

Discussion Paper: 2009/02

Information flows around the globe: Predicting opening gaps from overnight foreign stock price patterns

Jan G. de Gooijer, Cees G.H. Diks and Lukasz T. Gatarek

www.fee.uva.nl/ke/UvA-Econometrics

Amsterdam School of Economics

Department of Quantitative Economics
Roetersstraat 11
1018 WB AMSTERDAM
The Netherlands

UvA  UNIVERSITEIT VAN AMSTERDAM



Information Flows Around the Globe: Predicting Opening Gaps from Overnight Foreign Stock Price Patterns

Jan G. De Gooijer*

*Department of Quantitative Economics
University of Amsterdam*

Cees G. H. Diks†

*Department of Quantitative Economics
University of Amsterdam*

Lukasz T. Gatarek‡

*Econometric Institute
Erasmus University Rotterdam*

Abstract

This paper describes a forecasting exercise of close-to-open returns on major global stock indices, based on price patterns from foreign markets that have become available overnight. As the close-to-open gap is a scalar response variable to a functional variable, it is natural to focus on functional data analysis. Both parametric and non-parametric modeling strategies are considered, and compared with a simple linear benchmark model. The overall best performing model is nonparametric, suggesting the presence of nonlinear relations between the overnight price patterns and the opening gaps. This effect is mainly due to the European and Asian markets. The North-American and Australian markets appear to be informationally more efficient in that linear models using only the last available information perform well.

Keywords: Close-to-open gap forecasting; Functional data analysis; International stock markets; Nonparametric modeling.

JEL Classification: C14; C53; F37; G17

*Department of Quantitative Economics and Tinbergen Institute, University of Amsterdam, Roetersstraat 11, NL-1018 WB Amsterdam, The Netherlands. E-mail: J.G.DeGooijer@uva.nl

†Corresponding author: Department of Quantitative Economics and Tinbergen Institute, University of Amsterdam, Roetersstraat 11, NL-1018 WB Amsterdam, The Netherlands. E-mail: C.G.H.Diks@uva.nl

‡Econometric Institute and Tinbergen Institute, Erasmus University Rotterdam, P.O. Box 1738, NL-3000 DR Rotterdam, The Netherlands. E-mail: gatarek@mail.tinbergen.nl

1 Introduction

Empirical research in finance has traditionally focused on the analysis of daily stock returns, usually measured as changes in closing prices. However, since trading can be thought of as a continuous-time process, it is also natural to consider returns over other than daily intervals. Recently, some interest has been developed into dividing daily returns into overnight (close-to-open) returns and daytime returns. There is considerable empirical evidence that return dynamics are different over non-trading periods than during trading periods (French and Roll, 1986; Lockwood and Linn, 1990; Hasbrouck, 1991, 1993; and Madhavan *et al.*, 1997; George and Wang, 2001; Cliff *et al.*, 2008). Accordingly, a number of models have been proposed to quantify this phenomenon, often using stocks traded on a particular stock market; see, e.g., Oldfield and Rogalski (1980) and Hong and Wang (2000).

The information revealed in consecutive overnight and day-time returns can also be employed for prediction. In this vein, Zhong (2007) considered predicting daytime volatility of stock prices based on the preceding overnight returns. As far as we know, there have been no attempts to explore the price evolution in a set of foreign stock markets as a result of the information content revealed during non-trading periods of a home market. With this in mind, one of the aims of this study is to predict the overnight return on an individual stock index of a home market, based on the information content revealed in a set of foreign markets during non-trading hours of the home market. Additionally, we investigate if global markets are informationally efficient in the sense that adding information from clusters of stock indices traded further in the past into the information set does not improve predictive ability. To this end we employ linear regression, as well as parametric and nonparametric Functional Data Analysis (P-FDA and NP-FDA) techniques. FDA is a natural alternative to linear regression in this setting, because overnight foreign price patterns can be viewed as continuous functions of time. Within empirical finance, where the object of interest often depends on some continuous parameter (e.g. continuous time), functional data often arise. For instance, Benko (2006) applies parametric FDA techniques to the analysis of implied volatility functions and yield curve dynamics.

The paper is organized as follows. Section 2 formalizes the prediction problem. Section 3 describes the various FDA methods considered in this paper, as well as their corresponding predictive intervals (PIs). The measures used for evaluating the out-of-sample predictions are described in Section 4. Section 6 describes the results obtained, and Section 7 briefly discusses the results and concludes. To avoid confusion we like to stress that the adjective ‘functional’

refers to the form of the data and ‘parametric/nonparametric’ to the form of the constraints imposed on the model.

2 The Prediction Problem

To formalize the prediction problem some notation and definitions are introduced. Let $\chi_{i,s}$ denote the price pattern (observed curves) of stock index $i \in \{1, \dots, M\}$ across trading session (day) $s \in \{1, \dots, S\}$, with M the number of international stock indices under study, and S the number of trading sessions under consideration. The value of stock index i at within-trading session time $t \in \mathbb{R}$ (measured in five-minute units) is $\chi_{i,s}(t)$, $t \in (0, T_i)$ where $T_i \times 5\text{min.}$ is the duration of each trading session in market i . In practice only discretized versions of $\chi_{i,s}(t)$ can be used. Here discretized data (at regular five-minute intervals) are considered, denoted by $x_{i,s}(t)$, $t \in 1, \dots, T_i$.

To specify the information set on which predictions of close-to-open gaps are to be based, it is convenient to introduce a universal ‘background’ time variable that measures time globally, as opposed to the within-trading session time variable. Time is measured in five-minute units again, and in addition we assume that the universal clock does not run in weekends, between Central European Time (CET) Sat 00:00 and CET Mon 00:00; a period during which all markets are closed simultaneously. Since a 24 hour day contains 288 five-minute intervals, each observation can be represented as $x_{i,s}(t) = \tilde{x}_i(288 \times s + c_i + t)$, where $288 \times s + c_i + t$ is the universal time corresponding to a quote in session s of market i at trading-session time t . The shift $c_i \in 0, \dots, 287$ represents the opening time of market i , again in five-minute units. Note that $\tilde{x}_i(\cdot)$ is only defined for universal times t at which market i is open, and is not available otherwise.

The prediction variable of interest is the close-to-open return in market i for trading session s , given by $y_{i,s} = (x_{i,s}(1) - x_{i,s-1}(T_i)) / x_{i,s-1}(T_i)$, where $s-1$ denotes the last session prior to session s during which market i was open. In universal time, $y_{i,s}$ materializes at time $t_{i,s}^{\text{open}} = 288 \times s + c_i$.

