

บทที่ 4

ข้อมูลและการสำรวจภาคสนาม

ในบทที่ 4 นี้จะแบ่งออกเป็น 4 ส่วนหลักๆ โดยส่วนแรกจะกล่าวถึงวิธีการเก็บรวบรวมข้อมูลและการสำรวจภาคสนาม ส่วนที่สองจะนำเสนอผลจากการทำ Pre-survey และการกำหนดระดับราคา (Bid) ที่เหมาะสม ส่วนที่สาม คือ ผลจาก Final-survey และค่าสถิติเบื้องต้นของตัวแปรต้น และส่วนสุดท้าย คือ การเลือกแบบจำลองที่เหมาะสม

4.1 วิธีการเก็บรวบรวมข้อมูล และการสำรวจภาคสนาม

ในส่วนแรกวิธีการเก็บรวบรวมข้อมูลและการสำรวจภาคสนามแบ่งออกเป็น 4 หัวข้อย่อย ได้แก่ ประชากร ขนาดของตัวอย่าง การสุ่มตัวอย่างและการสำรวจภาคสนาม และการออกแบบสอบถาม ดังนี้

(1) ประชากร

ประชากร หมายถึง กลุ่มสมาชิกทั้งหมดของสิ่งต่างๆ ที่ต้องการศึกษา หรือต้องการสรุปผล ซึ่งประชากรในการศึกษานี้ คือ พ่อบ้านหรือแม่บ้าน ซึ่งหมายถึงผู้ที่เคยซื้อเนื้อหมูในซูเปอร์มาร์เก็ต มีหน้าที่รับผิดชอบในการเลือกซื้ออาหารสำหรับสมาชิกในครัวเรือนเป็นประจำ มีอายุตั้งแต่ 15 ปีขึ้นไป และอาศัยอยู่ในเขตกรุงเทพมหานคร

จากรายงานการสำรวจสำมะโนประชากรและเคหะในเขตกรุงเทพมหานครโดยสำนักงานสถิติแห่งชาติล่าสุดในปี 2543 พบว่า จำนวนครัวเรือนในกรุงเทพมหานครมีทั้งสิ้น 1,676,172 ครัวเรือน ดังนั้น เมื่อสมมติให้แต่ละครัวเรือนมีพ่อบ้านหรือแม่บ้านเพียง 1 คน จำนวนประชากรทั้งหมดจึงเท่ากับ 1,676,172 คน

(2) ขนาดของกลุ่มตัวอย่าง

กลุ่มตัวอย่าง หมายถึง กลุ่มสมาชิกของประชากรที่ถูกเลือกมาด้วยวิธีการต่างๆ เพื่อศึกษาวิเคราะห์ แล้วนำผลที่ได้ไปใช้อ้างอิงถึงประชากร

เนื่องจากในการศึกษานี้มีข้อจำกัดทั้งในเรื่องของเวลาและงบประมาณที่ใช้ในการศึกษา ดังนั้นการศึกษาถึงลักษณะของประชากรจำเป็นที่จะต้องศึกษาโดยใช้ตัวแทนของประชากรหรือกลุ่มตัวอย่าง ทั้งนี้ ตัวอย่างจะสามารถนำมาใช้เป็นตัวแทนของประชากรได้ดีจะต้องมีขนาดที่เหมาะสม ซึ่งในที่นี้ใช้สูตรการคำนวณหาขนาดของกลุ่มตัวอย่างตาม Yamane (1973) (อ้างถึงใน สุวรรณ, 2545) ดังนี้

$$n = \frac{N}{1 + N(e)^2}$$

โดยที่ n คือ ขนาดของกลุ่มตัวอย่างที่เหมาะสม

N คือ ขนาดของประชากร ในที่นี้เท่ากับ 1,676,172 คน

e คือ ค่าความคลาดเคลื่อนของกลุ่มตัวอย่าง ในการศึกษานี้กำหนดให้เท่ากับ 0.05¹

จากจำนวนประชากร 1,676,172 คน คำนวณโดยใช้สูตรข้างต้นจะได้เท่ากับ 399.90 ดังนั้น ขนาดของกลุ่มตัวอย่างที่จะใช้ในการศึกษานี้มีจำนวนทั้งสิ้น 400 คน

(3) วิธีการสุ่มตัวอย่างและการเก็บข้อมูลภาคสนาม

การสุ่มตัวอย่าง (Sampling) หมายถึง การเลือกตัวอย่างมาเป็นตัวแทนในการศึกษา โดยสมาชิกของกลุ่มตัวอย่างที่เลือกขึ้นมาจะมีโอกาสได้รับการเลือกเท่าๆกัน หรือถูกเลือกขึ้นมาโดยปราศจากความเอนเอียง (Unbiased) เพื่อให้ค่าสถิติที่คำนวณได้จากกลุ่มตัวอย่างจะมีค่า

¹ ในงานศึกษานี้กำหนดให้ค่าความคลาดเคลื่อนของกลุ่มตัวอย่างเท่ากับ 0.05 เนื่องจากจะทำให้ได้ขนาดของกลุ่มตัวอย่างที่เหมาะสมที่จะใช้ในการศึกษา เมื่อเทียบกับค่าความคลาดเคลื่อน 0.10 และ 0.01 ซึ่งจะทำให้ได้ขนาดของกลุ่มตัวอย่างเท่ากับ 99 และ 9,940 ตัวอย่างตามลำดับ

