

Perkembangan Structural Equation Model (SEM) Pada Analisis Technology Acceptance Model (TAM)

Pendekatan Bayesian pada Data Sampel Kecil

Margaretha Ari Anggorowati

Mahasiswa Program S3 Jurusan Statistika-FMIPA
Institut Teknologi Sepuluh Nopember (ITS)
Surabaya, Indonesia
m.ari@bps.go.id

Dr. Suhartono, M.Sc.

Jurusan Statistika-FMIPA
Institut Teknologi Sepuluh Nopember (ITS)
Surabaya, Indonesia
suhartono@statistika.its.ac.id

Prof. Drs. Nur Iriawan, M.Ikomp., Ph.D.

Jurusan Statistika-FMIPA
Institut Teknologi Sepuluh Nopember (ITS)
Surabaya, Indonesia
Nur_i@statistika.its.ac.id

Dr. Hayim Gautama

.Kementrian Komunikasi dan Informasi
hasyim@gmail.com

Abstract—*Technology Acceptance Model (TAM)* adalah model yang digunakan dalam analisis penerimaan pengguna (*user acceptance*) untuk mau menerima dan menggunakan suatu teknologi baru. Metode statistik pada analisis TAM menjelaskan hubungan antar konstruk di dalam model. Penggunaan metode statistik dalam analisis TAM berkembang sesuai dengan beberapa batasan seperti bentuk kasus adopsi teknologi dan sampel kecil. Paper ini mempelajari perkembangan metode statistik yang digunakan dalam TAM untuk data dengan sampel kecil. *Structural Equation Model (SEM)* dengan pendekatan Bayesian adalah metode yang dapat digunakan dalam analisis TAM dengan sampel kecil.

Kata Kunci: TAM, SEM, Bayesian, MCMC

I. PENDAHULUAN

Technology Acceptance Model (TAM) adalah model yang digunakan untuk menjelaskan bagaimana pengguna (user) mau menerima dan menggunakan suatu teknologi baru [1]. TAM dikembangkan oleh Fred Davis pada tahun 1989 dan merupakan pengembangan dari Theory of Action Reaction (TRA) yang dikembangkan oleh Fisbein dan Azjen pada tahun 1975. Dalam menggambarkan penerimaan user pada suatu teknologi baru, TAM menjelaskan dengan struktur eksternal variabel yang dipengaruhi oleh internal variabel yaitu *believe*, *behavior* dan *intention*. Pada TAM terdapat dua konstruk utama yaitu *perceived ease of use* (PEOU) dan *perceived ease of use* (PU).

Dalam tiga dasawarsa TAM banyak digunakan untuk mengukur berbagai proses adopsi teknologi. Implementasi TAM dalam berbagai lingkungan dan berbagai jenis teknologi membawa perkembangan TAM secara teoritikal. Perkembangan TAM secara teoritikal yaitu berkembangnya eksternal variabel pada struktur TAM. Selain struktur TAM,

perkembangan lain pada TAM adalah metode statistik yang digunakan di dalam analisis hubungan antar konstruk di dalam struktur. Berdasarkan studi literatur, *Structural Equation Model (SEM)* adalah metode statistik yang banyak digunakan dalam analisis TAM. Beberapa penelitian TAM yang menggunakan metode Statistik SEM adalah [2] menggunakan SEM dengan *nested model*, [3], [4] menggunakan data *self report*, dan [5] yang menggunakan SEM pada *longitudinal study*.

Berdasarkan studi literatur yang dilakukan pada jurnal-jurnal penelitian TAM, terdapat beberapa permasalahan yang muncul. Permasalahan tersebut diantaranya adalah: kondisi multi kultur (*multi cultural condition*), sampel yang homogen (*homogeny sample*), dan sampel kecil (*small sample size*).