Below, several different specifications are considered, which use various amounts of the information available. Based on trading hours, three global clusters of markets can be distinguished; in order, the Asian-Pacific markets, the European markets, and the American markets. Since we think of price developments as an information discovery process, and hence of more recent available prices being more relevant, we focus on forecasting the opening gap of markets belonging to a specific cluster based on the most recent price patterns in the preceding cluster. For comparison we also consider predictions based on price patterns from the two most recent

preceding clusters. For instance, we consider prediction of the opening gap for the US based on the price patterns in the European markets, and also based on the price patterns in the Asian-Pacific and European markets jointly. In cases of missing data as a result of a holiday in either one of the explanatory variables or the home market of interest, the corresponding explanatory data and opening gap are excluded from the analysis.

The original sample for each market is split into two sub-samples: a learning sample containing the units $\{(\mathbf{x}_{i,s}, y_{i,s})_{s=1,\dots,k_i}\}$, and a testing sample containing the units $\{(\mathbf{x}_{i,s}, y_{i,s})_{s=k_i+1,\dots,S}\}$ where k_i ($i \in 1, \dots, M$) denotes the number of observations in the learning sample. The learning sample allows us to build a functional kernel estimator with optimal smoothing parameter(s); both the $\mathbf{x}_{i,s}$'s and the corresponding $y_{i,s}$'s are used at this stage. The testing sample is used for making actual predictions and evaluating predictive performance.

3 Functional Data Analysis

In our description of functional data analysis we consider predicting the opening gap for a specific market i . For notational convenience, the subscript i is dropped from the respective random variables. Let $(\boldsymbol{\chi}_s, Y_s)_{s=1,\dots,k}$ be $k = k_i$ pairs of random variables, identically distributed as $(\boldsymbol{\chi}_i, Y_i)$ but not necessarily independent, and taking values in $\mathcal{E} \times \mathbb{R}$, where $(\mathcal{E}, d)$ is a semi-metric space with semi-metric d . In addition, it is assumed that $\{\boldsymbol{\chi}_s, Y_s\}$ is strictly stationary. The aim is to predict the unobserved scalar response variable Y_s from the curve(s) $\boldsymbol{\chi}_s$ (covariates). The idea behind FDA is to find close (with respect to a certain norm) covariates among the available past observed covariates. This problem can be viewed as follows. Suppose that there exists a function $r(\cdot)$ modeling the relationship between Y and $\boldsymbol{\chi}$ and that $r(\cdot)$ is defined through the conditional distribution. Given a convex loss function $\ell(\cdot)$ with a unique minimum, define $r(\cdot)$ such that it minimizes the mean $\mathbb{E}(\ell(Y - a) | \boldsymbol{\chi} = \chi)$ with respect to a .

3.1 Nonparametric FDA

A nonparametric estimator of $r(\cdot)$ provides a nonparametric predictor $\hat{Y}$ in terms of $\boldsymbol{\chi}$. Using this principle, we consider three nonparametric predictors, based on different loss functions. Assuming $(\boldsymbol{\chi}_s, Y_s)$ is α -mixing, Ferraty and Vieu (2005, 2006) proved almost complete convergence of the three nonparametric functional predictors considered here. Also these authors established the rates of convergence of the predictors.

Conditional mean

It is well-known that taking $\ell(u) = u^2$ leads to the conditional mean function $r(\chi) = \mathbb{E}(Y|\chi = \chi)$. Recall that the model is to be based on the observed k pairs $(\mathbf{x}_s, y_s)_{s=1, \dots, k}$ of identically distributed random variables, where $\mathbf{x}_s$ is a discretized version of the pattern χ_s . Let $\mathbf{x}$ be an observed curve (overnight foreign price pattern) at which the regression is estimated. Then, using the Nadaraya-Watson kernel density estimator, the one-step-ahead prediction (measured in five-minute units) is defined as $\hat{y}^{\text{mean}} = \sum_{s=1}^k y_s W(\mathbf{x}_s, \mathbf{x})$. Here $W(\cdot)$, the so-called kernel weight, is defined as $W(\mathbf{x}_s, \mathbf{x}) = K(d(\mathbf{x}_s, \mathbf{x})/h) / \sum_{r=1}^k K(d(\mathbf{x}_r, \mathbf{x})/h)$, where h denotes the bandwidth, $K(\cdot)$ the kernel function, and $d(\mathbf{x}_s, \mathbf{x})$ is any semi-metric between $\mathbf{x}_s$ and $\mathbf{x}$.

Conditional median

In this case the loss function is given by $\ell(u) = |u|$. Then the conditional median function is given by $r(\chi) = \inf\{y : F(y|\chi = \chi) \geq 1/2\}$, where $F(\cdot|\cdot)$ is the conditional distribution function of Y given $\chi = \chi$. Consequently, the one-step-ahead nonparametric functional predictor of the conditional median is defined as $\hat{y}^{\text{med}} = \inf\{y : \hat{F}(y|\chi) \geq 1/2\}$, where $\hat{F}(y|\chi) = \sum_{s=1}^k W(\mathbf{x}_s, \mathbf{x}) \mathbf{1}_{\{y_s \leq y\}}$, with $\mathbf{1}_{\{A\}}$ denoting the indicator function of set $\{A\}$, is the estimated conditional cumulative distribution function (CDF) of Y given $\chi = \chi$.

Conditional mode

In this case we have a non-convex loss function with a unique minimum $\ell(u) = 0$ when $u = 0$, and $\ell(u) = 1$ otherwise. The loss function becomes $r(\chi) = \arg \max_{y \in \mathbb{R}} f(y|\chi = \chi)$, where $f(\cdot|\cdot)$ denotes the conditional density function of Y given $\chi = \chi$. Hence, given the observed data, the nonparametric functional predictor of the conditional mode is given by $\hat{y}^{\text{mode}} = \arg \max_{y \in \mathbb{R}} \sum_{s=1}^k K(|y - y_s|/h) W(\mathbf{x}_s, \mathbf{x})$, where, for ease of notation, we assume that the same kernel function $K(\cdot)$ and bandwidth h apply in the y direction.

3.2 Functional Parametric Regression

A functional linear regression establishes a relationship between a functional covariate $\chi_s(t)$ and the response variable Y_s as follows

$$Y_s = \beta_0 + \int_0^T \chi_s(t) \beta(t) dt + \varepsilon, \quad (1)$$

where $\{\varepsilon\}$ is a sequence of i.i.d. random variables such that $\mathbb{E}(\varepsilon|\chi_s(t)) = 0$ and $\mathbb{E}(\varepsilon^2|\chi_s(t)) = \sigma^2 < \infty$.