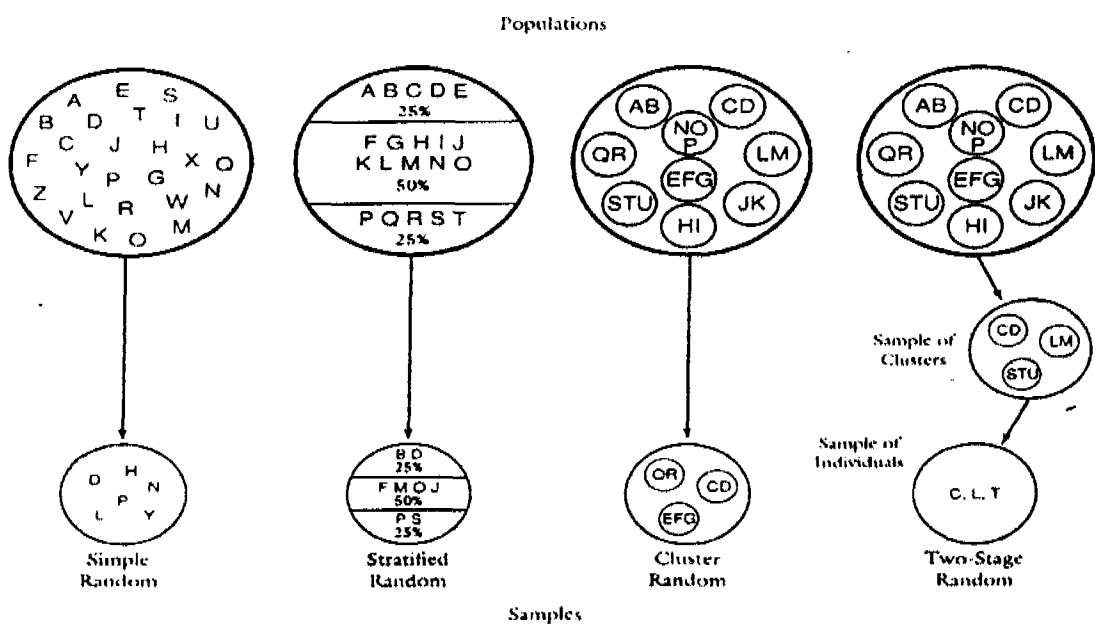
ใกล้เคียงกับค่าพารามิเตอร์ของประชากร ซึ่งกลุ่มตัวอย่างมีขนาดเล็กกว่าประชากร ดังนั้น จำเป็นจะต้องเลือกมาศึกษาเพียงบางส่วนของประชากร ในกรณีนี้จำนวนตัวอย่างเท่ากับ 400 คน จากประชากร 1,676,172 คน

การเลือกกลุ่มตัวอย่างเพื่อใช้ในการศึกษาแทนประชากรนั้น นอกจากจะต้องมีขนาดที่เหมาะสมแล้ว ตัวอย่างที่เลือกศึกษานั้นยังต้องสะท้อนถึงลักษณะของประชากรทั้งหมดได้ด้วย ซึ่งประชากรในกรณีนี้ คือ พ่อบ้านหรือแม่บ้านที่อาศัยอยู่ในกรุงเทพมหานคร ดังนั้นลักษณะที่แตกต่างกันที่สำคัญประการหนึ่งของประชากร คือ พื้นที่หรือเขตที่อยู่อาศัย โดยกรุงเทพมหานครแบ่งออกเป็น 50 เขต การเลือกตัวอย่างเฉพาะเขตใดเขตหนึ่งมาศึกษาอาจไม่สะท้อนถึงลักษณะของประชากรทั้งหมด

ดังนั้น ในการศึกษาจึงเลือกที่จะใช้วิธีการสุ่มตัวอย่างประเภทแบ่งกลุ่มหลายขั้นตอน (Multi-stage cluster sampling) ซึ่งเป็นการสุ่มตัวอย่างแบบอาศัยความน่าจะเป็น (Probability sampling) ประเภทหนึ่ง ซึ่งประเภทอื่นๆ ได้แก่ การสุ่มอย่างง่าย (Simple random sampling), การสุ่มอย่างเป็นระบบ (Systematic random sampling), การสุ่มแบบแบ่งชั้น (Stratified random sampling), และการสุ่มแบบหลายขั้นตอน (Multi-stage sampling) (ภาพที่ 4.1)

ภาพที่ 4.1

ประเภทของการสุ่มตัวอย่างแบบอาศัยความน่าจะเป็น (Probability sampling)



ที่มา: ศิริลักษณ์ (2538)

การสุ่มตัวอย่างทำได้โดยการแบ่งกรุงเทพมหานครออกเป็น 50 เขต แล้วสุ่มเขตอำเภอ 6 เขตมาเป็นกลุ่มตัวอย่าง โดยเลือกมาจากความมาก-น้อยของครัวเรือนในแต่ละเขต ซึ่งแบ่งออกเป็น 3 ขนาด² คือ จำนวนครัวเรือนมาก (17 เขต) จำนวนครัวเรือนปานกลาง (16 เขต) และจำนวนครัวเรือนน้อย (17 เขต) ซึ่งในงานศึกษานี้จะเลือกสุ่มเขตที่จะใช้เป็นกลุ่มตัวอย่างมา 2 เขต จากครัวเรือนในแต่ละขนาด หรือคิดเป็นอัตราส่วนประมาณ 1 ต่อ 8

ในการเลือกสุ่มเขตที่ใช้ในการศึกษาจะใช้วิธีการสุ่มตัวอย่างแบบสะดวกสบาย (Accidental sampling³) (โดยเลือกมา 2 เขต จาก 3 ขนาดครัวเรือน) ซึ่ง 6 เขต ที่จะใช้ในการศึกษา ได้แก่ เขตคลองเตย และเขตจตุจักร (เป็นตัวแทนของเขตที่มีจำนวนครัวเรือนมาก) เขตหลักสี่ และเขตปทุมวัน (เป็นตัวแทนของเขตที่มีจำนวนครัวเรือนปานกลาง) เขตพระนคร และเขตราชเทวี (เป็นตัวแทนของเขตที่มีจำนวนครัวเรือนน้อย) (ดูรายละเอียดได้ในภาคผนวก ค.)

