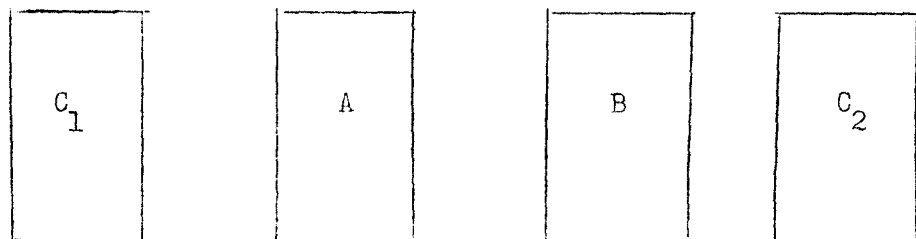


$$1C_1 + 0C_2$$

$$\frac{2}{3}C_1 + \frac{1}{3}C_2$$

$$\frac{1}{3}C_1 + \frac{2}{3}C_2$$

$$0C_1 + 1C_2$$



28

26.66

25.33

24

$$C_1 = 28, C_2 = 24$$

Check ที่จะใช้เทียบกับพื้นที่ A

$$= \frac{2C_1}{3} + \frac{1C_2}{3} = \frac{2}{3} \times 28 + \frac{1}{3} \times 24 = 26.66$$

Check ที่จะใช้เทียบกับพื้นที่ B

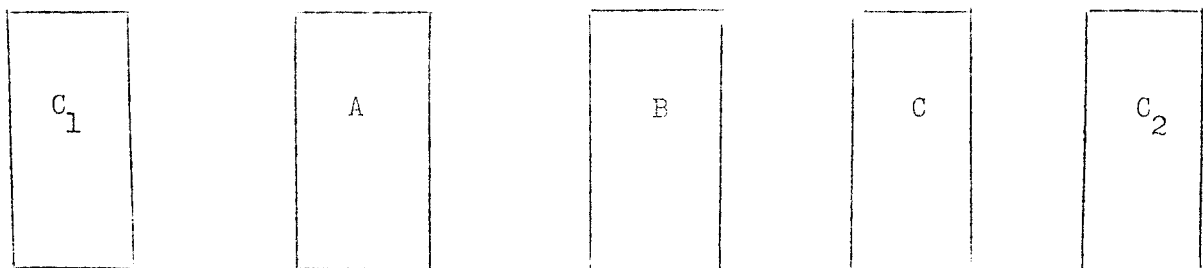
$$= \frac{1C_1}{3} + \frac{2C_2}{3} = \frac{1}{3} \times 28 + \frac{2}{3} \times 24 = 25.33$$

เราจะเห็นได้ว่า ผลลัพธ์ของพื้นที่ Check ที่ Grade แล้วเรียงลำดับดังนี้คือ

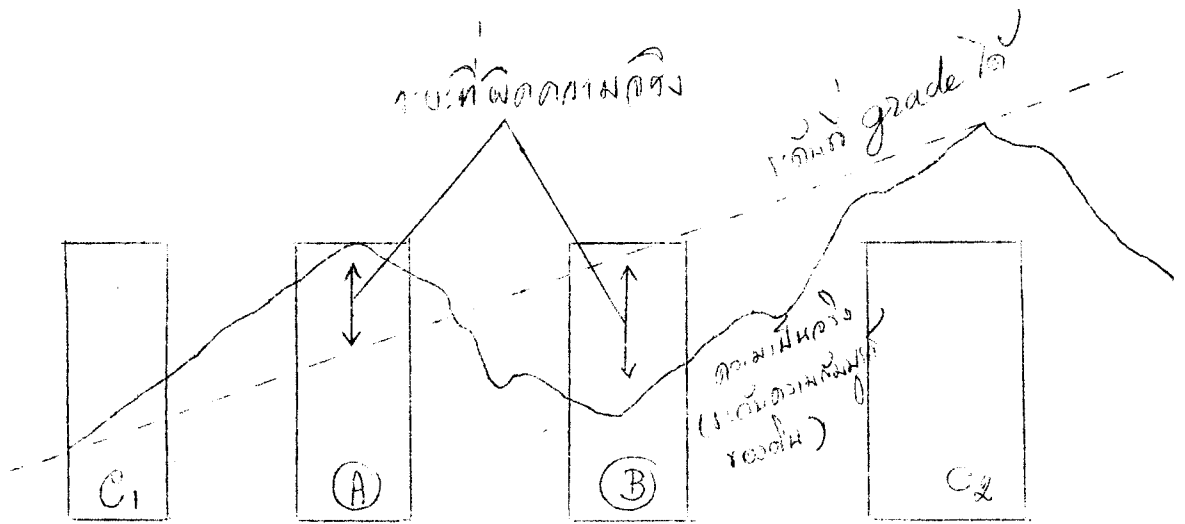
C_1	C for A	C for B	C_2
28	26.66	25.33	24