A popular approach to reduce the number of degrees of freedom in (1) is to use a truncated functional basis expansion, similar to the truncation applied in the NP-FDA case. There are three prominent examples of functional bases: Fourier, Polynomial and B-spline. Here, following Ramsay and Silverman (2005, Chapter 15), we adopt a set of Fourier (orthonormal) basis function $\theta_k(t)$, i.e.

$$\beta(t) = \sum_{k=1}^{K_\beta} b_k \theta_k(t) = \mathbf{b}' \boldsymbol{\theta}(t), \quad (2)$$

where K_β denotes the length of the set, and where $\boldsymbol{\theta}(t) = (\theta_1(t), \dots, \theta_{K_\beta}(t))'$ and $\mathbf{b} = (b_1, \dots, b_{K_\beta})'$. Similarly, $\chi_s(t)$ can be expanded in another set of Fourier basis function $\psi_{k,s}(t)$ of length K_z as follows

$$\chi_s(t) = \sum_{k=1}^{K_z} c_{s,k} \psi_k(t) = \mathbf{c}_s' \boldsymbol{\psi}(t), \quad (3)$$

where $\boldsymbol{\psi}(t) = (\psi_1(t), \dots, \psi_{K_z}(t))'$ and $\mathbf{c}_s = (c_{s,1}, \dots, c_{s,K_z})'$. Inserting (2) and (3) into (1), and using the data in the learning sample, yields

$$Y_s = \beta_0 + \mathbf{C}_s \mathbf{J} \mathbf{b} + \varepsilon_s, \quad (s \in 1, \dots, k),$$

where $\mathbf{J}$ is a $K_z \times K_\beta$ matrix defined by $\mathbf{J} = \int \boldsymbol{\psi}(t) \boldsymbol{\theta}'(t) dt$, and where the $k \times K_z$ matrix is given by $\mathbf{C} = \{c_{s,k} : s = 1, \dots, k, k = 1, \dots, K_z\}$. The notation can be further simplified by defining a $(K_\beta + 1)$ -vector $\boldsymbol{\xi} = (\beta_0 \ \mathbf{b}')'$ and a $k \times (K_\beta + 1)$ matrix $\mathbf{Z} = [\mathbf{1} \ \mathbf{C} \mathbf{J}]$. Thus, the resulting functional regression model has the same structure as the classical linear regression model. Consequently, the augmented parameter vector $\boldsymbol{\xi}$ can be estimated by least squares, i.e. $\hat{\boldsymbol{\xi}} = (\mathbf{Z}' \mathbf{Z})^{-1} \mathbf{Z}' \mathbf{y}$. Clearly, the above setup can be easily modified into a specification with more than one functional covariate. The choice of the numbers of basis functions, K_z and K_β , is a trade-off between information loss and computational costs. In the present study K_z and K_β were set equal to 15. For the specific data used in the present study we verified that there was no gain in performance by increasing the number of basis functions beyond this number.

Given $\hat{\boldsymbol{\xi}}$ and the set of basis functions $\boldsymbol{\theta}(t)$, estimates $\hat{\beta}(t)$ of $\beta_i(t)$ can be obtained. Then, using (1), the predictor for Y_s may be constructed as

$$\hat{y}_s^0 = \hat{\beta}_0 + \int_0^T \chi_s(t) \hat{\beta}(t) dt,$$

where s runs over the collection of available out-of-sample sessions, denoted as $\mathcal{S}$. The number of evaluation sample points available (size) will be denoted by $|\mathcal{S}|$. One feature of above setup is that the conditioning takes place on the same information set as used in predicting a response variable via NP-FDA.

3.3 Linear regression model

To have a benchmark model to compare the FDA results with, we consider a linear regression model which uses as the explanatory variables a constant plus the total overnight returns of the foreign stock indices. The use of overnight returns rather than the complete set of 5-minute observations ensures that the model is parsimonious. Although there are other ways to construct a parsimonious linear model, we have chosen to focus on overnight returns, since within an informationally efficient market these returns should contain all relevant information regarding the development of fundamentals underlying the indices of interest.

3.4 Predictive intervals

Following De Gooijer and Gannoun (2000) we consider two types of PIs: the conditional percentile interval (CPI) and the shortest conditional modal interval (SCMI). The CPI with nominal coverage probability γ is given by $\left(\xi_{\frac{1-\gamma}{2}}(\chi), \xi_{\frac{1+\gamma}{2}}(\chi)\right)$, where $\xi_{\alpha}(\chi)$ denotes the α -th quantile of the conditional distribution of Y given $\mathbf{X} = \chi$ i.e. the solution of $F(\xi_{\alpha}(y|\mathbf{x})) = \alpha$ with respect to y . A natural estimator for the CPI is $\left(\hat{\xi}_{\frac{1-\gamma}{2}}(\mathbf{x}), \hat{\xi}_{\frac{1+\gamma}{2}}(\mathbf{x})\right)$, where the estimated quantiles satisfy $\hat{F}(\hat{\xi}_{\alpha}(\mathbf{x})|\mathbf{x}) = \alpha$.

The SCMI with nominal coverage probability γ is $(a, b) = \arg \min_{(c,d)} \{d - c \mid F(d|\mathbf{x}) - F(c|\mathbf{x}) \geq \gamma\}$, and a natural estimator for the SCMI is $(\hat{a}, \hat{b}) = \arg \min_{(c,d)} \left\{d - c \mid \hat{F}(d|\mathbf{x}) - \hat{F}(c|\mathbf{x}) \geq \gamma\right\}$, where, as before, the estimated conditional CDF is given by $\hat{F}(y|\chi) = \sum_{s=1}^k W(\mathbf{x}_s, \mathbf{x}) \mathbf{1}_{\{y_s \leq y\}}$. The SCMI is particularly suitable when the predictive density is asymmetric. For symmetric and unimodal distributions SCMI reduces to CPI.

