

Published in final edited form as:

*Nat Biotechnol.* 2010 August ; 28(8): 827–838. doi:10.1038/nbt.1665.

# The MicroArray Quality Control (MAQC)-II study of common practices for the development and validation of microarray-based predictive models

MAQC Consortium\*

## Abstract

Gene expression data from microarrays are being applied to predict preclinical and clinical endpoints, but the reliability of these predictions has not been established. In the MAQC-II project, 36 independent teams analyzed six microarray data sets to generate predictive models for classifying a sample with respect to one of 13 endpoints indicative of lung or liver toxicity in rodents, or of breast cancer, multiple myeloma or neuroblastoma in humans. In total, >30,000 models were built using many combinations of analytical methods. The teams generated predictive models without knowing the biological meaning of some of the endpoints and, to mimic clinical reality, tested the models on data that had not been used for training. We found that model performance depended largely on the endpoint and team proficiency and that different approaches generated models of similar performance. The conclusions and recommendations from MAQC-II should be useful for regulatory agencies, study committees and independent investigators that evaluate methods for global gene expression analysis.

As part of the United States Food and Drug Administration's (FDA's) Critical Path Initiative to medical product development (<http://www.fda.gov/oc/initiatives/criticalpath/>), the MAQC consortium began in February 2005 with the goal of addressing various microarray reliability concerns raised in publications<sup>1–9</sup> pertaining to reproducibility of gene signatures. The first phase of this project (MAQC-I) extensively evaluated the technical performance of microarray platforms in identifying all differentially expressed genes that would potentially constitute biomarkers. The MAQC-I found high intra-platform reproducibility across test sites, as well as inter-platform concordance of differentially expressed gene lists<sup>10–15</sup> and confirmed that microarray technology is able to reliably identify differentially expressed

© 2010 Nature America, Inc. All rights reserved.

Correspondence should be addressed to L.S. ([leming.shi@fda.hhs.gov](mailto:leming.shi@fda.hhs.gov) or [leming.shi@gmail.com](mailto:leming.shi@gmail.com)).

\*A full list of authors and affiliations appears at the end of the paper.

**Accession codes.** All MAQC-II data sets are available through GEO (series accession number: GSE16716), the MAQC Web site (<http://www.fda.gov/nctr/science/centers/toxicoinformatics/maqc/>), ArrayTrack (<http://www.fda.gov/nctr/science/centers/toxicoinformatics/ArrayTrack/>) or CEBS (<http://cebs.niehs.nih.gov/>) accession number: 009-00002-0010-000-3.

Supplementary information is available on the Nature Biotechnology website.

### Publisher's Disclaimer: DISCLAIMER

This work includes contributions from, and was reviewed by, individuals at the FDA, the Environmental Protection Agency (EPA) and the NIH. This work has been approved for publication by these agencies, but it does not necessarily reflect official agency policy. Certain commercial materials and equipment are identified in order to adequately specify experimental procedures. In no case does such identification imply recommendation or endorsement by the FDA, the EPA or the NIH, nor does it imply that the items identified are necessarily the best available for the purpose.

### COMPETING FINANCIAL INTERESTS

The authors declare competing financial interests: details accompany the full-text HTML version of the paper at <http://www.nature.com/naturebiotechnology/>.

Reprints and permissions information is available online at <http://npg.nature.com/reprintsandpermissions/>.

genes between sample classes or populations<sup>16,17</sup>. Importantly, the MAQC-I helped produce companion guidance regarding genomic data submission to the FDA (<http://www.fda.gov/downloads/Drugs/GuidanceComplianceRegulatoryInformation/Guidances/ucm079855.pdf> ).

Although the MAQC-I focused on the technical aspects of gene expression measurements, robust technology platforms alone are not sufficient to fully realize the promise of this technology. An additional requirement is the development of accurate and reproducible multivariate gene expression-based prediction models, also referred to as classifiers. Such models take gene expression data from a patient as input and as output produce a prediction of a clinically relevant outcome for that patient. Therefore, the second phase of the project (MAQC-II) has focused on these predictive models<sup>18</sup>, studying both how they are developed and how they are evaluated. For any given microarray data set, many computational approaches can be followed to develop predictive models and to estimate the future performance of these models. Understanding the strengths and limitations of these various approaches is critical to the formulation of guidelines for safe and effective use of preclinical and clinical genomic data. Although previous studies have compared and benchmarked individual steps in the model development process<sup>19</sup>, no prior published work has, to our knowledge, extensively evaluated current community practices on the development and validation of microarray-based predictive models.