ซึ่งแสดงว่าดินเปลี่ยนแปลงจากที่แลวก้อย ๆ เลวกลงเรื่อย ๆ แต่ถ้าวเราเอาสามพื้นที่ สอดเข้าระหว่าง Check ก็จะต้องใช้สูตรนี้ คือ

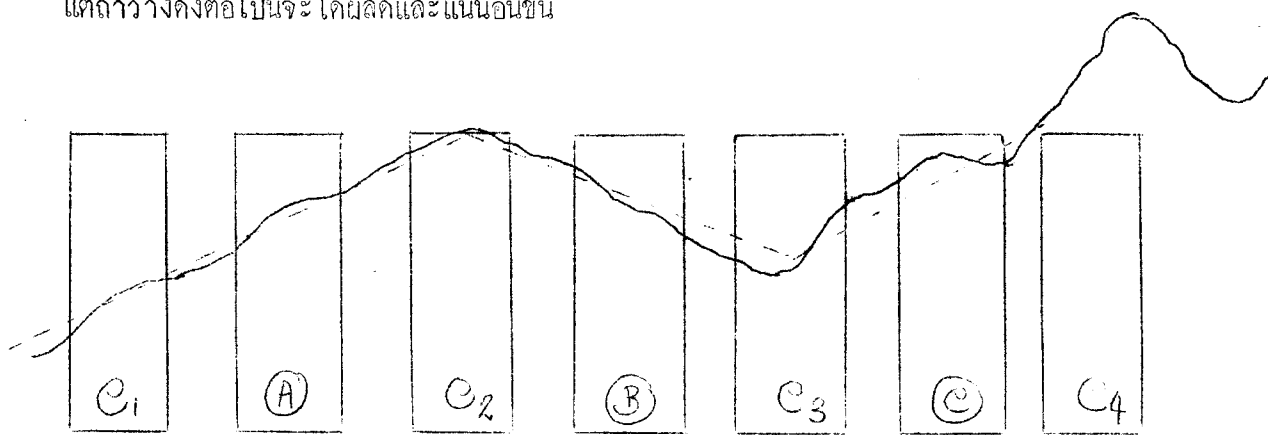
$$\begin{array}{ccccccc} C_1 = 1 & \longrightarrow & C_1 = \frac{3}{4} & \longrightarrow & C_1 = \frac{2}{4} & \longrightarrow & C_1 = \frac{1}{4} \longrightarrow C_1 = 0 \\ C_2 = 0 & \longleftarrow & C_2 = \frac{1}{4} & \longleftarrow & C_2 = \frac{2}{4} & \longleftarrow & C_2 = \frac{3}{4} \longleftarrow C_2 = 1 \end{array}$$



เราจะสอดจำนวนพื้นที่เข้าไประหว่าง Check มากน้อยเท่าใดนั้นขึ้นอยู่กับสิ่งที่จะต้องพิจารณาหลายประการด้วยกัน เช่นระยะการเปลี่ยนแปลงของดินในบริเวณแปลงทดลองนั้น ถ้าดินมีระยะการเปลี่ยนแปลงกระชั้นชิดมาก หากเราจะทิ้งระยะระหว่าง Check กว้างเกินไปก็อาจทำให้ผลการ Grade ผิดไปจากความเป็นจริง ซึ่งย่อมหมายถึงผลการทดลองผิดความเป็นจริงด้วย ดังนี้.—



แต่ถ้าวางดังต่อไปนี้จะได้ผลดีและแน่นอนขึ้น



จะเห็นได้ว่าแทนที่จะสอดสองพื้นที่เราแบ่งเป็นหนึ่งพื้นที่เท่านั้น แต่ได้ Check ใหญ่มากขึ้น

แต่ถ้าหากการเปลี่ยนแปลงของดินไม่กระชั้นชิด คือมีการเปลี่ยนแปลงทีละน้อย ๆ ในระยะกว้าง การได้ Check ก็เกินไปก็เป็นการยุ่งยากและหมดเปลือง ทำให้ผลงาน สะเอียดละออจนเกินความจำเป็น

นอกจากเหตุผลดังกล่าวแล้ว ขนาดแปลงก็อาจทำให้เกิดความกระทบกระเทือนได้เหมือนกัน ถ้าแปลงมีขนาดกว้างมากก็จะทำให้ช่องระหว่าง Check ทางมาก ทำให้ต้องลดจำนวนพันธุ์ที่จะสอดเข้าไปในระหว่าง Check ด้วย สิ่งดังกล่าวมาล้วนนี้เป็นประเด็นสำคัญซึ่งจะต้องนำมาพิจารณาในการวางระยะ Check ให้พอเหมาะกับสภาพการเปลี่ยนแปลงของดิน เพื่อให้งานได้ผลแน่นอนแต่เป็นไปในทางประหยัดและตัดความยุ่งยากสิ้นทั้งมวล