สถานที่หลักที่จะไปเก็บข้อมูลภาคสนาม คือ ซุปเปอร์มาร์เก็ตในแต่ละเขต และเพื่อลดความเอนเอียง ซุปเปอร์มาร์เก็ตที่จะไปสอบถามนั้นไม่จำเป็นจะต้องเป็นซุปเปอร์มาร์เก็ต ชนิดเดียวกันในแต่ละเขต

ทั้งนี้ ช่วงเวลาที่ทำการสำรวจ Pre-survey คือ ช่วงเดือนพฤษภาคม-มิถุนายน 2548 ในขณะที่ช่วงเวลาที่สำรวจ Final-survey คือ ช่วงเดือนธันวาคม 2548-มกราคม 2549 โดยกลุ่มตัวอย่างที่จะเลือกไปออกแบบสอบถามในซุปเปอร์มาร์เก็ต จะเป็นลูกค้าที่นั่งอยู่ ณ ศูนย์อาหารของซุปเปอร์มาร์เก็ตที่เลือกศึกษา เนื่องจากคาดว่าจะในกลุ่มลูกค้าที่สะดวกในการตอบคำถามมากกว่าลูกค้าที่กำลังเดินเลือกซื้อสินค้า

² เขตที่มีครัวเรือนมาก คือ เขตที่มีจำนวนครัวเรือนตั้งแต่ 38,000 ครัวเรือนขึ้นไป เขตที่มีจำนวนครัวเรือนปานกลาง คือ เขตที่มีจำนวนครัวเรือน 28,400-38,000 ครัวเรือน และเขตที่มีจำนวนครัวเรือนน้อย คือ เขตที่มีจำนวนครัวเรือนน้อยกว่า 28,400 ครัวเรือน (ดูรายละเอียดได้ในภาคผนวก ค.)

³ การสุ่มตัวอย่างแบบสะดวกสบาย (Accidental sampling) เป็นรูปแบบหนึ่งของการสุ่มตัวอย่างแบบไม่อาศัยความน่าจะเป็น (Non-probability sampling) ซึ่งมีอีก 3 รูปแบบ ได้แก่ การสุ่มตัวอย่างแบบเจาะจง (Judgmental sampling), การสุ่มตัวอย่างแบบโควตา (Quota sampling) และการสุ่มตัวอย่างแบบลูกโซ่ (Snowball sampling) ซึ่งสามารถดูรายละเอียดได้ใน *ศิริลักษณ์ (2538)*

(4) การออกแบบสอบถาม

การออกแบบสอบถามในงานศึกษานี้แบ่งออกเป็น 2 รอบ คือ รอบแรก (Pre-survey) จะให้การถามแบบปลายเปิด (Open-ended) ซึ่งจะให้จำนวนตัวอย่างเท่ากับ 50 ตัวอย่าง ผลที่ได้จาก Pre-survey จะถูกนำไปใช้ในการออกแบบระดับราคา (Bid) ให้เหมาะสมเพื่อใช้ในการถามรอบที่สอง คือ Final-survey (จำนวน 400 ตัวอย่าง) ต่อไป

เนื้อหาในแบบสอบถามจะแบ่งออกเป็น 3 ส่วน คือ ส่วนแรกจะถามถึงข้อมูลเกี่ยวกับสภาพเศรษฐกิจและสังคมของผู้บริโภค (Socio-economic variable) ได้แก่ อายุ เพศ การศึกษารายได้ เป็นต้น ส่วนที่สองเป็นเรื่องเกี่ยวกับ Belief and perception variables เช่น ทศนคติต่อสินค้าเกษตรอินทรีย์ ความรู้ความเข้าใจเกี่ยวกับสินค้าเกษตรอินทรีย์ ส่วนสุดท้ายจะถามเรื่องความเต็มใจที่จะจ่าย ซึ่งในส่วนสุดท้ายใน Pre-survey จะเป็นคำถามปลายเปิด (Open-ended question) ซึ่งจะถามว่า “ท่านจะมีความเต็มใจที่จะจ่ายเพิ่มขึ้นสูงสุดกี่บาทสำหรับเนื้อหมูอินทรีย์ 1 กิโลกรัม? (เมื่อเทียบกับเนื้อหมูธรรมดา 1 กิโลกรัม)” แต่เมื่อเป็น Final-survey ส่วนสุดท้ายจะเป็น Double Bounded Close-ended question ซึ่งมีรายละเอียดดังต่อไปนี้

ในการหาค่าเฉลี่ยของความเต็มใจที่จะจ่ายส่วนต่างราคาสูงสุด (Mean of MWTP) ของผู้บริโภคใน Final-survey จะออกแบบสอบถามโดยใช้คำถามปลายปิด 2 รอบ (Double Bounded Close-ended) เช่น เมื่อเผชิญกับระดับราคาส่วนต่าง 20 บาทต่อกิโลกรัม⁴ ผู้บริโภคเต็มใจที่จะจ่ายหรือไม่ หากผู้ตอบเลือกเต็มใจที่จะจ่ายในการถามรอบที่ 2 จะเพิ่มระดับราคาขึ้น (เช่น 30 บาทต่อกิโลกรัม⁵) แต่หากผู้บริโภคไม่เต็มใจที่จะจ่ายในระดับราคาส่วนต่าง 20 บาทต่อกิโลกรัม