Microarray-based gene expression data and prediction models are increasingly being submitted by the regulated industry to the FDA to support medical product development and testing applications<sup>20</sup>. For example, gene expression microarray-based assays that have been approved by the FDA as diagnostic tests include the Agendia MammaPrint microarray to assess prognosis of distant metastasis in breast cancer patients<sup>21,22</sup> and the Pathwork Tissue of Origin Test to assess the degree of similarity of the RNA expression pattern in a patient's tumor to that in a database of tumor samples for which the origin of the tumor is known<sup>23</sup>. Gene expression data have also been the basis for the development of PCR-based diagnostic assays, including the xDx Allomap test for detection of rejection of heart transplants<sup>24</sup>.

The possible uses of gene expression data are vast and include diagnosis, early detection (screening), monitoring of disease progression, risk assessment, prognosis, complex medical product characterization and prediction of response to treatment (with regard to safety or efficacy) with a drug or device labeling intent. The ability to generate models in a reproducible fashion is an important consideration in predictive model development.

A lack of consistency in generating classifiers from publicly available data is problematic and may be due to any number of factors including insufficient annotation, incomplete clinical identifiers, coding errors and/or inappropriate use of methodology<sup>25,26</sup>. There are also examples in the literature of classifiers whose performance cannot be reproduced on independent data sets because of poor study design<sup>27</sup>, poor data quality and/or insufficient cross-validation of all model development steps<sup>28,29</sup>. Each of these factors may contribute to a certain level of skepticism about claims of performance levels achieved by microarray-based classifiers.

Previous evaluations of the reproducibility of microarray-based classifiers, with only very few exceptions<sup>30,31</sup>, have been limited to simulation studies or reanalysis of previously published results. Frequently, published benchmarking studies have split data sets at random, and used one part for training and the other for validation. This design assumes that the training and validation sets are produced by unbiased sampling of a large, homogeneous population of samples. However, specimens in clinical studies are usually accrued over

years and there may be a shift in the participating patient population and also in the methods used to assign disease status owing to changing practice standards. There may also be batch effects owing to time variations in tissue analysis or due to distinct methods of sample collection and handling at different medical centers. As a result, samples derived from sequentially accrued patient populations, as was done in MAQC-II to mimic clinical reality, where the first cohort is used for developing predictive models and subsequent patients are included in validation, may differ from each other in many ways that could influence the prediction performance.

The MAQC-II project was designed to evaluate these sources of bias in study design by constructing training and validation sets at different times, swapping the test and training sets and also using data from diverse preclinical and clinical scenarios. The goals of MAQC-II were to survey approaches in genomic model development in an attempt to understand sources of variability in prediction performance and to assess the influences of endpoint signal strength in data. By providing the same data sets to many organizations for analysis, but not restricting their data analysis protocols, the project has made it possible to evaluate to what extent, if any, results depend on the team that performs the analysis. This contrasts with previous benchmarking studies that have typically been conducted by single laboratories. Enrolling a large number of organizations has also made it feasible to test many more approaches than would be practical for any single team. MAQC-II also strives to develop good modeling practice guidelines, drawing on a large international collaboration of experts and the lessons learned in the perhaps unprecedented effort of developing and evaluating >30,000 genomic classifiers to predict a variety of endpoints from diverse data sets.

MAQC-II is a collaborative research project that includes participants from the FDA, other government agencies, industry and academia. This paper describes the MAQC-II structure and experimental design and summarizes the main findings and key results of the consortium, whose members have learned a great deal during the process. The resulting guidelines are general and should not be construed as specific recommendations by the FDA for regulatory submissions.

## RESULTS

### Generating a unique compendium of >30,000 prediction models

The MAQC-II consortium was conceived with the primary goal of examining model development practices for generating binary classifiers in two types of data sets, preclinical and clinical (Supplementary Tables 1 and 2). To accomplish this, the project leader distributed six data sets containing 13 preclinical and clinical endpoints coded A through M (Table 1) to 36 voluntary participating data analysis teams representing academia, industry and government institutions (Supplementary Table 3). Endpoints were coded so as to hide the identities of two negative-control endpoints (endpoints I and M, for which class labels were randomly assigned and are not predictable by the microarray data) and two positive-control endpoints (endpoints H and L, representing the sex of patients, which is highly predictable by the microarray data). Endpoints A, B and C tested teams' ability to predict the toxicity of chemical agents in rodent lung and liver models. The remaining endpoints were predicted from microarray data sets from human patients diagnosed with breast cancer (D and E), multiple myeloma (F and G) or neuroblastoma (J and K). For the multiple myeloma and neuroblastoma data sets, the endpoints represented event free survival (abbreviated EFS), meaning a lack of malignancy or disease recurrence, and overall survival (abbreviated OS) after 730 days (for multiple myeloma) or 900 days (for neuroblastoma) post treatment or diagnosis. For breast cancer, the endpoints represented estrogen receptor status, a common diagnostic marker of this cancer type (abbreviated 'erpos'), and the success of

treatment involving chemotherapy followed by surgical resection of a tumor (abbreviated 'pCR'). The biological meaning of the control endpoints was known only to the project leader and not revealed to the project participants until all model development and external validation processes had been completed.