ในการวางแผนคันควาแบบนี้ โดยทั่ว ๆ ไปถ้าใช้กับพันธุ์พืชซึ่งทำแปลง (plot) เล็ก ๆ ปลูกประมาณสามแถวแล้ว มักใช้จำนวนพันธุ์ที่จะสอดเข้าไประหว่าง Check ประมาณ ๒ - ๕ พันธุ์ต่อหมู่ แต่โดยปกติมักใช้ ๓ พันธุ์ ถ้ามีทั้งหมด ๑๕ พันธุ์ ก็แบ่งออกเป็น ๖ หมู่ และทั้งหกหมู่ถือเป็นหนึ่งซ้ำหรือหนึ่ง Block

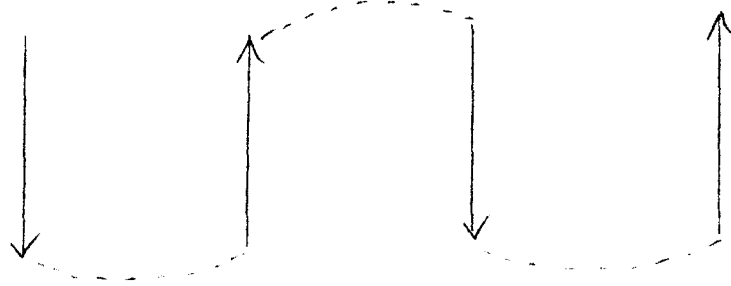
การกระทำซ้ำในการวางแผนแบบนี้มีผลกระทบประมาณสิบห้าปีว่าเหมาะเพราะเราต้องการความแน่นอนสูงมาก เนื่องจากแต่ละ Block ใช้วิธีเรียงลำดับพันธุ์ ไม่ได้มีการสลับ ซึ่งเป็นข้อเสียของวิธีการแบบนี้

การเรียงลำดับพันธุ์ดังกล่าวมาแล้วนั้นใช้วิธีเรียงไปทางเดียวกัน เมื่อครบหมดทุกพันธุ์แล้วก็ขึ้นพันธุ์ใหม่ เป็นซ้ำที่สอง หรือ Replication ที่สอง แต่ใช้วิธีเรียงวนกลับหรือสวนทางกันกับซ้ำก่อน มิใช่เรียงไปในแนวเดียวกัน ดังจะแสดงให้เห็นในแผนผังต่อไปนี้

ซ้ำที่ ๑	ซ้ำที่ ๒	ซ้ำที่ ๓	ซ้ำที่ ๔
๑	๘	๑	๘
๒	๗	๒	๗
๓	๖	๓	๖
๔	๕	๔	๕
๕	๔	๕	๔
๖	๓	๖	๓
๗	๒	๗	๒
๘	๑	๘	๑

๑๕๖

แสดงการเรียงลำดับพื้นที่ในแต่ละซ้ำ และแต่ละ Block



สมมติว่าจะใช้พื้นที่ Check สอดทุก ๆ สองพื้นที่ ดังนั้นแต่ละซ้ำจึงมีพื้นที่ Check ทั้งหมด ๓ แปลง ต่อไปนี้จะขอยกตัวอย่างการทดลองเปรียบเทียบพื้นที่ต่างไปทางประเทศจังหวัด เชียงใหม่ ซึ่งได้วางแผนการค้นคว้าและดำเนินการทดลองด้วยตนเองเมื่อปี พ.ศ.๒๕๕๑ ถึง พ.ศ.๒๕๕๒

Replications							
I	II	III	IV	V	VI	VII	VIII
C - 1	C - 5	C - 1	C - 5	C - 1	C - 5	C - 1	C - 5
1	8	1	8	1	8	1	8
2	7	2	7	2	7	2	7
C - 2	C - 4	C - 2	C - 4	C - 2	C - 4	C - 2	C - 4
3	6	3	6	3	6	3	6
4	5	4	5	4	5	4	5
C - 3	C - 3	C - 3	C - 3	C - 3	C - 3	C - 3	C - 3
5	4	5	4	5	4	5	4
6	3	6	3	6	3	6	3
C - 4	C - 2	C - 4	C - 2	C - 4	C - 2	C - 4	C - 2
7	2	7	2	7	2	7	2
8	1	8	1	8	1	8	1
C - 5	C - 1	C - 5	C - 1	C - 5	C - 1	C - 5	C - 1

๑๗๖

C = Check (King & Queen)