Analisis TAM dengan data sampel kecil menjadi permasalahan tersendiri karena berkaitan dengan metode statistik yang akan digunakan di dalam analisis hipotesis hubungan antar konstruk di dalam struktur TAM. Permasalahan ini muncul khususnya jika dalam analisis TAM digunakan SEM. Seperti [6] mengatakan bahwa SEM adalah metode tingkat lanjut yang *powerful* dalam analisis TAM. SEM tidak hanya dapat digunakan untuk validasi dari teoritikal model tetapi dapat mereduksi indikator. Tetapi, di sisi lain, SEM membutuhkan terpenuhinya beberapa asumsi, diantaranya adalah normalitas dan linieritas. Untuk dapat memenuhi asumsi normalitas, maka persyaratan sampel besar menjadi salah satu keharusan dalam analisis TAM dengan menggunakan SEM. Peneliti pada [7] mengatakan untuk analisis CFA setidaknya dibutuhkan sampel sebanyak 200. Sedangkan [8] mengatakan bahwa sampel sebesar 161 terlalu kecil untuk estimasi empat konstruk pada struktur TAM.

II. STRUCTURAL EQUATION MODEL (SEM)

A. SEM Klasik (Standar)

Structural Equation Model (SEM) adalah keluarga model statistik yang digunakan untuk menjelaskan relasi antara banyak variabel. Dalam hal ini SEM menguji struktur relasi yang ada dalam bentuk beberapa persamaan seperti pada regresi berganda [9]. SEM banyak digunakan dalam analisis pada penelitian sosial, perilaku, pendidikan, kesehatan, pemasaran dan ekonomi. Salah satu alasan mengapa SEM banyak digunakan adalah SEM memungkinkan peneliti untuk memiliki metode yang komprehensif dalam menghitung dan menguji model teori yang ada [10]. Menurut [9] hal yang membedakan SEM dengan metode multivariat yang lain adalah bahwa SEM memisahkan relasi pada tiap-tiap variabel bebas, dan estimasi dilakukan secara terpisah tetapi saling bergantung dengan regresi berganda pada beberapa persamaan secara serentak (dengan menggunakan program statistic). Karakteristik lain dari SEM adalah :

- SEM umumnya digunakan pada variabel yang tidak dapat diukur secara langsung,
- SEM menghitung error pada semua variabel terukur,
- Model umumnya sesuai dengan matrik varian kovarian dari variabel terukur.

SEM menjadi metode analisis yang komprehensif dikarenakan analisis pada struktur model akan melibatkan beberapa proses, yaitu : Analisis Jalur, Model Analisis Faktor Konfirmatori dan Model Regresi.

Suatu model SEM merupakan suatu sistem persamaan simultan. Persamaan ini terdiri dari variabel random, parameter struktur dan mungkin juga untuk variabel bukan random. Random variabel sendiri terdiri dari variabel laten, variabel terukur dan variabel error. Variabel laten mengacu pada suatu konsep yang abstrak dan tidak dapat diukur secara langsung. Model dari variabel laten yang dijabarkan dalam persamaan struktural yang mencakup seluruh relasi antar variabel laten [11].

Pada SEM dikenal dua bentuk variabel laten, yaitu variabel laten endogen dan variabel laten eksogen. Variabel laten eksogen sebagai variabel bebas dilambangkan dengan ξ dan variabel laten endogen sebagai variabel terikat dilambangkan dengan η . Sebagai contoh jika dalam model terdapat dua variabel endogen dan satu variabel eksogen maka variabel laten dalam model dapat digambarkan :

$$\eta_1 = \gamma_{11}\xi_1 + \zeta_1$$

$$\eta_2 = \beta_{21}\eta_1 + \gamma_{21}\xi_1 + \zeta_2$$

Dalam notasi matrik persamaan (2.1) dapat dituliskan :

$$\begin{bmatrix} \eta_1 \\ \eta_2 \end{bmatrix} = \begin{bmatrix} 0 & 0 \\ \beta_{21} & 0 \end{bmatrix} \begin{bmatrix} \eta_1 \\ \eta_2 \end{bmatrix} + \begin{bmatrix} \gamma_{11} \\ \gamma_{21} \end{bmatrix} \begin{bmatrix} \xi_1 \end{bmatrix} + \begin{bmatrix} \zeta_1 \\ \zeta_2 \end{bmatrix}$$

dan persamaan struktural dari model tersebut adalah

$$\eta = B\eta + \Gamma\xi + \zeta$$

dengan :