3.5 Practical issues

For general NP-FDA prediction R/S+-routines are available at the website: <http://www.lsp.upstlse.fr/staph/npfda>; see also Ferraty and Vieu (2006, Chapter 7) for some details. We modified these routines for our purpose. The resulting R-codes, the datasets, and a brief description can be obtained from the authors. Two relatively “simple” practical aspects, concern the choice of the kernel function and the associated bandwidth. Throughout the analysis we employed the quadratic kernel: $K(u) = \frac{1}{2}(1 - u^2)_{[0,1]}(u)$. The bandwidth choice follows the data-driven procedure as described in Ferraty *et al.* (2005, Section 4.4), i.e. h is chosen in order to minimize $\sum_{s=1}^k \left| \hat{y}_s^{(\cdot)} - y_s^{(\cdot)} \right|$, where $\hat{y}_s^{(\cdot)}$ denotes the value of a predictor based on one of the FDA methods discussed above.

FPCA builds upon ideas from classical PCA. In fact, assuming $\mathbb{E}(\int \chi^2(t)dt) < \infty$, it can be shown that the functional random variable χ can be written as $\chi = \sum_{k=1}^{\infty} (\int \chi_s(t)e_k(t)dt)e_k$, where e_k are orthonormal eigenfunctions of the covariance operator $\Gamma_{\chi}(t, t') = \mathbb{E}(\chi(t)\chi(t'))$ associated with the eigenvalues $\lambda_1 \geq \lambda_2 \geq \dots \geq 0$. A truncated version of this expansion forms the basis of the FPCA semi-metric. In particular, the empirical version of this semi-metric is defined, in the case $\mathbf{x}$ is an observed pattern of a single variable consisting of T consecutive observations, as

$$\mathbf{d}_q(\mathbf{x}_s, \mathbf{x}) = \sqrt{\sum_{k=1}^q \left(\sum_{j=1}^T (\mathbf{x}_s(j) - \mathbf{x}(j))[\hat{e}_k]_j \right)^2}, \quad (4)$$

where, q denotes the number of retained principal components in the FPCA expansion, with q much smaller than T . It is straightforward to generalize (4) into a semi-metric suitable for multiple covariates. In that case, the parameter q need to be chosen. Comparing the prediction performance using the evaluation measures introduced in the next section, we noticed that values of $q \geq 6$ did not alter the results. Hence, we fixed q at 5 for each market.

Another practical issue is that price levels may differ considerably across trading days. To obtain price patterns that are comparable across trading days all price patterns in $\mathbf{x}$ are expressed relative to the opening price of that day.

4 Prediction Evaluation Measures

Four prediction evaluation criteria will be adopted. The first two criteria are measures for evaluating point predictions, while the latter two are concerned with evaluation of the PIs. The first measure is the mean-squared prediction error, given by $\text{MSE} = |\mathcal{S}|^{-1} \sum_{s \in \mathcal{S}} (\hat{y}_s^{(\cdot)} - y_s)^2$, where $\hat{y}_s^{(\cdot)}$ denotes the value of a predictor based on the respective NP-FDA, P-FDA approaches discussed above, or based on predictions obtained from the benchmark linear (multivariate) regression model.

From a practical point of view there often is some interest in predicting the sign of a return on the index rather than its precise value. For instance, many trading strategies are based on sign predictions. The following measure evaluates the point predictions by comparing the predicted and realized signs of the close-to-open gaps: $d^{\text{sgn}} = |\mathcal{S}|^{-1} \sum_{s \in \mathcal{S}} \mathbf{1}_{\{\text{sgn}(\hat{y}_s^{(\cdot)}) \neq \text{sgn}(y_s)\}}$. This measures the fraction of cases where the sign of the predicted close-to-open return and the actual return differ. Up to a constant factor, this is a generalization of the mean-squared prediction error applied to the signs of $\hat{y}_{i,s}^{(\cdot)}$ and $y_{i,s}$ rather than the actual values.

The remaining two criteria are used to evaluate the PIs. We compute the average width of PIs, as well as the empirical coverage probability. Let $\hat{\ell}_i$ and $\hat{u}_i$ denote the estimated lower and upper interval limits, respectively, obtained with one of the specific methods described in section 3.4 (CPI or SCMI). The (root-mean-squared) average width of the corresponding PIs is calculated as $v = \sqrt{|\mathcal{S}|^{-1} \sum_{s \in \mathcal{S}} (\hat{u} - \hat{\ell})^2}$. The empirical coverage probability is computed as $p_c = |\mathcal{S}|^{-1} \sum_{s \in \mathcal{S}} \mathbf{1}_{\{y_s \in (\hat{\ell}, \hat{u})\}}$. Ideally, a PI has coverage probability equal to the nominal coverage probability, while having the smallest possible average width. As an overall measure of the capability of the intervals to ‘capture much probability’ while having a small width, we also calculate the average PI length divided by the average coverage probability, $q = \bar{v}/\bar{p}_c$.

5 Data

The data consist of intra-day quotations of the following nine ($M = 9$) major stock market indices: the All Ordinary Composite Stock Index (AU), the Nikkei 225 Stock Index (JP), the Hang Seng Stock Index (HK), the FTSE 100 Share Index (UK), the Frankfurt DAX 30 Composite Stock Index (DE), the CAC 40 Composite Stock Index (FR), the Zurich Swiss Market Composite Index (CH), the Dow Jones Industrial Average (US), and the Toronto 300 Composite Stock Index (CA). All indices are retrieved from the Bloomberg databank. The period covered is from 24th September 2007 to 8th May 2008. Bloomberg offers one, five, and 15-minute quotations. In the case of one-minute quotations the market microstructure noise is more pronounced. Hence, we decided to use five-minute quotes.

For each of the y -variables, i.e. a close-to-open gap of one of the 9 stock indices, we consider various specifications, differing in terms of the information included in $\mathbf{x}$. To limit the number of possible specifications, information is added to the $\mathbf{x}$ -variable cluster-wise, where the three global clusters are the Asian cluster (JP, AU, HK), the European cluster (DE, FR, CH, UK), and the North-American cluster (US, CA). The first specification only contains the stock index patterns from the ‘previous’ cluster, for instance using the Asian cluster to predict the opening gap of the CAC 40. This specification is referred to as ‘Cluster(-1)’. The second specification only uses information from the before-last cluster. For instance, using the North-American cluster to predict the opening gap of the CAC 40. This is denoted by ‘Cluster(-2)’. Finally, specification ‘Cluster(-1)–Cluster(-2)’ contains the patterns from the last two clusters in the $\mathbf{x}$ -variable.

For specifications with one explanatory functional variable the dataset is organized in the form of a matrix. The predictions are based on the explanatory functional variable $(x_{i,s}(t)/x_{i,s}(1)) \times$

100 which resulted in MSEs that were at least twice as small as for three alternative transformations.