1 = Cole's Early

5 = Klondike

2 = Dixie Queen

6 = Klondike R 7

3 = Early Cansas

7 = Stone Mountain

4 = Florida Giant

8 = Tom Watson

Varieties	Replications								Total
	I	II	III	IV	V	VI	VII	VIII	
C-1	52.3	91.2	52.6	79.7	51.7	137.7	41.5	99.0	587.7
1	70.3	58.4	69.7	93.6	112.5	121.8	77.1	61.2	664.6
2	105.7	77.1	82.3	91.9	128.4	57.7	100.8	100.7	744.6
C-2	93.3	86.6	94.6	94.0	102.7	130.2	106.7	104.0	812.1
3	44.3	50.0	50.2	41.8	54.5	57.6	49.1	34.5	377.0
4	64.8	62.9	85.0	87.7	56.3	155.7	109.1	107.2	728.7
C-3	71.1	84.7	70.2	97.5	121.8	198.6	184.3	126.6	954.8
5	6.3	22.1	12.6	4.1	40.1	26.3	13.9	12.6	138.0
6	68.8	74.4	34.0	97.0	34.0	69.1	119.6	42.5	539.4
C-4	110.4	128.2	70.7	98.9	99.7	147.8	131.0	88.5	875.2
7	74.5	98.1	60.9	72.1	76.4	72.2	120.2	73.4	647.8
8	58.3	90.3	68.1	78.8	95.7	87.4	116.9	36.8	632.3
C-5	95.7	45.3	95.1	41.1	142.6	85.5	119.6	19.3	647.2

จากตารางที่แสดงผลผลิตไว้โดยละเอียดแล้วในหน้าที่แล้ว (หน้า ๑๗๕-๑๖)

เราก็ดำเนินการหาผลเฉลี่ย (M & S.E.M) ของแต่ละพันธุ์ด้วยวิธีธรรมดาทั้งที่ได้กล่าวมาแล้วในบทนี้ ๆ เสร็จแล้วจึงดำเนินการ Grade พันธุ์ Check สำหรับใช้ตรวจสอบเปรียบเทียบพันธุ์แต่ละพันธุ์ทั้ง ๘ พันธุ์นั้น โดยเหตุที่ในการวางแผนครั้งนี้เราเอาพันธุ์ Check สอดเข้าไปทุก ๆ สองพันธุ์ ดังนั้นสูตรที่จะใช้สำหรับ Grade พันธุ์ Check จึงจำเป็นต้องใช้สูตรดังนี้

เฉลย

$$\begin{array}{cccc}
 \frac{2}{3} C_1 & \frac{2}{3} C_1 & \frac{1}{3} C_1 & 0C_1 \\
 + & + & + & + \\
 0C_2 & \frac{1}{3} C_2 & \frac{2}{3} C_2 & 1C_2
 \end{array}$$

ตัวอย่างเช่นการคำนวณหาพันธ์ สำหรับเปรียบเทียบพันธ์ที่ ๑ เราก็คะทำดังนี้

$$\frac{2}{3} C_1 + \frac{1}{3} C_2 = \frac{2}{3} \times 73.46 + \frac{1}{3} \times 101.51 = 82.80$$

เมื่อกำหนดพันธ์ Check ได้ครบหมดทุกพันธ์แล้ว แต่พันธ์ Check ที่คำนวณได้นั้นยังไม่มี Standard error จึงจำเป็นต้องคำนวณหา S.E._M ของพันธ์ Check เพื่อใช้ในการเปรียบเทียบ คือเอาพันธ์ Check ที่ได้จากการทดลองจริง ๆ ทั้งหมดซึ่งมี S.E._M อยู่แล้วมารวมกันหมด และเอา S.E._M รวมกันด้วย เสร็จแล้วคิดคำนวณว่า S.E._M ของพันธ์ Check นั้น เป็นกี่เปอร์เซ็นต์ของผลผลิต ดังนี้

$$M \quad 73.46 + 101.51 + 119.35 + 109.40 + 80.90 = 484.62$$

$$S.E._M \quad 12.38 + 4.73 + 17.41 + 8.89 + 14.73 = 58.14$$

$$\text{Standard error of check in percent} = \frac{58.14 \times 100}{484.62} = 11.99 \%$$

เมื่อทราบว่า S.E._M เป็นกี่เปอร์เซ็นต์แล้วเราก็คำนวณหาได้ว่า Graded check แต่ละค่ามี S.E._M เป็น ๑๑.๙๙ % จะมีค่า S.E._M เป็นเท่าใด เสร็จแล้วจึงดำเนินการเปรียบเทียบกันเป็นคู่ ๆ ตามวิธีกรรมคาโดยใช้ "t" ratio เป็นมาตรฐาน ดังเช่นเดียวกันกับที่ได้อธิบายและยกตัวอย่างมาแล้ว