η = variabel laten endogen

$B = m \times m$ koefisien matrik

$\Gamma = m \times n$ koefisien matrik

ξ = variabel laten eksogen

$\zeta = p \times 1$ vektor error

Persamaan dari model terukur dapat dituliskan :

$$x = \Lambda_x \xi + \delta$$

$$y = \Lambda_y \eta + \epsilon$$

Pada SEM klasik (standar) estimasi dilakukan berdasarkan matrik varian kovarian, sehingga dibutuhkan sampel besar agar terpenuhi asumsi normalitas.

Elemen kunci dari SEM adalah parameter yang disebut sebagai parameter model, dan parameter ini umumnya belum diketahui. Parameter yang belum diketahui diestimasi dengan menggunakan matrik sampel varian kovarian, sehingga model menjadi sesuai (fit). Beberapa metode estimasi utama pada SEM adalah Unweighted Least Squares (ULS), Generalized Least Squares (GLS), Maximum Likelihood (ML), dan Asymptotically Distribution Free (ADF). Estimasi pada SEM membutuhkan kondisi dimana model dapat diidentifikasi melalui jumlah parameter dan jumlah data dalam model tersebut. Jika t adalah jumlah parameter dan q adalah jumlah variabel terukur maka jumlah parameter harus lebih kecil dari jumlah data yang ada sehingga model dapat dikatakan teridentifikasi (identified) dan sebaliknya jika jumlah parameter lebih besar dari jumlah data maka dikatakan bahwa model tidak dapat diidentifikasi (unidentified).

$$t \leq \frac{(p+q)(q+1)}{2}$$

p = jumlah variabel terukur

t = jumlah parameter

Hipotesis dari model persamaan struktural adalah :

$$H_0 : \Sigma = \Sigma(\theta)$$

$$H_1 : \Sigma \neq \Sigma(\theta)$$

Dengan Σ adalah matrik varian kovarian populasi sedangkan $\Sigma(\theta)$ adalah matrik varian kovarian sebagai fungsi dari model parameter θ . Matrik varian kovarian yang dapat dibentuk adalah matrik varian kovarian x , matrik varian kovarian y , matrik varian kovarian xy . Untuk mendapatkan ketiga matrik varian kovarian tersebut, maka harus diketahui matrik varian kovarian antar variabel laten eksogen, matrik varian kovarian antar variabel laten endogen dan matrik varian kovarian antar variabel laten eksogen dan endogen.

Matrik varian kovarian x dan y adalah

$$\Sigma = \begin{bmatrix} \Sigma_{yy} & \Sigma_{yx} \\ \Sigma_{xy} & \Sigma_{xx} \end{bmatrix}$$

$$\Sigma = \begin{bmatrix} \Lambda_y(I-B)^{-1}(\Gamma\Phi\Gamma^T + \Psi)((I-B)^{-1})^T\Lambda_y^T + \Theta_\epsilon & \Lambda_y\Phi\Gamma^T((I-B)^{-1})^T\Lambda_x^T \\ \Lambda_x\Phi\Gamma^T((I-B)^{-1})^T\Lambda_y^T & \Lambda_x\Phi\Lambda_x^T + \Theta_\delta \end{bmatrix}$$

B. Prosedur dalam SEM

Secara garis besar prosedur dalam analisis SEM standar meliputi tahapan-tahapan sebagai berikut:

- **Spesifikasi model**
Spesifikasi model struktur dilakukan sebelum proses estimasi. Model struktur ini diformulasikan berdasarkan suatu teori atau penelitian sebelumnya.
- **Identifikasi**
Identifikasi model adalah tahapan dimana investigasi apakah suatu model dapat diestimasi atau tidak. Hal ini berkaitan dengan kemungkinan parameter yang akan diestimasi di dalam model.
- **Estimasi**
Estimasi adalah tahapan untuk mendapatkan nilai parameter di dalam model.
- **Uji kecocokan**
Uji kecocokan model adalah uji yang dilakukan antara model dengan data. Kriteria uji kecocokan disebut sebagai Goodness Of Fit (GOF).