Table A.1, provided in Appendix A, provides an overview of the respective trading times (expressed in CET). In addition, Table A.1 shows information on the total number T_i of five-minute quotes per trading session when predictions are based on one explanatory functional variable. In the case of two- or more explanatory functional variables, the total number of five-minute quotes varies with specifications and trading times. To save space, we have not included this information in the paper. Further, note that Table A.1 contains the total number k_i of five-minute quotes (in parentheses) in the learning sample. The testing sample for each specification contains 35 days (curves). The complete dataset was prepared with great care, taking into account national holidays in all markets by considering overnight returns.

6 Results

6.1 MSE and sign

Table 1 shows the out-of-sample MSEs observed for each of these specifications, as well as their standard errors. For each Y -variable, the specifications that performed best in terms of the MSE criterion is indicated by an underlined entry. To facilitate interpretation of the results, the Table also provides aggregate MSE values, consisting of averages for similar specifications across the various y -variables. Global aggregate MSE values are provided, as well as individual aggregates for the three ‘clusters’ Asia, Europe and North-America. The standard errors of the aggregates are calculated from the standard errors of the individual MSEs, where these were assumed to be uncorrelated. The idea behind presenting aggregate MSEs across certain groups of $(\mathbf{x}, y)$ -pairs is that they provide a measure of a method’s average accuracy across $(\mathbf{x}, y)$ -pairs randomly selected from that group.

To interpret the results, it is convenient to start at the aggregate level and then look for particular differences between clusters and individual stock indices. At the overall aggregate level, the best specification in terms of MSE turned out to be the mean-based NP-FDA method using information from both other clusters (i.e. all information that has become available overnight). Notably, the average MSE (0.47) observed for that specification is considerably smaller than the average MSE values observed for the linear model specifications, which suggests the presence of a nonlinear relation between $\mathbf{x}$ and y . For the three NP-FDA methods one can observe that the specification using only information from the oldest cluster performs worse than using that from

the most recent cluster. The fact that this can be improved upon again by considering both clusters suggests that the information revealed by the markets that were open most recently is not reflecting all available information regarding the opening gap.

The observed pattern at the European aggregate level coincides with that just described at the global aggregate level. Deviations become apparent when looking at the aggregate results for North-America. Although the mean-based NP-FDA is again performing very well, the linear model is performing just as well, but based on a different information set (information from Europe only, rather than from Europe and Asia). A possible interpretation might be that although there is extra information in the patterns of the European markets that could be exploited for prediction, the linear model is more parsimonious and therefore able to achieve equal out-of-sample performance for the small dataset considered here. The aggregate NP-FDA and P-FDA results for Asia suggest that the opening gaps in the Asian markets are determined by the North-American stock index patterns as well as the European. When a linear specification is used, the best model seems to suggest that Asia is only affected by the patterns in the European markets, which would be highly counterintuitive. The presence of nonlinear dependence of the Asian opening gaps on the observed patterns may explain this. Indeed, the mean-based NP-FDA and the P-FDA method achieve smaller out-of-sample MSEs based on trading patterns in both European and North-American markets.

The results for the individual European markets show structure that roughly coincides with that of the (global as well as European) aggregate. The best model is the mean-based NP-FDA, except for Germany, for which P-FDA performs slightly, but insignificantly, better. The best-performing NP-FDA specification is that using information from both clusters, while the linear model performs best with a parsimonious specification, based on information from the latest available cluster only. The best performing model for the US opening gap is the linear regression model, based on the information revealed by the European markets overnight, in line with what one would expect for informationally efficient markets. A similar result holds for Canada, although in that case the mean-based NP-FDA performed slightly (very insignificantly) better. The opening in Japan seems to be affected by the North-American markets only, both in terms of the linear benchmark and the NP-FDA. The P-FDA results might indicate that Europe also has some effect on Japan, but this is insignificant. The best performing model for Australia is the linear model based on information from Europe. However, the observed MSEs for several of the other specifications are almost as small, and well within the standard error. Likewise, for Hong Kong one of the linear models, one of the NP-FDA and one of the P-FDA methods perform

practically equally well. Note that for none of the y -variables the median-based NP-FDA or the mode-based NP-FDA method performed best in terms of the MSE criterion.

Interestingly, this picture changes rather substantially if we consider the other performance measure, d^{sgn} , the results of which are given in Table 2. The globally aggregated results show that the linear model based on information from the last cluster performs best on average, closely (with an insignificant difference) followed by the mode-based NP-FDA using information of the before-last cluster.

The close-to-open gap in the European stock indices is mainly determined by the patterns in the North-American markets, and the mode-based NP-FDA method picks up this structure best. Across the different prediction methods, the sign of the opening gap of the North-American indices appears to be determined by the Asian as well as the European patterns, although the prediction method that performed best (mode-based NP-FDA) did so using the Asian stock index patterns only. The Asian aggregate results show that both the linear model and the mean-based NP-FDA perform well, using patterns from the North-American indices only.

The results for the individual indices roughly follow the structure already reflected by the aggregate results. An exception is the CAC 40, as it is the only European index for which the sign of the opening gap is determined by the Asian patterns only. For the other European indices the sign is determined by the North-American index patterns.

6.2 CPI and SCMI

Table 3 shows the results obtained for the CPI and the SCMI predictive intervals. For ease of presentation only the aggregate results are provided, which closely coincide with the individual results. It can be observed that all coverage probabilities are smaller than the nominal value of 90%. In all cases the coverage probabilities of CPI are better in the sense that they are closer to the nominal value. On the other hand, on average the SCMIs are shorter than the CPIs. This indicates that the CPI is more sensitive to the position in the state-space from which predictions are being made than the SCMI. The overall quality measure q corresponding with the ratio of the average length and the average coverage probabilities are very similar for both types of intervals.

7 Summary and Conclusion

The aggregate results for the MSE show that the best FDA specification, mean-based NP-FDA with Cluster(-1)-Cluster(-2), on average performs much better than any of the linear models. This suggests that the NP-FDA method successfully exploits nonlinearities in the relation between $\mathbf{x}$ and y . This result is in line with the huge empirical evidence for nonlinear dependence in daily stock returns. In a recent systematic model-based prediction exercise, Guidolin *et al.* (2009), found that stock and bond returns from the G7 countries, and in particular UK and US, appear to require nonlinear modelling.