$$M = 100 \quad S.E._M = 11.99$$

$$M = 82.80 \quad S.E._M = \frac{11.99 \times 82.80}{100} = 9.99$$

$$M = 100 \quad S.E._M = 11.99$$

$$M = 92.15 \quad S.E._M = \frac{11.99 \times 92.15}{100} = 11.04$$

$$M = 100 \quad S.E._M = 11.99$$

$$M = 107.45 \quad S.E._M = \frac{11.99 \times 107.45}{100} = 12.88$$

$$M = 100 \quad S.E._M = 11.99$$

$$M = 113.39 \quad S.E._M = \frac{11.99 \times 113.39}{100} = 13.59$$

	Total	$M \pm S.E._M$	Calculated Check by graded method	$D \pm S.E._D$	Calculated "t"
C ₁	587.7	73.46 $\pm$ 12.38			
1	664.6	83.7 $\pm$ 8.39	82.80 $\pm$ 9.92	0.27 $\pm$ 12.98	0.020
2	744.6	93.07 $\pm$ 7.50	92.15 $\pm$ 11.04	0.92 $\pm$ 13.34	0.068
C ₂	812.1	101.51 $\pm$ 4.73			
3	377.0	47.12 $\pm$ 2.32	107.45 $\pm$ 12.88	-60.33 $\pm$ 13.08	4.612
4	728.7	91.09 $\pm$ 11.57	113.39 $\pm$ 13.59	-22.30 $\pm$ 17.84	1.250
C ₃	954.8	119.35 $\pm$ 17.41			
5	138.0	17.25 $\pm$ 4.17	116.02 $\pm$ 13.91	-98.77 $\pm$ 14.51	6.807
6	539.4	67.42 $\pm$ 10.78	112.71 $\pm$ 13.51	-45.29 $\pm$ 17.27	2.622
C ₄	875.2	109.40 $\pm$ 8.89			
7	647.8	80.97 $\pm$ 6.69	99.89 $\pm$ 11.97	-18.92 $\pm$ 13.71	1.380
8	632.3	79.04 $\pm$ 8.70	90.39 $\pm$ 10.83	-11.35 $\pm$ 13.89	0.817
C ₅	647.2	80.90 $\pm$ 14.73			

Standard error of check in percent = 11.99 %

ตารางที่แสดงอยู่ข้างบนนี้เป็นตารางต่อจากตารางในหน้า ๑๓๕ และคำอธิบายวิธีการคำนวณเพื่อให้ได้ตัวเลขซึ่งปรากฏอยู่ในตารางข้างบนนี้ก็ได้อธิบายไว้แล้วและได้ยกตัวอย่างไว้ด้วยในหน้า ๑๓๖ คือหน้าที่เราอ่านเอง และสิ่งที่ปรากฏในตารางข้างบนนี้ก็เป็นสิ่งที่ได้ทำการศึกษามาแล้วในบทนี้ ๆ ทั้งสิ้น เพียงแต่ว่าวิธีการวางแผนได้ถูกจัดแบ่งให้คิดแปลออกมาโดยมีเหตุผลซึ่งยอมรับกันทั่ว ๆ ไป แต่วิธีการคำนวณนั้นใช้ของเดิมทั้งสิ้น การเปรียบเทียบระหว่างพื้นที่นำมาทำการทดลองกับพื้นที่ Check ที่ Grade แล้วก็กระทำไปโดยใช้ระดับมาตรฐาน "t" เป็นเครื่องตัดสิน เราจะเห็นได้ว่าในช่องที่แสดงผลต่างของคู่เปรียบเทียบ ($D \pm S.E._D$) นั้น มีเครื่องหมายติดลบอยู่ด้วยในบางค่า แสดงว่าผลผลิตของพื้นที่นั้น ๆ ต่ำกว่าผลผลิตของ

๑๓๖
พันธ์ Check. เสียอีก แต่การที่เราจะสรุปได้ว่าพันธ์ไหนดีหรือเลวกว่าพันธ์ Check จริงหรือไม่
นั้นเราจำเป็นจะต้อง ดูตารางมาตรฐาน "t" เพื่อใช้เทียบกับ Calculated t ที่เรากำนวณ
ได้ในของสุดท้ายของตาราง แล้วเราก็จะสรุปผลได้

วิธีการวางแผนคนควาแบบนั้น เราจะเห็นได้ว่า ในกรณีที่เราปริมาณพันธ์พืชเป็น
จำนวนมาก ๆ เป็นเรือนพัน ๆ หรือหมื่น ๆ พันธุ์ เช่นในการคัดพันธุ์บริสุทธิ์ของข้าว ซึ่งเราจำเป็น
จะต้องไปเก็บรวงข้าวมาเป็นหมื่น ๆ รวงแล้วปลูกรวงคอแถว เพื่อทำการคัดสายพันธุ์ (Strain)
วิธีการนั้นน่าจะสะดวกมากเพราะเรามีโอกาสที่จะตรวจตราเปรียบเทียบในปีแรกได้โดยใช้สายตา
ธรรมดา ถ้าพันธ์หรือสายพันธุ์ใดเลวกว่า Check อย่างเห็นได้ชัดเจนนัยตาแล้ว เราก็คัดทิ้ง
ไปได้เลย ซึ่งเป็นการคัดเบื้องต้น เมื่อถึงปีที่ ๒ ที่ ๓ แล้ว พันธุ์ที่เหลือจากการคัดย่อมจะมี
คุณสมบัติใกล้เคียงกันมาก เราก็ดำเนินการเปรียบเทียบกันทางสถิติเพื่อความแน่นอนและกันการ
ผิดพลาดได้.