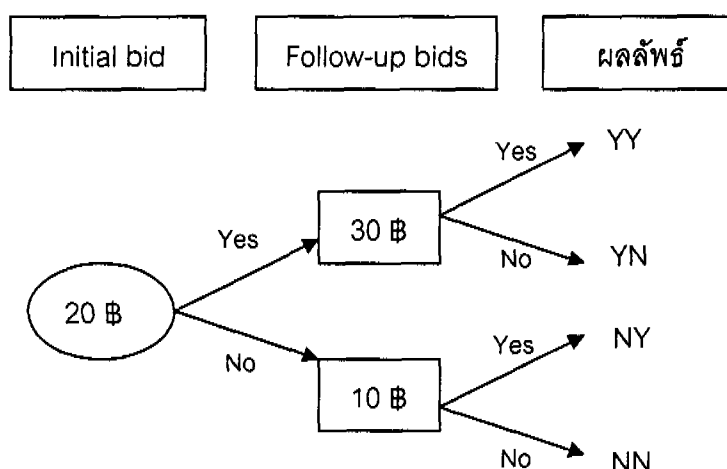
⁴ ราคาส่วนต่าง 20 บาทต่อกิโลกรัมในการถามรอบแรก (Initial bid) นี้เป็นราคาที่สมมติขึ้น ซึ่งในการศึกษาจริงนั้น ราคา Initial bid นี้จะเท่ากับค่ามัธยฐานของ Open-ended MWTP ที่ได้จากการทำ Pre-survey

⁵ ราคาส่วนต่าง 30 บาทต่อกิโลกรัมนี้ คือ Higher bid ซึ่งในการศึกษาจริงตามหลัก C-Optimal bid design นั้นค่า Higher bid จะได้จากเปอร์เซ็นไทล์ ที่ 75 ของ Open-ended MWTP ที่ได้จากการทำ Pre-survey

ในคำถามรอบแรก การถามรอบที่ 2 จะลดราคาลง (เช่น 10 บาทต่อกิโลกรัม⁶) ตัวอย่างรูปแบบการสอบถามจะเป็นดังภาพที่ 4.2

ภาพที่ 4.2

ตัวอย่างการออกแบบสอบถามแบบ Double Bounded Close-ended CVM



ที่มา: จากการรวบรวม

ระดับราคาส่วนต่าง 20 บาท (Initial bid: B_0), 30 บาท (Higher bid: B_H) และ 10 บาท (Lower bid: B_L) สามารถหาได้จากการทำ Pre-survey ดังที่ได้กล่าวไปแล้วในบทที่ 2 เรื่อง การกำหนดระดับราคา (Bid) ที่เหมาะสมว่า การออกแบบระดับราคาที่เหมาะสมสำหรับ Double Bounded Logit Model นั้น คือ การออกแบบให้ค่า Initial bid เท่ากับ Median WTP (เปอร์เซ็นต์ไทล์ที่ 50) และ Follow-up bid อยู่ที่ระดับเปอร์เซ็นต์ไทล์ที่ 25 สำหรับ B_L และเปอร์เซ็นต์ไทล์ที่ 75 สำหรับ B_H ตามหลัก C-Optimal design ซึ่งในการศึกษาครั้งนี้จะทำการทำ Pre-survey เท่ากับ 50 ตัวอย่าง และจะถามใน Final-survey เท่ากับ 400 ตัวอย่าง (ซึ่งสอดคล้องกับข้อเสนอของ Hanemann and Kanninen (1996) ที่เสนอแนะว่าจำนวนตัวอย่างในรอบ Final-survey ควรมากกว่า Pre-survey 2-3 เท่าตัว)

⁶ ราคาส่วนต่าง 10 บาทต่อกิโลกรัมนี้ คือ Lower bid ซึ่งในการศึกษาจริงตามหลัก C-Optimal bid design นั้น Lower bid จะได้จากเปอร์เซ็นต์ไทล์ที่ 25 ของ Open-ended MWTP ที่ได้จากการทำ Pre-survey

ทั้งนี้ ในงานศึกษาครั้งนี้ได้คำนึงถึงปัญหาในด้านความถูกต้องแม่นยำและความน่าเชื่อถือของการหาค่าความเต็มใจที่จะจ่ายโดยใช้วิธี CVM จึงได้พยายามออกแบบสอบถาม และใช้วิธีการศึกษาที่เหมาะสมเพื่อลดปัญหาดังกล่าวให้ได้มากที่สุด เช่น การใช้ Close-ended CVM เพื่อลดปัญหา Strategic bias รวมทั้งใช้ตารางประกอบการสอบถามเพื่อลดปัญหา Hypothetical bias เป็นต้น (ดูรายละเอียดได้ในบทที่ 2 หัวข้อ ความถูกต้องแม่นยำและความน่าเชื่อถือของวิธีการ CVM)

4.2 ผลจาก Pre-survey และการกำหนดระดับราคา (Bid) ที่เหมาะสม

ภายหลังจากการสำรวจภาคสนามรอบแรก (Pre-survey) โดยใช้วิธีการออกแบบสอบถามด้วยคำถามปลายเปิด (Open-ended) จำนวน 50 คน พบว่าข้อมูลทางสถิติเบื้องต้นของค่าความเต็มใจที่จะจ่ายราคาส่วนต่างสูงสุดระหว่างเนื้อสุกรอินทรีย์กับเนื้อสุกรธรรมดา 1 กิโลกรัม (MWTP) มีลักษณะดังนี้ (ตารางที่ 4.1)

ตารางที่ 4.1

ค่า Open-ended MWTP จาก Pre-survey และการกำหนดระดับราคาที่เหมาะสม

MWTP (บาทต่อกิโลกรัม)	ความถี่	ความถี่สะสม
3.5	1	1
5	3	4
10	17	21
15	3	24
20	11	35
25	1	36
30	9	45
40	3	48
50	2	50
Mean (บาทต่อกิโลกรัม)	19.37	
เปอร์เซ็นไทล์ที่ 25 (บาทต่อกิโลกรัม)	10 = Lower bid	

เปอร์เซ็นต์ไทล์ที่ 50 (บาทต่อกิโลกรัม)	20 = Initial bid
เปอร์เซ็นต์ไทล์ที่ 75 (บาทต่อกิโลกรัม)	30 = Upper bid