The three clusters seem to be governed by different types of dynamics. When considering only the linear specification, The European and Asian markets appear to be informationally efficient in that the specification Cluster(-1) gives the smallest MSE among the linear models. For the US this is also the overall best performing specification, supporting the idea that the US index is informationally efficient. However, for all European markets the MSE obtained with the mean-based NP-FDA using the Cluster(-1)-Cluster(-2) specification was substantially smaller than those obtained with the best linear model. This suggests that, although impossible to see using linear models, the European markets are not informationally efficient after all; the best predictor is nonparametric and rather than information from the latest cluster only, it uses information from the last two clusters. The best specification for JP is, as might be expected, based on Cluster(-1) only. However, also there a substantial improvement in the MSE is obtained in going from the linear to the mean-based NP-FDA specification, suggesting the presence of a nonlinear relation between the North-American price patterns and JP. Further, the results have shown that in none of the cases P-FDA outperforms NP-FDA. This holds for the MSE as well as for the d^{sgn} measure. Among the NP-FDA methods considered, the best MSEs were obtained with the mean-based NP-FDA, while the mode-based NP-FDA performed best in terms of d^{sgn} in many cases.

Finally, as far as we are aware, exploring full information in the intra-day stock price patterns in foreign markets to predict the opening of an index in a home-market, using NP-FDA, has not been a topic of earlier research. Clearly the present study recognizes the fact that traders in a home market use any information revealed overnight in a foreign market due to fast transmission of information worldwide. Hence, our approach is closer to the underlying process of information flows than studies based on daily returns, volatility of returns, or closing prices.

References

- Benko, M. (2006), *Functional Data Analysis with Applications in Finance*, PhD Humboldt-Universität, Berlin.
- Cliff, M., Cooper, M.J. and Gulen, H. (2008), “Return differences between trading and non-trading hours: Like night and day”. Working paper, <http://ssrn.com/abstract=1004081>.
- De Gooijer, J.G. and Gannoun, A. (2000), “Nonparametric conditional predictive regions for time series”, *Computational Statistics & Data Analysis*, 33, 359–275.
- Ferraty, F. and Vieu, P. (2005), “Conditional quantiles for dependent functional data with application to the climatic *El Niño* phenomenon”, *Sankhyā: The Indian Journal of Statistics*, 67(2), 378–398.
- Ferraty, F. and Vieu, P. (2006), *Nonparametric Functional Data Analysis*, Springer-Verlag, New York.
- French, K.R. and Roll, R. (1986), “Stock return variances: The arrival of information and the reaction of traders”, *Journal of Financial Economics*, 17, 5–26.
- George, T.J. and Hwang, C.Y. (2001), “Information flow and pricing errors: A unified approach to estimation and testing”, *Review of Financial Studies*, 14, 979–1020.
- Guidolin, M., Hyde, S., McMillan, D. and Ono, S. (2009), “Non-linear predictability in stock and bond returns: When and where is it exploitable?”, *International Journal of Forecasting*, 25, 373–399.
- Hasbrouck, J. (1991), “Measuring the information content of stock trades”, *Journal of Finance*, 46, 179–207.
- Hasbrouck, J. (1993), “Assessing the quality of a security market: A new approach to transaction-cost measurement”, *Review of Financial Studies*, 6, 191–212.
- Hong, H. and Wang, J. (2000), “Trading and return under periodic market closures”, *Journal of Finance*, 55, 297–354.
- Lockwood, L.J. and Lin, S. (1965), “An examination of stock market return volatility during overnight and intraday periods, 1964-1989”, *Journal of Finance*, 45, 591–601.
- Madhavan, A., Richardson, M. and Roomans, M. (1997), “Why do security prices change? A transaction level analysis of NYSE stocks”, *Review of Financial Studies*, 10, 1035–1064.
- Oldfield, G.S. and Rogalski, R.J. (1980), “A theory of common stock returns over trading and non-trading periods”, *Journal of Finance*, 35, 729–751.
- Ramsay, J.O. and Silverman, B.W. (2005), *Functional Data Analysis* (2nd ed.) Springer Verlag,

New York.

Zhong, Z. (2007), “The predictability of overnight information”, M.Sc. thesis, Singapore Management University.

Table 1: MSEs with standard deviations in parentheses.