"r" TABLE
Correlation Coefficients at the 5 % and 1 % Levels
of Significance

Degree of Freedom	5 %	1 %
1	.997	1.000
2	.950	.990
3	.878	.959
4	.811	.917
5	.754	.874
6	.707	.834
7	.666	.798
8	.632	.765
9	.602	.735
10	.576	.708
11	.553	.684
12	.532	.661
13	.514	.641
14	.497	.623
15	.482	.606
16	.468	.590
17	.456	.575
18	.444	.561
19	.433	.549
20	.423	.537
21	.413	.526
22	.404	.515
23	.396	.505
24	.388	.496
25	.381	.487
26	.374	.478
27	.367	.470
28	.361	.463
29	.355	.456
30	.349	.449
35	.325	.418
40	.304	.393
45	.288	.372
50	.273	.354
60	.250	.325
70	.232	.302
80	.217	.283
90	.215	.267
100	.195	.254
125	.174	.228
200	.138	.181
300	.113	.148
400	.098	.128
500	.088	.115

1	2	3	4	5	6	8	12	24	∞	Values of t
---	---	---	---	---	---	---	----	----	---	-------------

Degrees of freedom for smaller mean square

13	4.67	3.80	3.41	3.18	3.02	2.92	2.77	2.60	2.42	2.21	2.160
	9.07	6.70	5.74	5.20	4.86	4.62	4.30	3.96	3.59	3.16	3.012
14	4.60	3.74	3.34	3.11	2.96	2.85	2.70	2.53	2.32	2.13	2.145
	8.86	6.51	5.56	5.03	4.69	4.46	4.14	3.80	3.43	3.00	2.977
15	4.54	3.68	3.29	3.06	2.90	2.79	2.64	2.48	2.29	2.07	2.131
	8.68	6.36	5.42	4.89	4.56	4.32	4.00	3.67	3.20	2.87	2.947
16	4.49	3.63	3.24	3.01	2.85	2.74	2.59	2.42	2.24	2.01	2.120
	8.53	6.23	5.29	4.77	4.44	4.20	3.89	3.55	3.18	2.75	2.921
17	4.45	3.59	3.20	2.96	2.81	2.70	2.55	2.38	2.19	1.96	2.110
	8.40	6.11	5.18	4.67	4.34	4.10	3.79	3.45	3.08	2.65	2.898
18	4.41	3.55	3.16	2.93	2.77	2.66	2.51	2.34	2.15	1.92	2.101
	8.28	6.01	5.09	4.58	4.25	4.01	3.71	3.37	3.01	2.57	2.878
19	4.38	3.52	3.13	2.90	2.74	2.63	2.48	2.31	2.11	1.88	2.093
	8.18	5.93	5.01	4.50	4.17	3.94	3.63	3.30	2.92	2.49	2.861
20	4.35	3.49	3.10	2.87	2.71	2.60	2.45	2.28	2.08	1.84	2.086
	8.10	5.85	4.94	4.43	4.10	3.87	3.56	3.23	2.86	2.42	2.845
21	4.32	3.47	3.07	2.84	2.68	2.57	2.42	2.25	2.05	1.81	2.080
	8.02	5.78	4.87	4.37	4.04	3.81	3.51	3.07	2.80	2.36	2.831
22	4.30	3.44	3.05	2.82	2.66	2.55	2.40	2.23	2.03	1.78	2.074
	7.94	5.72	4.82	4.31	3.99	3.75	3.45	3.12	2.75	2.30	2.819
23	4.28	3.42	3.03	2.80	2.64	2.53	2.38	2.20	2.00	1.76	2.069
	7.88	5.66	4.76	4.26	3.94	3.71	3.41	3.07	2.70	2.26	2.807
24	4.26	3.40	3.01	2.78	2.62	2.51	2.36	2.18	1.98	1.73	2.064
	7.82	5.61	4.72	4.22	3.90	3.67	3.36	3.03	2.66	2.21	2.797

4.00200,02118

40
59
190

Table 9.