C. SEM Non Standar

SEM non standar adalah kondisi dimana asumsi normalitas dan linearitas pada SEM tidak dapat dipenuhi. SEM non standar dapat disebabkan oleh karena data sampel kecil, adanya efek non-linier pada model karena adanya hubungan non linier antar variabel laten atau adanya model *mixture*. Pada kondisi SEM non standar, metode estimasi dengan SEM kalsik tidak dapat digunakan.

Salah satu metode yang digunakan untuk mengatasi permasalahan SEM non standar khususnya pada data sampel kecil adalah Partial Least Square (PLS). Menurut [12] PLS dapat mengatasi dua permasalahan di dalam SEM, yaitu sampel kecil dan data tidak lengkap (*missing data*). Namun demikian [12] mengatakan bahwa PLS tidak memiliki kemampuan secara khusus untuk mengatasi permasalahan sampel kecil, tetapi PLS menjadi metode yang bisa digunakan ketika jumlah data sampel tidak cukup untuk analisis berbasis matrik kovarian (covariance based structural equation model /CBSEM).

III. SEM BAYESIAN

SEM dengan pendekatan Bayesian dikembangkan oleh [11]. Berbeda dengan SEM standar, SEM Bayesian tidak bekerja berdasarkan matrik varian kovarian, tetapi berdasarkan raw data observasi. Hal ini membuat analisis menjadi sederhana karena estimasi dilakukan pada momen pertama. Pada SEM standar dengan metode maximum likelihood, estimasi dilakukan pada momen kedua. Hal lain yang menjadi kemampuan SEM Bayesian adalah dilibatkannya informasi prior dalam estimasi. Informasi prior menunjukkan informasi dari data yang digunakan. Dengan informasi prior analisis dapat dilakukan lebih akurat.

Pada estimasi dengan pendekatan Bayesian, parameter θ diperlakukan sebagai random variabel yang memiliki distribusi dan disebut sebagai distribusi prior. Jika M adalah persamaan

SEM dengan vektor parameter θ yang belum diketahui maka fungsi kepadatan peluang dari θ adalah $p(\theta|M)$. Jika Bayesian inferen dilakukan pada data observasi Y maka distribusi bersama dari Y dan θ pada M adalah $p(Y, \theta|M)$.

Perilaku θ pada data Y digambarkan dengan distribusi bersyarat θ oleh Y , dan dikenal sebagai distribusi posterior $p(\theta|Y, M)$. Dengan

$$p(Y, \theta|M) = p(Y|\theta, M) p(\theta) = p(\theta|Y, M) p(Y|M)$$

dan bahwa $p(Y|M)$ tidak bergantung pada θ maka :

$$\log p(\theta|Y, M) \propto \log p(Y|\theta, M) + \log p(\theta)$$

dimana : $p(Y|\theta, M)$ adalah fungsi likelihood.

A. Konsep Data Augmentation

Konsep data augmentation digunakan pada simulasi posterior. Data augmentation dilakukan dengan memperlakukan laten sebagai hipotesa data tidak lengkap (hypotetical missing data) dan kemudian menambahkan data observasi dengan laten sehingga dapat dilakukan analisis distribusi posterior dengan data lengkap. Menurut [13], [14], dan [15] strategi ini sangat berguna dalam metode statistik. Dengan konsep data augmentation dengan menggunakan Markov Chain Monte Carlo (MCMC) dan algoritma Gibbs Sampler maka permasalahan data sampel kecil dapat teratasi.