x_s	y_s	Mean	Median	Mode	PFDA	Lin. reg
HK-AU-JP	FR	0.60 (0.18)	0.79 (0.19)	0.68 (0.22)	0.75 (0.12)	0.67 (0.25)
US-CA		0.71 (0.25)	0.97 (0.27)	0.71 (0.25)	0.85 (0.14)	0.73 (0.23)
HK-AU-JP-US-CA		<u>0.44</u> (0.11)	0.65 (0.22)	0.64 (0.22)	0.61 (0.09)	0.75 (0.22)
HK-AU-JP	UK	0.44 (0.11)	0.39 (0.08)	0.44 (0.11)	0.44 (0.10)	0.43 (0.12)
US-CA		0.50 (0.16)	0.93 (0.29)	0.59 (0.19)	0.57 (0.09)	0.49 (0.14)
HK-AU-JP-US-CA		<u>0.33</u> (0.08)	0.47 (0.11)	0.44 (0.11)	0.31 (0.07)	0.51 (0.14)
HK-AU-JP	CH	0.63 (0.21)	0.72 (0.27)	0.73 (0.27)	0.79 (0.11)	0.72 (0.30)
US-CA		0.68 (0.29)	0.67 (0.25)	0.75 (0.29)	0.74 (0.10)	0.73 (0.30)
HK-AU-JP-US-CA		<u>0.45</u> (0.18)	0.84 (0.30)	0.57 (0.27)	0.58 (0.10)	0.71 (0.30)
HK-AU-JP	DE	0.54 (0.19)	0.45 (0.12)	0.59 (0.18)	0.59 (0.12)	0.56 (0.17)
US-CA		0.49 (0.14)	0.51 (0.13)	0.57 (0.16)	0.58 (0.10)	0.56 (0.13)
HK-AU-JP-US-CA		0.33 (0.19)	0.46 (0.12)	0.47 (0.18)	<u>0.26</u> (0.12)	0.56 (0.17)
DE-FR-CH-UK	US	0.39 (0.07)	0.83 (0.20)	0.40 (0.08)	0.48 (0.07)	<u>0.35</u> (0.09)
HK-AU-JP		0.56 (0.15)	1.05 (0.26)	0.53 (0.13)	0.57 (0.07)	0.51 (0.13)
HK-AU-JP-DE-FR-UK-CH		0.37 (0.09)	0.48 (0.10)	0.42 (0.09)	0.40 (0.08)	0.57 (0.16)
DE-FR-CH-UK	CA	0.50 (0.15)	0.71 (0.19)	0.49 (0.19)	0.57 (0.08)	0.42 (0.14)
HK-AU-JP		0.55 (0.20)	1.18 (0.29)	0.53 (0.18)	0.49 (0.07)	0.47 (0.16)
HK-AU-JP-DE-FR-UK-CH		<u>0.41</u> (0.17)	0.56 (0.14)	0.47 (0.17)	0.55 (0.11)	0.57 (0.18)
US-CA	JP	<u>0.35</u> (0.08)	0.63 (0.15)	0.44 (0.10)	0.45 (0.10)	0.50 (0.09)
DE-FR-CH-UK		0.67 (0.10)	1.07 (0.22)	0.70 (0.12)	0.57 (0.14)	0.56 (0.08)
DE-FR-CH-UK-US-CA		0.50 (0.08)	0.63 (0.12)	0.65 (0.13)	0.41 (0.14)	0.63 (0.12)
US-CA	AU	0.32 (0.10)	0.42 (0.16)	0.41 (0.13)	0.33 (0.07)	0.28 (0.08)
DE-FR-CH-UK		0.37 (0.08)	0.52 (0.11)	0.39 (0.09)	0.39 (0.12)	<u>0.26</u> (0.07)
DE-FR-CH-UK-US-CA		0.29 (0.08)	0.28 (0.07)	0.28 (0.07)	0.35 (0.07)	0.32 (0.08)
US-CA	HK	1.34 (0.26)	2.28 (0.49)	1.70 (0.29)	2.17 (0.93)	1.83 (0.42)
DE-FR-CH-UK		1.76 (0.36)	3.20 (0.70)	2.25 (0.56)	1.31 (0.95)	1.20 (0.36)
DE-FR-CH-UK-US-CA		1.14 (0.26)	2.68 (0.52)	1.95 (0.41)	<u>1.07</u> (1.03)	2.10 (0.45)
Overall aggregate						
Cluster(-1)		0.57 (0.05)	0.80 (0.08)	0.65 (0.06)	0.73 (0.11)	0.64 (0.07)
Cluster(-2)		0.70 (0.07)	1.12 (0.11)	0.78 (0.09)	0.67 (0.11)	0.61 (0.07)
Cluster(-1)-cluster(-2)		<u>0.47</u> (0.05)	0.78 (0.08)	0.65 (0.07)	0.50 (0.12)	0.75 (0.08)
Europe						
Cluster(-1)		0.55 (0.09)	0.59 (0.09)	0.61 (0.10)	0.64 (0.06)	0.60 (0.11)
Cluster(-2)		0.60 (0.11)	0.77 (0.12)	0.66 (0.11)	0.69 (0.05)	0.63 (0.11)
Cluster(-1)-cluster(-2)		<u>0.39</u> (0.07)	0.61 (0.10)	0.53 (0.10)	0.44 (0.05)	0.63 (0.11)
North-America						
Cluster(-1)		0.45 (0.08)	0.77 (0.14)	0.45 (0.10)	0.53 (0.05)	<u>0.39</u> (0.08)
Cluster(-2)		0.56 (0.13)	1.12 (0.19)	0.53 (0.11)	0.53 (0.05)	0.49 (0.10)
Cluster(-1)-cluster(-2)		<u>0.39</u> (0.10)	0.52 (0.09)	0.45 (0.10)	0.48 (0.07)	0.57 (0.12)
Asia						
Cluster(-1)		0.67 (0.10)	1.11 (0.18)	0.85 (0.11)	0.98 (0.31)	0.87 (0.15)
Cluster(-2)		0.93 (0.13)	1.60 (0.25)	1.11 (0.19)	0.76 (0.32)	0.67 (0.13)
Cluster(-1)-cluster(-2)		0.64 (0.09)	1.20 (0.18)	0.96 (0.15)	<u>0.61</u> (0.35)	1.02 (0.16)

Table 2: d^{sgn} with standard deviations in parentheses.