Values of P (Top line) for Various Values of Chi-Square (Vertical Columns) and for Various Degree of Freedom (N')

The degrees of freedom one less than the number of classes

(From R.A. Fisher, "Statistical Methods for Research Workers")

N'	P = 0.99	0.98	0.95	0.90	0.80	0.70	0.50	0.30	0.20	0.10	0.05	0.02	0.01
1	0.00016	0.00063	0.0039	0.016	0.064	0.148	0.455	1.074	1.642	2.706	3.841	5.412	6.635
2	0.0201	0.0404	0.103	0.211	0.446	0.713	1.386	2.408	3.219	4.605	5.991	7.824	9.210
3	0.115	0.185	0.352	0.584	1.005	1.424	2.366	3.665	4.642	7.251	7.815	9.837	11.341
4	0.297	0.429	0.711	1.064	1.649	2.195	3.357	4.878	5.989	7.779	9.488	11.668	13.277
5	0.554	0.752	1.145	1.610	2.343	3.000	4.351	6.064	7.289	9.236	11.070	13.388	15.086
6	0.872	1.134	1.635	2.204	3.070	3.828	5.348	7.231	8.558	10.645	12.592	15.033	16.812
7	1.239	1.564	2.167	2.833	3.822	4.671	6.346	8.383	9.803	12.017	14.067	16.622	18.475
8	1.646	2.032	2.733	3.490	4.594	5.527	7.344	9.524	11.030	13.362	15.507	18.168	20.090
9	2.088	2.532	3.325	4.168	5.380	6.393	8.343	10.656	12.242	14.684	16.919	19.679	21.666
10	2.558	3.059	3.940	4.865	6.179	7.267	9.342	11.781	13.442	15.987	18.017	21.161	23.209

Table 9.

Degree of freedom for greater mean square

	1	2	3	4	5	6	8	12	24	10	Values of t
1	161.45 4052.10	199.50 4999.03	215.72 5420.349	224.57 5625.14	230.17 5764.08	233.97 5859.39	238.89 5981.34	243.91 6105.83	249.04 6234.16	254.32 6366.48	12.736 63.657
2	18.51 98.49	19.00 99.01	19.16 99.17	19.25 99.25	19.30 99.30	19.33 99.33	19.37 99.36	19.41 99.42	19.45 99.46	19.50 99.50	4.303 9.925
3	10.13 14.12	9.55 30.81	9.28 29.46	9.12 28.71	9.01 28.24	8.94 27.91	8.88 27.49	8.74 27.05	8.64 26.60	8.53 26.12	3.182 5.841
4	7.71 21.20	6.94 18.00	6.59 16.69	6.39 15.98	6.26 15.52	6.16 15.21	6.04 14.80	5.91 14.37	5.77 13.93	5.63 13.46	2.776 4.604
5	6.61 16.26	5.79 13.27	5.41 12.06	5.19 11.39	5.05 10.97	4.95 10.67	4.82 10.27	4.68 9.89	4.53 9.47	4.36 9.02	2.571 4.032
6	5.99 13.74	5.14 10.92	4.76 9.78	4.53 9.15	4.39 8.75	4.28 8.47	4.15 8.10	4.00 7.72	3.84 7.31	3.67 6.88	2.447 3.707
7	5.59 12.25	4.74 9.55	4.35 8.45	4.12 7.85	3.97 7.46	3.87 7.19	3.73 6.84	3.57 6.47	3.41 6.07	3.23 5.65	2.365 3.499
8	5.32 11.26	4.46 8.65	4.07 7.59	3.84 7.01	3.69 6.63	3.58 6.37	3.44 6.03	3.28 5.67	3.12 5.28	2.93 4.86	2.306 3.355
9	5.12 10.56	4.26 8.02	3.86 6.99	3.63 6.42	3.48 6.06	3.37 5.80	3.23 5.47	3.07 5.11	2.90 4.73	2.71 4.31	2.262 3.250
10	4.96 10.04	4.10 7.56	3.71 6.55	3.48 5.99	3.33 5.64	3.22 5.39	3.07 5.06	2.91 4.71	2.74 4.33	2.54 3.91	2.228 3.196
11	4.84 9.65	3.98 7.20	3.59 6.22	3.36 6.67	3.20 5.32	3.09 5.07	2.95 4.74	2.79 4.40	2.61 4.02	2.40 3.60	2.201 3.106
12	4.75 9.33	3.88 6.93	3.49 5.95	3.26 5.41	3.11 5.06	3.00 4.82	2.85 4.50	2.69 4.16	2.50 3.78	2.40 3.36	2.170 3.055