ที่มา: จากการสำรวจภาคสนามรอบ Pre-survey

จากตารางที่ 4.1 พบว่า ค่า MWTP ของผู้บริโภคมีค่าเฉลี่ย (Mean) เท่ากับ 19.37 บาทต่อกิโลกรัม ในขณะที่ค่ามัธยฐาน (Median) หรือเปอร์เซ็นต์ไทล์ที่ 50 มีค่าเท่ากับ 20 บาทต่อกิโลกรัม ในขณะที่เปอร์เซ็นต์ไทล์ที่ 25 และ 75 มีค่าเท่ากับ 10 และ 30 บาทต่อกิโลกรัมตามลำดับ ดังนั้นตามหลัก C-Optimal design (ดูรายละเอียดได้ในบทที่ 2 หัวข้อ เรื่องการออกแบบระดับราคาที่เหมาะสมสำหรับ Double Bounded Logit Model) จะได้ Initial bid เท่ากับ 20 บาทต่อกิโลกรัม รวมทั้ง Lower bid และ Higher bid เท่ากับ 10 และ 30 บาทต่อกิโลกรัม ตามลำดับ

4.3 ผลจาก Final-survey และค่าสถิติเบื้องต้นของตัวแปรต้น

ภายหลังจากได้ Initial, Lower และ Higher bid จากการทำ Pre-survey แล้ว ในการสำรวจภาคสนามรอบ Final-survey ได้ใช้วิธีการออกแบบสอบถามแบบ Close-ended CVM จำนวนทั้งสิ้น 400 ตัวอย่าง

ในการทำ Final-survey ใช้เวลาในช่วงเดือนธันวาคม 2548-มกราคม 2549 ช่วงเวลาประมาณ 10.00-18.00 น. ของทุกวัน ได้มีปัญหาและอุปสรรคบางประการในการออกภาคสนามตามห้างและซูเปอร์มาร์เก็ตบางแห่ง เนื่องจากทางห้างและซูเปอร์มาร์เก็ตบางแห่งไม่อนุญาตให้เข้าไปทำแบบสอบถามได้ ประกอบกับหากเป็นวันจันทร์-ศุกร์ที่ไม่ใช่เวลาพักเที่ยงจะมีกลุ่มตัวอย่างค่อนข้างน้อยมาก (แม้ช่วงเวลาพักเที่ยงจะมีกลุ่มตัวอย่างมาก แต่อัตราการตอบแบบสอบถามค่อนข้างน้อย เนื่องจากมีระยะเวลาในการตอบค่อนข้างน้อย ผู้ตอบไม่ค่อยจะสะดวกที่จะตอบแบบสอบถามมากนัก เนื่องจากต้องรับประทานอาหารกลางวันและรีบกลับไปทำงาน) ดังนั้น ช่วงเวลาที่เหมาะสมสำหรับการออกภาคสนามที่ห้างหรือซูเปอร์มาร์เก็ต คือ วันธรรมดา ช่วงเย็น หรือช่วงวันหยุดเสาร์-อาทิตย์

ดังนั้น เนื่องจากข้อจำกัดและระยะเวลาที่ใช้ในการศึกษา ในงานศึกษานี้จึงได้ทำ Final-survey ในสถานที่อื่นๆ ที่ไม่ใช่ห้าง หรือซูเปอร์มาร์เก็ตด้วย เช่น สวนสาธารณะ (สวนจตุจักร) พนักงานมหาวิทยาลัย (มหาวิทยาลัยธรรมศาสตร์ และมหาวิทยาลัยธุรกิจบัณฑิต) พนักงานบริษัทหรือองค์กรที่มีคนรู้จักของผู้ทำการศึกษาอยู่ด้วย (ตาราง 4.2)

ตารางที่ 4.2
เขตและสถานที่ที่ไปสู่มตัวอย่าง

เขต - สถานที่ที่ไปสู่มตัวอย่าง	จำนวนตัวอย่าง	สัดส่วน
เขตคลองเตย – โลตัส พระราม 4	43	10.75%
เขตจตุจักร – เซ็นทรัลลาดพร้าว, สวนจตุจักร	74	18.50%
เขตหลักสี่ - มหาวิทยาลัยธุรกิจบัณฑิต	65	16.25%
เขตปทุมวัน – โลตัส พระราม 1	68	17.00%
เขตพระนคร - มหาวิทยาลัยธรรมศาสตร์	66	16.50%
อื่นๆ	84	21.00%
รวม	400	100.00%

ที่มา: จากการสำรวจภาคสนามรอบ Final-survey

เมื่อพิจารณาจากค่าเฉลี่ยของตัวอย่างทั้ง 400 ตัวอย่าง พบว่า Socio-economic variables มีลักษณะดังนี้ คือ มีอายุเฉลี่ยประมาณ 32 ปี ส่วนใหญ่เป็นเพศหญิง มีการศึกษาโดยเฉลี่ยที่ระดับปริญญาตรี มีรายได้เฉลี่ยประมาณ 15,001-25,000 บาทต่อเดือน มีรายได้ของครอบครัวโดยเฉลี่ย 45,001-55,000 บาทต่อเดือน มีจำนวนสมาชิกในครอบครัวประมาณ 4 คน