$\mathbf{x}_s$	y_s	Mean	Median	Mode	PFDA	Lin. reg
HK-AU-JP	FR	<u>0.17</u> (0.06)	0.29 (0.08)	0.54 (0.08)	0.34 (0.08)	0.17 (0.06)
US-CA		0.31 (0.08)	0.34 (0.08)	0.26 (0.07)	0.31 (0.08)	0.31 (0.08)
HK-AU-JP-US-CA		0.29 (0.08)	0.26 (0.07)	0.34 (0.08)	0.40 (0.08)	0.37 (0.08)
HK-AU-JP	UK	0.23 (0.07)	0.23 (0.07)	0.49 (0.08)	0.29 (0.08)	0.20 (0.07)
US-CA		0.37 (0.08)	0.37 (0.08)	<u>0.06</u> (0.04)	0.40 (0.08)	0.34 (0.08)
HK-AU-JP-US-CA		0.29 (0.08)	0.17 (0.06)	0.31 (0.08)	0.43 (0.08)	0.31 (0.08)
HK-AU-JP	CH	0.23 (0.07)	0.31 (0.08)	0.57 (0.08)	0.34 (0.08)	<u>0.20</u> (0.07)
US-CA		0.34 (0.08)	0.37 (0.08)	<u>0.20</u> (0.07)	0.31 (0.08)	0.37 (0.08)
HK-AU-JP-US-CA		0.34 (0.08)	0.31 (0.08)	0.34 (0.08)	0.37 (0.08)	0.46 (0.08)
HK-AU-JP	DE	0.23 (0.07)	0.17 (0.06)	0.54 (0.08)	0.29 (0.08)	0.20 (0.07)
US-CA		0.23 (0.07)	0.34 (0.08)	<u>0.09</u> (0.05)	0.31 (0.08)	0.34 (0.08)
HK-AU-JP-US-CA		0.23 (0.07)	0.17 (0.06)	0.31 (0.08)	0.29 (0.08)	0.34 (0.08)
DE-FR-CH-UK	US	0.51 (0.08)	0.43 (0.08)	0.51 (0.08)	0.51 (0.08)	0.49 (0.08)
HK-AU-JP		0.37 (0.08)	0.46 (0.08)	<u>0.29</u> (0.08)	0.43 (0.08)	0.40 (0.08)
HK-AU-JP-DE-FR-UK-CH		<u>0.29</u> (0.08)	0.34 (0.08)	0.54 (0.08)	0.34 (0.08)	0.46 (0.08)
DE-FR-CH-UK	CA	0.34 (0.08)	0.34 (0.08)	0.31 (0.08)	0.37 (0.08)	0.29 (0.08)
HK-AU-JP		0.49 (0.08)	0.43 (0.08)	0.23 (0.07)	0.40 (0.08)	0.37 (0.08)
HK-AU-JP-DE-FR-UK-CH		0.29 (0.08)	0.40 (0.08)	<u>0.11</u> (0.05)	0.34 (0.08)	0.46 (0.08)
US-CA	JP	<u>0.20</u> (0.07)	0.26 (0.07)	0.29 (0.08)	0.20 (0.07)	0.23 (0.07)
DE-FR-CH-UK		0.51 (0.08)	0.51 (0.08)	0.54 (0.08)	0.37 (0.08)	0.37 (0.08)
DE-FR-CH-UK-US-CA		0.43 (0.08)	0.37 (0.08)	0.57 (0.08)	0.29 (0.08)	0.31 (0.08)
US-CA	AU	<u>0.26</u> (0.07)	0.37 (0.08)	0.34 (0.08)	0.29 (0.08)	0.31 (0.08)
DE-FR-CH-UK		0.34 (0.08)	0.31 (0.08)	0.43 (0.08)	0.43 (0.08)	0.37 (0.08)
DE-FR-CH-UK-US-CA		0.29 (0.08)	0.34 (0.08)	0.43 (0.08)	0.43 (0.08)	0.40 (0.08)
US-CA	HK	<u>0.29</u> (0.08)	0.31 (0.08)	0.31 (0.08)	0.29 (0.08)	0.26 (0.07)
DE-FR-CH-UK		0.29 (0.08)	0.40 (0.08)	0.46 (0.08)	0.34 (0.08)	0.26 (0.07)
DE-FR-CH-UK-US-CA		0.31 (0.08)	0.37 (0.08)	0.40 (0.08)	0.31 (0.08)	0.34 (0.08)
Overall aggregate						
Cluster(-1)		0.27 (0.02)	0.31 (0.03)	0.40 (0.03)	0.32 (0.03)	<u>0.26</u> (0.02)
Cluster(-2)		0.36 (0.03)	0.39 (0.03)	0.28 (0.02)	0.37 (0.03)	0.35 (0.03)
Cluster(-1)-cluster(-2)		0.31 (0.03)	0.30 (0.03)	0.37 (0.03)	0.36 (0.03)	0.38 (0.03)
Europe						
Cluster(-1)		0.22 (0.03)	0.25 (0.04)	0.54 (0.04)	0.32 (0.04)	0.19 (0.03)
Cluster(-2)		0.31 (0.04)	0.36 (0.04)	<u>0.15</u> (0.03)	0.33 (0.04)	0.34 (0.04)
Cluster(-1)-cluster(-2)		0.29 (0.04)	0.23 (0.04)	0.33 (0.04)	0.37 (0.04)	0.37 (0.04)
North-America						
Cluster(-1)		0.43 (0.06)	0.39 (0.06)	0.41 (0.06)	0.44 (0.06)	0.39 (0.06)
Cluster(-2)		0.43 (0.06)	0.45 (0.06)	<u>0.26</u> (0.05)	0.42 (0.06)	0.39 (0.06)
Cluster(-1)-cluster(-2)		0.29 (0.05)	0.37 (0.06)	0.33 (0.05)	0.34 (0.06)	0.46 (0.06)
Asia						
Cluster(-1)		<u>0.25</u> (0.04)	0.31 (0.05)	0.31 (0.05)	0.26 (0.04)	0.27 (0.04)
Cluster(-2)		0.38 (0.05)	0.41 (0.05)	0.48 (0.05)	0.38 (0.05)	0.33 (0.05)
Cluster(-1)-cluster(-2)		0.34 (0.05)	0.36 (0.05)	0.47 (0.05)	0.34 (0.05)	0.35 (0.05)

Table 3: Coverage probabilities, predictive interval widths, and overall predictive interval quality measure q . Nominal coverage is 0.90.

	Cov. prob.		Width		q	
	CPI	SCMI	CPI	SCMI	CPI	SCMI
Overall aggregate						
Cluster(-1)	0.76	0.72	2.02	1.87	2.65	2.58
Cluster(-2)	0.72	0.71	1.89	1.79	2.62	2.52
Cluster(-1)-cluster(-2)	0.67	0.63	1.57	1.47	<u>2.35</u>	<u>2.35</u>
Europe						
Cluster(-1)	0.80	0.73	1.96	1.75	2.45	2.39
Cluster(-2)	0.67	0.69	1.50	1.43	2.25	2.07
Cluster(-1)-cluster(-2)	0.67	0.65	1.31	1.27	1.96	<u>1.95</u>
North-America						
Cluster(-1)	0.74	0.76	1.96	1.83	2.66	2.43
Cluster(-2)	0.83	0.80	2.03	1.92	2.45	2.40
Cluster(-1)-cluster(-2)	0.83	0.77	1.68	1.56	<u>2.02</u>	2.03
Asia						
Cluster(-1)	0.73	0.69	2.13	2.03	<u>2.94</u>	2.96
Cluster(-2)	0.73	0.68	2.33	2.18	3.21	3.23
Cluster(-1)-cluster(-2)	0.73	0.50	1.82	1.82	3.30	3.34

A Appendix

Table A.1: Trading times expressed in CET. Total number T_i of 5-minute quotes per trading session s , when predictions are based on a single explanatory variable, and (in parentheses) the total number k_i of 5-minute quotes in the learning sample.

$\mathbf{x}_{i,s}$	CET	$y_{i,s}$								
		AU	JP	HK	UK	DE	FR	CH	US	CA
AU	00:00-06:05				74 (118)	74 (116)	74 (118)	74 (115)	74 (115)	74 (117)
JP	01:00-06:35				57 (110)	57 (109)	57 (109)	57 (110)	57 (108)	57 (110)
HK	03:00-09:00				50 (114)	50 (113)	50 (115)	50 (113)	51 (111)	51 (113)
UK	09:00-17:30	103 (86)	103 (84)	103 (82)					66 (117)	66 (119)
DE	09:00-17:35	104 (86)	104 (83)	104 (82)					66 (115)	66 (117)
FR	09:00-17:25	102 (86)	102 (84)	102 (82)					66 (117)	66 (119)
CH	09:00-17:30	101 (85)	101 (83)	101 (81)					66 (114)	66 (116)
US	14:30-21:00	79 (85)	79 (84)	79 (81)	79 (86)	79 (85)	79 (85)	79 (85)		
CA	14:30-21:05	80 (86)	80 (84)	80 (82)	80 (87)	80 (86)	80 (86)	80 (87)		