To evaluate the reproducibility of the models developed by a data analysis team for a given data set, we asked teams to submit models from two stages of analyses. In the first stage (hereafter referred to as the 'original' experiment), each team built prediction models for up to 13 different coded endpoints using six training data sets. Models were 'frozen' against further modification, submitted to the consortium and then tested on a blinded validation data set that was not available to the analysis teams during training. In the second stage (referred to as the 'swap' experiment), teams repeated the model building and validation process by training models on the original validation set and validating them using the original training set.

To simulate the potential decision-making process for evaluating a microarray-based classifier, we established a process for each group to receive training data with coded endpoints, propose a data analysis protocol (DAP) based on exploratory analysis, receive feedback on the protocol and then perform the analysis and validation (Fig. 1). Analysis protocols were reviewed internally by other MAQC-II participants (at least two reviewers per protocol) and by members of the MAQC-II Regulatory Biostatistics Working Group (RBWG), a team from the FDA and industry comprising biostatisticians and others with extensive model building expertise. Teams were encouraged to revise their protocols to incorporate feedback from reviewers, but each team was eventually considered responsible for its own analysis protocol and incorporating reviewers' feedback was not mandatory (see Online Methods for more details).

We assembled two large tables from the original and swap experiments (Supplementary Tables 1 and 2, respectively) containing summary information about the algorithms and analytic steps, or 'modeling factors', used to construct each model and the 'internal' and 'external' performance of each model. Internal performance measures the ability of the model to classify the training samples, based on cross-validation exercises. External performance measures the ability of the model to classify the blinded independent validation data. We considered several performance metrics, including Matthews Correlation Coefficient (MCC), accuracy, sensitivity, specificity, area under the receiver operating characteristic curve (AUC) and root mean squared error (r.m.s.e.). These two tables contain data on >30,000 models. Here we report performance based on MCC because it is informative when the distribution of the two classes in a data set is highly skewed and because it is simple to calculate and was available for all models. MCC values range from +1 to -1, with +1 indicating perfect prediction (that is, all samples classified correctly and none incorrectly), 0 indicates random prediction and -1 indicating perfect inverse prediction.

The 36 analysis teams applied many different options under each modeling factor for developing models (Supplementary Table 4) including 17 summary and normalization methods, nine batch-effect removal methods, 33 feature selection methods (between 1 and >1,000 features), 24 classification algorithms and six internal validation methods. Such diversity suggests the community's common practices are well represented. For each of the models nominated by a team as being the best model for a particular endpoint, we compiled the list of features used for both the original and swap experiments (see the MAQC Web site at <http://edkb.fda.gov/MAQC/>). These comprehensive tables represent a unique resource. The results that follow describe data mining efforts to determine the potential and limitations of current practices for developing and validating gene expression-based prediction models.

## Performance depends on endpoint and can be estimated during training

Unlike many previous efforts, the study design of MAQC-II provided the opportunity to assess the performance of many different modeling approaches on a clinically realistic blinded external validation data set. This is especially important in light of the intended clinical or preclinical uses of classifiers that are constructed using initial data sets and validated for regulatory approval and then are expected to accurately predict samples collected under diverse conditions perhaps months or years later. To assess the reliability of performance estimates derived during model training, we compared the performance on the internal training data set with performance on the external validation data set for each of the 18,060 models in the original experiment (Fig. 2a). Models without complete metadata were not included in the analysis.

We selected 13 ‘candidate models’, representing the best model for each endpoint, before external validation was performed. We required that each analysis team nominate one model for each endpoint they analyzed and we then selected one candidate from these nominations for each endpoint. We observed a higher correlation between internal and external performance estimates in terms of MCC for the selected candidate models ( $r = 0.951$ ,  $n = 13$ , Fig. 2b) than for the overall set of models ( $r = 0.840$ ,  $n = 18,060$ , Fig. 2a), suggesting that extensive peer review of analysis protocols was able to avoid selecting models that could result in less reliable predictions in external validation. Yet, even for the hand-selected candidate models, there is noticeable bias in the performance estimated from internal validation. That is, the internal validation performance is higher than the external validation performance for most endpoints (Fig. 2b). However, for some endpoints and for some model building methods or teams, internal and external performance correlations were more modest as described in the following sections.