ในส่วนของตัวแปร Beliefs and perceptions variables พบว่า ตัวอย่างประมาณครึ่งหนึ่งจะกังวลในการบริโภคสารเคมีตกค้าง (PEST) (มีค่าเฉลี่ยเท่ากับ 54%) ส่วนใหญ่มีความกังวลเรื่องสุขภาพ (HEAL) (มีค่าเฉลี่ยเท่ากับ 75%) ส่วนใหญ่คิดว่าตนเองมีความรู้ความเข้าใจในสินค้าเกษตรอินทรีย์น้อย (KNOW) (มีค่าเฉลี่ยเท่ากับ 95%) ส่วนใหญ่มีความเห็นว่าข้อมูลข่าวสารเกี่ยวกับสินค้าเกษตรอินทรีย์ในปัจจุบันไม่เพียงพอ (INFO) (มีค่าเฉลี่ยเท่ากับ 92%) ส่วนใหญ่ไม่เชื่อมั่นในระบบสินค้าเกษตรและอาหารของไทยในปัจจุบัน (BELIEF) (มีค่าเฉลี่ยเท่ากับ 80%) ส่วนใหญ่ไม่เคยหรือซื้อสินค้าอาหารปลอดภัยเพียงบางครั้ง (มีค่าเฉลี่ยเท่ากับ 71%) ส่วนใหญ่มีความรู้ความเข้าใจในสินค้าเกษตรอินทรีย์ปานกลาง (KNOW2) (มีค่าเฉลี่ยเท่ากับ 7.41) และมีทัศนคติต่อสินค้าเกษตรอินทรีย์ค่อนข้างดี (ATT) (มีค่าเฉลี่ยเท่ากับ 24.85) (ตารางที่ 4.3 และ 4.4)

ตารางที่ 4.3
ความถี่และสัดส่วนของตัวแปรต้น

ตัวแปรและความหมายของตัวแปร	ความถี่	สัดส่วน (%)
<i>AGE</i> = อายุ		
- น้อยกว่า 20 ปี	12	3.00
- 20 – 40 ปี	294	73.50
- มากกว่า 40 ปี	94	23.50
<i>SEX</i> = เพศ		
- เพศชาย = 0	59	14.75
- เพศหญิง = 1	341	85.25
<i>EDU</i> = ระดับการศึกษาสูงสุด		
- ต่ำกว่าหรือเท่ากับมัธยมศึกษาตอนต้น	26	6.50
- มัธยมศึกษาตอนปลาย-ปริญญาตรี	324	81.00
- สูงกว่าปริญญาตรี	50	12.50
<i>INC</i> = รายได้ต่อเดือน		
- น้อยกว่าหรือเท่ากับ 25,000 บาท	312	78.00
- 25,001-55,000 บาท	75	18.75
- มากกว่า 55,000 บาท	13	3.25
<i>INC2</i> = รายได้ครอบครัวต่อเดือน		
- น้อยกว่าหรือเท่ากับ 45,000 บาท	256	64.00
- 45,001-75,000 บาท	70	17.50
- มากกว่า 75,000 บาท	74	18.50
<i>FAM</i> = จำนวนสมาชิกในครอบครัว (รวมผู้ตอบแบบสอบถาม)		
- มีสมาชิกในครอบครัว 1-3 คน	149	37.25
- มีสมาชิกในครอบครัว 4-7 คน	239	59.75
- มีสมาชิกในครอบครัวมากกว่า 7 คน	12	3.00
<i>PEST</i> = ความกังวลในเรื่องสารเคมีตกค้างจากการบริโภคอาหาร (Pesticide Concern)		
- ไม่กังวล = 0 (ไม่ตอบ "สารเคมีตกค้าง" เป็นอันดับที่ 1)	186	46.50
- กังวล = 1 (ตอบ "สารเคมีตกค้าง" เป็นอันดับที่ 1)	214	53.50
<i>HEAL</i> = ความกังวลเกี่ยวกับสุขภาพจากการบริโภคอาหาร (Health Concern)		
- ไม่กังวล = 0 (ไม่ตอบ "ความสะอาดถูกหลักอนามัย" เป็นอันดับที่ 1)	102	25.50
- กังวล = 1 (ตอบ "ความสะอาดถูกหลักอนามัย" เป็นอันดับที่ 1)	298	74.50

ตัวแปรและความหมายของตัวแปร	ความถี่	สัดส่วน (%)
<i>KNOW</i> = ความรู้ความเข้าใจเกี่ยวกับสินค้าเกษตรอินทรีย์ที่ผู้บริโภคประเมินตนเอง (Self-reported knowledge)		
- รู้ค่อนข้างน้อย หรือไม่รู้เลย = 0	382	95.50
- รู้ค่อนข้างมาก หรือรู้มาก = 1	18	4.50
<i>INFO</i> = ความเพียงพอของข้อมูลข่าวสารเกี่ยวกับสินค้าเกษตรอินทรีย์ในปัจจุบัน		
- ไม่เพียงพอ = 0	369	92.25
- เพียงพอ = 1	31	7.75
<i>BELIEF</i> = ความเชื่อมั่นต่อระบบสินค้าเกษตรและอาหารของไทยในปัจจุบัน (เช่น การจัดการฟาร์ม มาตรฐาน การตรวจสอบรับรอง ฯลฯ)		
- เชื่อมั่นเล็กน้อย หรือไม่เชื่อมั่นเลย = 0	320	80.00
- เชื่อมั่นค่อนข้างมาก หรือเชื่อมั่นมาก = 1	80	20.00
<i>BUY</i> = ความบ่อยครั้งในการซื้อสินค้าอาหารปลอดภัย (เช่น ผักปลอดสารพิษ หรือเนื้อสุกรอนามัย ฯลฯ)		
- ซื้อบ้างบางครั้ง หรือไม่เคยซื้อเลย = 0	285	71.25
- ซื้อค่อนข้างบ่อย หรือซื้อเป็นประจำ = 1	115	28.75
<i>KNOW2</i> = ความรู้ความเข้าใจเกี่ยวกับสินค้าเกษตรอินทรีย์ที่ได้จากการประเมิน		
- มีความรู้ความเข้าใจน้อย (มีคะแนนน้อยกว่า 5 คะแนน)	12	3.00
- มีความรู้ความเข้าใจปานกลาง (มีคะแนนอยู่ระหว่าง 5-9 คะแนน)	357	89.25
- มีความรู้ความเข้าใจมาก (มีคะแนนมากกว่า 9 คะแนน)	31	7.75
<i>ATT</i> = หัศนคติต่อสินค้าเกษตรอินทรีย์		
- มีทัศนคติไม่ดี (มีคะแนนน้อยกว่า 18 คะแนน)	14	3.50
- มีทัศนคติปานกลาง (มีคะแนนอยู่ระหว่าง 18-27 คะแนน)	282	70.50
- มีทัศนคติดี (มีคะแนนมากกว่า 27 คะแนน)	104	26.00