B. Distribusi Prior

Distribusi prior θ menggambarkan distribusi nilai parameter yang mungkin ada dalam suatu model. Terdapat dua jenis distribusi prior yaitu distribusi prior non-informatif dan distribusi prior informatif. Distribusi prior non-informatif terjadi pada situasi dimana distribusi prior tidak memiliki populasi, atau terdapat sedikit informasi prior. Pada keadaan demikian distribusi prior memiliki pengaruh yang kecil pada distribusi posterior. Sebaliknya pada distribusi prior informatif, informasi distribusi prior dapat diketahui. Distribusi prior yang umum digunakan pada pendekatan Bayesian adalah conjugate prior distribution. Berdasarkan [16], distribusi prior pada persamaan pengukuran yang digunakan adalah:

$$\begin{aligned} \psi_{ek}^{-1} &\stackrel{D}{=} \text{Gamma}[\alpha_{0ek}, \beta_{0ek}] \\ [\Lambda_k | \psi_{ek}] &\stackrel{D}{=} N[\Lambda_{0k}, \psi_{ek} H_{0yk}] \\ \Phi^{-1} &\stackrel{D}{=} W_q[R_0, \rho_0] \end{aligned}$$

Dimana $\text{Gamma}(\alpha, \beta)$ adalah distribusi dengan parameter $\alpha > 0$ dan invers parameter $\beta > 0$ $W[.,.]$ adalah distribusi Wishart dengan dimensi q , $\alpha_{0ek}, \beta_{0ek}, \Lambda_{0k}, \rho_0$ adalah matrik definit positif, sedangkan H_{0yk} dan R_0 adalah hiperparameter yang didapat dari penelitian sebelumnya.

C. Analisis Posterior

Estimasi θ pada bayesian didefinisikan sebagai rata-rata atau model dari $p(\theta|y)$ dan dinyatakan dalam persamaan

$$\log p(\theta|Y, M) \propto \log p(Y|\theta, M) + \log p(\theta)$$

dengan mencari nilai maksimum dari

$$\log p(Y|\theta, M) + \log p(\theta).$$

Secara teori rata-rata dari distribusi posterior ($\theta|y$) dapat dihasilkan dengan cara penggabungan /integrasi, tetapi sebagian besar situasinya adalah bahwa penggabungan tersebut tidak memiliki bentuk close form. Untuk SEM non standar, distribusi posterior ($\theta|y$) dapat bersifat lebih kompleks dan sulit untuk melakukan simulasi terhadap data observasi serta sulit mendapatkan bentuk distribusi dari model SEM tersebut. Pada [16] menyatakan bahwa salah satu cara untuk melakukan simulasi posterior adalah data augmentation (penambahan data) seperti yang diteliti oleh [17]. Pada pendekatan SEM-Bayesian analisis dilakukan dengan $p(\theta, \Omega|Y)$ dimana Ω adalah kumpulan variabel laten pada model. Pada kasus tertentu bentuk $p(\theta, \Omega|Y)$ umumnya belum memenuhi close form [11].

D. Gibbs Sampler pada Simulasi Posterior

Gibbs sampling adalah bagian dari Markov Chain Monte Carlo (MCMC) dimana transisi kernel dilakukan dengan distribusi bersyarat penuh [14]. Untuk data sampel kecil pendekatan metode klasik GLS dan ML menghadapi banyak kesulitan pada proses estimasi parameter θ . Untuk menangani kesulitan pada estimasi maka pada analisis posterior, data Y ditambahkan dengan matrik variabel laten $\Omega = (\omega_1 \dots \omega_n)$ (konsep data augmentation), sehingga algoritma Gibbs Sampler yang digunakan untuk menghasilkan obeservasi pada distribusi posterior $[\theta, \Omega|Y]$ adalah :

- (i) Bangkitkan $\theta^{(j+1)}$ dari $p(\theta | \Omega^{(j)}, Y)$
- (ii) Bangkitkan $\Omega^{(j+1)}$ dari $p(\Omega | \theta^{(j+1)}, Y)$