Degrees of freedom for smaller mean square

Degrees of freedom smaller mean square

	1	2	3	4	5	6	78	12	24	00	Values of t
25	4.24	3.38	2.99	2.76	2.60	2.49	2.34	2.16	1.96	1.71	2.060
	7.77	5.57	4.68	4.18	3.86	3.63	3.32	2.99	2.62	2.17	2.787
26	4.22	3.37	2.98	2.74	2.59	2.47	2.32	2.15	1.95	1.60	2.056
	7.72	5.53	4.64	4.14	3.82	3.59	3.29	2.96	2.58	2.13	2.779
27	4.21	3.35	2.96	2.73	2.57	2.46	2.30	2.13	1.93	1.67	2.052
	7.68	5.49	4.60	4.11	3.78	3.56	3.26	2.93	2.55	2.10	2.771
28	4.20	3.34	2.95	2.71	2.56	2.44	2.29	2.12	1.91	1.65	2.048
	7.64	5.45	4.57	4.07	3.75	3.63	3.23	2.90	2.52	2.06	2.763
29	4.18	3.33	2.98	2.70	2.54	2.43	2.28	2.10	1.90	1.64	2.045
	7.60	5.42	4.54	4.04	3.78	3.50	3.20	2.87	2.49	2.03	2.756
30	4.17	3.32	2.92	2.69	2.53	2.42	2.27	2.09	1.89	1.62	2.042
	7.56	5.39	4.51	4.62	3.78	3.47	3.17	2.84	2.47	2.01	2.750
35	4.12	3.26	2.87	2.64	2.48	2.37	2.22	2.04	1.83	1.57	2.030
	7.42	5.27	4.49	3.01	3.69	3.37	3.07	2.74	2.37	1.00	2.724
40	4.08	3.23	2.84	2.61	2.45	2.34	2.18	2.00	1.79	1.52	2.021
	7.31	5.18	4.31	3.83	3.51	3.29	2.99	2.66	2.29	1.82	2.704
45	4.06	3.21	2.81	2.58	2.42	2.31	2.15	1.97	1.76	1.48	2.014
	7.23	5.11	4.25	3.77	3.45	3.23	2.94	2.61	2.23	1.75	2.690
50	4.03	3.18	2.79	2.56	2.40	2.29	2.13	1.95	1.74	1.44	2.008
	7.17	5.06	4.20	3.72	3.41	3.19	2.89	2.56	2.18	1.68	2.678
60	4.00	3.15	2.76	2.52	2.37	2.25	2.10	1.92	1.70	1.39	2.000
	7.08	4.98	4.13	3.65	3.34	3.12	2.82	2.50	2.12	1.60	2.660
70	3.98	3.13	2.74	2.50	2.35	2.23	2.07	1.89	1.67	1.35	1.994
	7.01	4.92	4.07	3.60	3.29	3.07	2.78	2.45	2.07	1.53	2.648

4.00.

Degrees of freedom for smaller mean square

	1	2	3	4	5	6	8	12	24	∞	Values of t
80	3.96	3.11	2.72	2.49	2.33	2.21	2.06	1.88	1.65	1.31	1.990
	6.96	4.88	4.04	3.56	3.26	3.04	2.74	2.42	2.03	1.47	2.638
90	3.95	3.10	2.71	2.47	2.32	2.20	2.04	1.86	1.64	1.28	1.987
	6.92	4.85	4.01	3.53	3.23	3.01	2.72	2.39	2.00	1.43	2.632
100	3.94	3.09	2.70	2.46	2.30	2.19	2.03	1.85	1.53	1.26	1.984
	6.90	4.82	3.98	3.51	3.21	2.99	2.69	2.37	1.98	1.39	2.626
125	3.92	3.07	2.68	2.44	2.29	2.17	2.01	1.83	1.60	1.21	1.979
	6.84	4.78	3.94	3.47	3.17	2.95	2.66	2.33	1.94	1.32	2.616
150	3.90	3.06	2.66	2.43	2.27	2.16	2.00	1.82	1.59	1.18	1.976
	6.81	4.75	3.91	3.45	3.14	2.92	2.63	2.31	1.92	1.27	2.609
200	3.89	3.04	2.65	2.42	2.26	2.14	1.98	1.80	1.57	1.14	1.972
	6.76	4.71	3.88	3.41	3.11	2.89	2.60	2.28	1.88	1.21	2.601
300	3.87	3.03	2.64	2.41	2.25	2.13	1.97	1.79	1.55	1.10	1.968
	6.72	4.68	3.85	3.38	3.08	2.86	2.57	2.24	1.85	1.14	2.592
400	3.86	3.02	2.63	2.40	2.24	2.12	1.96	1.78	1.54	1.06	1.965
	6.70	4.66	3.83	3.37	3.06	2.85	2.56	2.23	1.84	1.11	2.588
500	3.86	3.01	2.62	2.39	2.23	2.11	1.96	1.77	1.54	1.06	1.965
	6.69	4.65	3.82	3.36	3.05	2.84	2.55	2.22	1.83	1.08	2.586
1000	3.85	3.00	2.61	2.38	2.22	2.10	1.95	1.76	1.53	1.03	1.965
	6.66	4.63	3.80	3.34	3.04	2.82	2.53	2.20	1.81	1.04	2.581
∞	3.84	2.99	2.60	2.37	2.21	2.09	1.94	1.75	1.52		1.960
	6.64	4.60	3.78	3.32	3.02	2.80	2.51	2.18	1.79		3.576

4:00200:02118