ที่มา: จากการออกภาคสนามรอบ Final-survey

ตารางที่ 4.4
ค่าสถิติเบื้องต้นของตัวแปรต้น

	AGE	SEX	EDU	INC	INC2	FAM	PEST	HEAL	KNOW	INFO	BELIEF	BUY	KNOW2	ATT
Min	17	0	1	1	1	1	0	0	0	0	0	0	4	11
Max	60	1	8	10	10	12	1	1	1	1	1	1	14	35
Mean	32.55	0.85	5.52	2.75	3.57	4.06	0.54	0.75	0.05	0.08	0.20	0.29	7.41	24.85
Median	30	1	6	2	2	4	1	1	0	0	0	0	7	25
S.D.	10.29	0.36	1.22	1.52	2.84	1.76	0.50	0.44	0.21	0.27	0.40	0.45	1.56	4.20

ที่มา: จากการคำนวณ

จากตารางที่ 4.4 จะเห็นว่าตัวแปรต้นนอกเหนือไปจาก Bid variables มีทั้งหมด 14 ตัว แบ่งออกเป็นตัวแปรหุ่น 7 ตัว และเป็นตัวแปรที่ไม่ใช่ตัวแปรหุ่น 7 ตัว ในที่นี้มีตัวแปร *BUY* และ *BELIEF* เพิ่มเข้ามาอีก 2 ตัว โดยที่ตัวแปร *BUY* หมายถึง ความบ่อยครั้งในการซื้อสินค้าอาหารประเภทที่มีความปลอดภัยของผู้บริโภค หากซื้อเป็นประจำหรือค่อนข้างบ่อยจะมีค่าเท่ากับ 1 แต่หากไม่ค่อยซื้อหรือไม่เคยซื้อเลยจะมีค่าเท่ากับ 0 ซึ่งคาดว่าจะมีผลในทิศทางเดียวกับความเต็มใจที่จะจ่ายด้วยเหตุผลเดียวกันกับตัวแปร *HEAL* ในขณะที่ตัวแปร *BELIEF* หมายถึง ความเชื่อมั่นในระบบสินค้าเกษตรและอาหารของไทยในปัจจุบัน มีค่าเท่ากับ 1 หากเชื่อมั่นมากหรือค่อนข้างมาก และมีค่าเท่ากับ 0 หากไม่เชื่อมั่นหรือเชื่อมั่นค่อนข้างน้อย คาดว่าจะมีผลในทิศทางตรงข้ามกับความเต็มใจที่จะจ่ายเนื่องจากหากผู้บริโภคเชื่อมั่นว่าระบบปัจจุบันมีความปลอดภัยสูงและมีการตรวจสอบรับรองที่ได้มาตรฐานอยู่แล้วก็ไม่จำเป็นที่จะต้องจ่ายแพงขึ้นเพื่อซื้อสินค้าเกษตรอินทรีย์แต่อย่างใด ทั้งนี้ ทั้ง 2 ตัวแปรได้ใส่เข้ามาเองในแบบสอบถามไม่มีในงานศึกษาในอดีตแต่อย่างใด

ถึงแม้ว่ากลุ่มตัวอย่างที่ได้จากการสำรวจภาคสนามรอบ Final-survey อาจมีการกระจุกตัวในบางตัวแปร เช่น มีเพศหญิงถึง 85% ในขณะที่กลุ่มตัวอย่าง 95% คิดว่าตนเองมีความรู้ความเข้าใจเกี่ยวกับสินค้าเกษตรอินทรีย์น้อยหรือไม่มีเลย และประมาณ 92% คิดว่าข้อมูลข่าวสารเกี่ยวกับสินค้าเกษตรอินทรีย์ในปัจจุบันไม่เพียงพอ เป็นต้น แต่เชื่อว่าการกลุ่มตัวอย่างดังกล่าวสามารถใช้เป็นตัวแทนของประชากรได้ในดีพอสมควร เนื่องจากเชื่อว่าการกระจุกดังกล่าวไม่

แตกต่างกันไปจากประชากรมากนัก เนื่องจากได้มีการกระจายพื้นที่ในการสำรวจพอสมควร (ตารางที่ 4.2)

ตารางที่ 4.5
ค่าสัมประสิทธิ์สหสัมพันธ์ระหว่างตัวแปรต้น

	AGE	SEX	EDU	INC	INC2	FAM	PEST	HEAL	KNOW	INFO	BELIEF	BUY	KNOW2	ATT
AGE	1.00	0.03	0.01	0.52	0.23	-0.09	0.20	-0.05	0.08	0.00	-0.19	0.00	0.11	0.05
SEX	0.03	1.00	-0.08	-0.07	-0.06	-0.06	-0.04	-0.02	0.04	-0.06	-0.06	0.01	-0.03	0.04
EDU	0.01	-0.08	1.00	0.33	0.37	-0.02	0.05	0.12	0.12	-0.12	-0.09	0.06	0.15	0.06
INC	0.52	-0.07	0.33	1.00	0.58	-0.03	0.22	0.04	0.07	-0.06	-0.13	0.06	0.08	0.16
INC2	0.23	-0.06	0.37	0.58	1.00	0.22	0.11	-0.02	0.08	-0.06	-0.13	0.14	0.00	0.10
FAM	-0.09	-0.06	-0.02	-0.03	0.22	1.00	-0.07	-0.01	0.06	0.04	-0.02	0.07	-0.04	-0.07
PEST	0.20	-0.04	0.05	0.22	0.11	-0.07	1.00	0.14	0.03	-0.07	-0.02	0.10	0.04	0.04
HEAL	-0.05	-0.02	0.12	0.04	-0.02	-0.01	0.14	1.00	-0.04	0.04	-0.04	0.13	0.00	0.03
KNOW	0.08	0.04	0.12	0.07	0.08	0.06	0.03	-0.04	1.00	0.12	-0.05	0.13	-0.03	0.11
INFO	0.00	-0.06	-0.12	-0.06	-0.06	0.04	-0.07	0.04	0.12	1.00	0.09	0.02	-0.08	-0.05
BELIEF	-0.19	-0.06	-0.09	-0.13	-0.13	-0.02	-0.02	-0.04	-0.05	0.09	1.00	0.04	-0.18	0.11
BUY	0.00	0.01	0.06	0.06	0.14	0.07	0.10	0.13	0.13	0.02	0.04	1.00	-0.07	0.18
KNOW2	0.11	-0.03	0.15	0.08	0.00	-0.04	0.04	0.00	-0.03	-0.08	-0.18	-0.07	1.00	-0.13
ATT	0.05	0.04	0.06	0.16	0.10	-0.07	0.04	0.03	0.11	-0.05	0.11	0.18	-0.13	1.00