E. Prosedur pada SEM Bayesian

Analisis SEM dengan pendekatan Bayesian dilakukan dalam tahap-tahap yang berbeda dengan SEM standar. Pada pendekatan Bayesian tahapan identifikasi tidak dilakukan. Permasalahan jumlah sampel kecil dapat diatasi dengan menerapkan konsep data augmentation menggunakan MCMC dan Gibbs Sampler, sehingga simulasi posterior dapat dilakukan berdasarkan data lengkap. Tahapan-tahapan dalam analisis SEM dengan pendekatan Bayesian adalah:

- Analisis Diskriptif
Analisis diskriptif adalah langkah awal yang dilakukan untuk melihat gambaran umum dari data. Dengan menganalisis sebaran data dan nilai skewness dari semua variabel yang digunakan dalam penelitian, dapat diputuskan penanganan yang tepat untuk data kategori dari setiap variabel. Selain itu, nilai frekuensi dari setiap kategori pada setiap variabel adalah informasi penting untuk menentukan nilai threshold.
- Penentuan Parameter dalam model
Setelah diketahui gambaran umum data observasi, maka diformulasikan stuktur SEM untuk model TAM yang akan dianalisis. Dengan diformulasikan model struktur SEM maka dapat ditentukan parameter-parameter yang akan diestimasi. Pada SEM Bayesian, parameter yang akan diestimasi adalah: θ ($\Phi, \psi, \Lambda, \Gamma, B$).

- Penentuan Threshold pada Model
Untuk data diskrit, dibutuhkan pendekatan yang lebih baik dengan menganggap data kategori diskrit berasal dari distribusi kontinyu normal. dengan mengaplikasikan suatu nilai *threshold*. Pendekatan *threshold* merupakan model yang fleksibel, dan memberikan interpretasi yang mudah untuk parameter model.
- Analisis Hubungan Antar Variabel Laten
Analisis hubungan antar variabel laten dengan menggunakan plot perlu dilakukan untuk melihat kemungkinan adanya hubungan nonlinier antar laten. Jika ternyata terdapat ploa hubungan non linier maka diperlukan tahap analisis yang berbeda.
- Penentuan Parameter pada Distribusi Prior
Penentuan distribusi prior merupakan bagian penting dalam analisis Bayesian. Penentuan distribusi prior akan ikut menentukan signifikansi dari parameter di dalam model. Distribusi prior dapat ditentukan berdasarkan pengetahuan peneliti-peneliti sebelumnya, dan dengan melakukan analisis khusus terhadap data penelitian. dalam SEM Bayesian distribusi prior yang digunakan adalah *conjugate prior* dengan nilai hiperparameter yang telah ditentukan.
- Estimasi parameter
Estimasi parameter dilakukan dengan melakukan simulasi posterior dengan menggunakan MCMC dan algoritma Gibbs Sampler.
- Analisis Residual
Analisis residual dilakukan untuk melihat kebaikan model. Model dengan residual terkecil akan menunjukkan bahwa model tersebut adalah model terbaik.

F. Studi Kasus

Analisis TAM dengan menggunakan SEM Bayesian sudah dilakukan pada studi kasus adopsi teknologi pengolahan data sensus penduduk di Badan Pusat Statistik [18]. Pada studi kasus tersebut data sampel yang tersedia sejumlah 37. Analisis TAM dengan SEM Bayesian dilakukan dengan menggunakan software WinBugs 14.

IV. KESIMPULAN

Perkembangan metode statistik dalam analisis TAM berkembang berdasarkan permasalahan dan limitasi yang muncul dalam riset-riset yang dilakukan oleh para peneliti. Salah satu permasalahan utama yang muncul dalam riset TAM adalah data sampel kecil. SEM dengan pendekatan Bayesian dapat digunakan untuk mengatasi permasalahan data sampel kecil.

Terdapat perbedaan tahapan dalam analisis pada SEM standar dan SEM Bayesian. Pada SEM Bayesian, tidak dibutuhkan tahapan identifikasi model. Data dengan sampel kecil dapat diatasi dengan menerapkan konsep data augmentation.

MCMC dan Gibbs Sampler digunakan dalam simulasi posterior. Simulasi posterior dilakukan dalam estimasi parameter model.

Distribusi prior salah satu hal penting dalam analisis posterior. Distribusi prior akan mempengaruhi akurasi dan signifikansi dari parameter yang diestimasi.

DAFTAR PUSTAKA

- [1] A. Alrafi, B. Noble, "Technology Acceptance model," unpublished, 2005
- [2] A. K. Gyampah, A.F. Salam, "An extension of technology acceptance model in an ERP implementation environment", in *Information and Management*, vol. 41, pp. 731-745, 2004.
- [3] J. Wu, Y. Chen, and L. Lin, "Empirical evaluation of revised end user computing acceptance model," in *Computers in Human Behaviour*, vol. 23, pp. 162-174, 2007.
- [4] J. Wu, Y. Chen, and L. Lin, "Empirical evaluation of revised end user computing acceptance model", in *Computers in Human Behaviour*, vol. 23, pp. 162-174, 2007.
- [5] T. Teo, C. Lee, C. Chai and S. Wong, "Assessing the intention to use technology among pre-service teachers in Singapore and Malaysia: A Multigroup invariance analysis of the Technology Acceptance Model (TAM)", in *Computers and Education*, vol. 53, pp. 1000-1009, 2009.
- [6] S. Hung, K. Tang, C. Chang, and C. Ke, "User acceptance of intergovernmental services: An example of electronic document management system," in *Government Information Quarterly*, vol. 26, pp. 387-397, 2009.
- [7] X. Deng, W. Doll, A. R. Hendrickson and J. A. Scazzero, "A multi-group analysis of structural invariance: an illustration using the technology acceptance model", in *Information and Management*, vol. 42, pp. 745-759, 2005.
- [8] I. Im, Y. Kim, and H. Han, "The effect of perceived risk and technology type on users' acceptance of technologies" in *Information and Management*, vol. 45, pp. 1-9, 2008.
- [9] J. F. Hair, W.C. Black, B.J. Babin, R.E. Anderson, and R.L. Tatham, *Multivariate data Analysis*, New Jersey: Pearson Education, Inc, 2006.
- [10] T. Raykov, G.A. Marcoulides, *A First Course in Structural Equation Modelling*, New Jersey: Lawrence Erlbaum Associates, Inc, 2006.
- [11] K. Bollen, *Structural Equation with Latent Variable*. New York: John Wiley & Sons, 1989.
- [12] Henseler, Jorg; Ringle, M. Christian, Sinkovics, R. Rudolf R., *The use of Partial Least Square path modeling in international marketing*, New Challenge to International marketing Advance in International Marketing. Emerald Group Publishing Limited.
- [13] D. B. Rubin, "EM and beyond", in *Psychometrika*, vol. 56, pp. 241-256, 1991.
- [14] S.L. Zeger, M.R. Karim, "Generalized linear models with random effects: A Gibbs Sampling approach", in *Journal of the American Statistical Association*, vol. 86, pp. 79-86, 1991.
- [15] J.H. Albert, S. Chib, "Bayesian analysis of binary and polychotomous response data", in *Journal of the American Statistical Association*, vol. 88, pp. 669-679, 1993.
- [16] S. Lee, *Structural Equation Modeling A Bayesian Approach*. England: John Wiley and Sons Ltd, 2007.
- [17] M.A. Tanner, W.H Wong, "The calculation of posterior distributions by data augmentation (with discussion)", in *Journal of the American Statistical Association*, vol. 82, pp. 528-550, 1987.
- [18] M.A. Anggorowati, N. Iriawan, Suhartono, dan H. Gautama, "Restructuring and Expanding Technology Acceptance Model: Structural Equation Model and Bayesian Approach", in *Journal of Applied Sciences*, vol. 9(4), pp. 496-504, 2012.