ที่มา: จากการคำนวณ

เมื่อพิจารณาค่าสัมประสิทธิ์สหสัมพันธ์ (Correlation coefficient) ระหว่างตัวแปรต้นทั้งหมด (ตารางที่ 4.5) พบว่า ตัวแปรต้นที่นำมาใช้ในแบบจำลองทุกตัวแทบจะไม่มีความสัมพันธ์ซึ่งกันและกันในระดับสูงกว่า 0.5 เลย ตัวแปรที่มีความสัมพันธ์ซึ่งกันและกันสูงที่สุด คือ INC และ INC2 ซึ่งมีความสัมพันธ์ในทิศทางเดียวกันเท่ากับ 0.58 ดังนั้น สามารถสรุปในเบื้องต้นได้ว่าข้อมูลที่ใช้ไม่มีปัญหา Multicollinearity

4.4 การเลือกแบบจำลองที่เหมาะสม

การเลือกแบบจำลองที่เหมาะสมในที่นี้จะใช้วิธีการใส่ตัวแปรต้นที่มีอยู่ทั้งหมด (14 ตัว) ไปในแบบจำลอง แล้วจึงค่อยเลือกตัดตัวแปรบางตัวทิ้ง โดยทำทุกกรณีเริ่มจากการตัดตัวแปรออกทีละตัวไปจนกระทั่งตัวตัวแปรต้น (ที่ไม่ใช่ Bid variable) ออกทุกตัว

เมื่อพิจารณาจาก Adjusted McFadden's Pseudo R^2 พบว่าแบบจำลองที่มีค่า Adjusted McFadden Pseudo R^2 สูงที่สุด มีค่าเท่ากับ 0.70 อย่างไรก็ตามแบบจำลองดังกล่าวได้ตัดตัวแปรสำคัญทางทฤษฎีตัวใดตัวหนึ่งทิ้งไป เช่น รายได้ (*INC*) ความกังวลเรื่องสุขภาพจากการบริโภคอาหาร (*HEAL*) ความรู้ความเข้าใจในสินค้าเกษตรอินทรีย์ที่ได้จากการประเมิน (*KNOW2*) และทัศนคติต่อสินค้าเกษตรอินทรีย์ (*ATT*) ดังนั้น ในการเลือกแบบจำลองที่เหมาะสมจึงไม่สามารถใช้เกณฑ์ของ Adjusted McFadden's Pseudo R^2 สูงสุดแต่เพียงอย่างเดียวมาตัดสินใจในการเลือกแบบจำลองได้

ดังนั้น ในการเลือกแบบจำลองที่เหมาะสมนั้น นอกจากจะดูเกณฑ์ของ Adjusted McFadden's Pseudo R^2 สูงสุดแล้ว ยังพิจารณาเฉพาะแบบจำลองที่มีตัวแปรสำคัญทางทฤษฎีครบถ้วนอีกด้วย โดยตัวแปรสำคัญทางทฤษฎีในที่นี้หมายถึง ตัวแปรได้จากการทบทวนวรรณกรรมปริทรรศน์ที่มีอยู่ในบทที่ 2 ตารางที่ 2.9 ที่มีเครื่องหมายเหมือนกัน (สามารถตั้งสมมติฐานได้) ซึ่งได้แก่ รายได้ (*INC*) ความกังวลในการบริโภคสารเคมีตกค้าง (*PEST*) ความกังวลเรื่องสุขภาพ (*HEAL*) ความรู้ความเข้าใจในสินค้าเกษตรอินทรีย์ (*KNOW2*) และทัศนคติต่อสินค้าเกษตรอินทรีย์ (*ATT*)

แบบจำลองที่เหมาะสมที่สุดมีค่า Adjusted McFadden's Pseudo R^2 เท่ากับ 0.682 ซึ่งสูงที่สุดเมื่อเทียบกับแบบจำลองอื่นที่มีตัวแปรสำคัญทางทฤษฎีครบถ้วน โดยแบบจำลองดังกล่าวได้ตัดตัวแปรทิ้งไป 4 ตัวแปร ได้แก่ อายุ (*AGE*), เพศ (*SEX*), ระดับการศึกษาสูงสุด (*EDU*), และความบ่อยครั้งในการซื้ออาหารปลอดภัย (*BUY*) ดังนั้น แบบจำลองที่เหมาะสมและนำไปอธิบายผลในบทที่ 5 ผลการศึกษาจะเป็นดังนี้

$$WTP = \alpha - \rho B + \beta_1 INC + \beta_2 INC2 + \beta_3 FAM_2 + \beta_4 PEST + \beta_5 HEAL + \beta_6 KNOW \\ \beta_7 INFO + \beta_8 BELIEF + \beta_9 KNOW_2 + \beta_{10} ATT + \varepsilon \quad (4.1)$$