To evaluate whether some endpoints might be more predictable than others and to calibrate performance against the positive- and negative-control endpoints, we assessed all models generated for each endpoint (Fig. 2c). We observed a clear dependence of prediction performance on endpoint. For example, endpoints C (liver necrosis score of rats treated with hepatotoxicants), E (estrogen receptor status of breast cancer patients), and H and L (sex of the multiple myeloma and neuroblastoma patients, respectively) were the easiest to predict (mean MCC  $> 0.7$ ). Toxicological endpoints A and B and disease progression endpoints D, F, G, J and K were more difficult to predict (mean MCC  $\sim 0.1$ – $0.4$ ). Negative-control endpoints I and M were totally unpredictable (mean MCC  $\sim 0$ ), as expected. For 11 endpoints (excluding the negative controls), a large proportion of the submitted models predicted the endpoint significantly better than chance (MCC  $> 0$ ) and for a given endpoint many models performed similarly well on both internal and external validation (see the distribution of MCC in Fig. 2c). On the other hand, not all the submitted models performed equally well for any given endpoint. Some models performed no better than chance, even for some of the easy-to-predict endpoints, suggesting that additional factors were responsible for differences in model performance.

## Data analysis teams show different proficiency

Next, we summarized the external validation performance of the models nominated by the 17 teams that analyzed all 13 endpoints (Fig. 3). Nominated models represent a team’s best assessment of its model-building effort. The mean external validation MCC per team over 11 endpoints, excluding negative controls I and M, varied from 0.532 for data analysis team (DAT)24 to 0.263 for DAT3, indicating appreciable differences in performance of the models developed by different teams for the same data. Similar trends were observed when AUC was used as the performance metric (Supplementary Table 5) or when the original training and validation sets were swapped (Supplementary Tables 6 and 7). Table 2

summarizes the modeling approaches that were used by two or more MAQC-II data analysis teams.

Many factors may have played a role in the difference of external validation performance between teams. For instance, teams used different modeling factors, criteria for selecting the nominated models, and software packages and code. Moreover, some teams may have been more proficient at microarray data modeling and better at guarding against clerical errors. We noticed substantial variations in performance among the many *K*-nearest neighbor algorithm (KNN)-based models developed by four analysis teams (Supplementary Fig. 1). Follow-up investigations identified a few possible causes leading to the discrepancies in performance<sup>32</sup>. For example, DAT20 fixed the parameter ‘number of neighbors’ *K* = 3 in its data analysis protocol for all endpoints, whereas DAT18 varied *K* from 3 to 15 with a step size of 2. This investigation also revealed that even a detailed but standardized description of model building requested from all groups failed to capture many important tuning variables in the process. The subtle modeling differences not captured may have contributed to the differing performance levels achieved by the data analysis teams. The differences in performance for the models developed by various data analysis teams can also be observed from the changing patterns of internal and external validation performance across the 13 endpoints (Fig. 3, Supplementary Tables 5–7 and Supplementary Figs. 2–4). Our observations highlight the importance of good modeling practice in developing and validating microarray-based predictive models including reporting of computational details for results to be replicated<sup>26</sup>. In light of the MAQC-II experience, recording structured information about the steps and parameters of an analysis process seems highly desirable to facilitate peer review and reanalysis of results.

### Swap and original analyses lead to consistent results

To evaluate the reproducibility of the models generated by each team, we correlated the performance of each team’s models on the original training data set to performance on the validation data set and repeated this calculation for the swap experiment (Fig. 4). The correlation varied from 0.698–0.966 on the original experiment and from 0.443–0.954 on the swap experiment. For all but three teams (DAT3, DAT10 and DAT11) the original and swap correlations were within  $\pm 0.2$ , and all but three others (DAT4, DAT13 and DAT36) were within  $\pm 0.1$ , suggesting that the model building process was relatively robust, at least with respect to generating models with similar performance. For some data analysis teams the internal validation performance drastically overestimated the performance of the same model in predicting the validation data. Examination of some of those models revealed several reasons, including bias in the feature selection and cross-validation process<sup>28</sup>, findings consistent with what was observed from a recent literature survey<sup>33</sup>.

Previously, reanalysis of a widely cited single study<sup>34</sup> found that the results in the original publication were very fragile—that is, not reproducible if the training and validation sets were swapped<sup>35</sup>. Our observations, except for DAT3, DAT11 and DAT36 with correlation  $< 0.6$ , mainly resulting from failure of accurately predicting the positive-control endpoint H in the swap analysis (likely owing to operator errors), do not substantiate such fragility in the currently examined data sets. It is important to emphasize that we repeated the entire model building and evaluation processes during the swap analysis and, therefore, stability applies to the model building process for each data analysis team and not to a particular model or approach. Supplementary Figure 5 provides a more detailed look at the correlation of internal and external validation for each data analysis team and each endpoint for both the original (Supplementary Fig. 5a) and swap (Supplementary Fig. 5d) analyses.

As expected, individual feature lists differed from analysis group to analysis group and between models developed from the original and the swapped data. However, when feature

lists were mapped to biological processes, a greater degree of convergence and concordance was observed. This has been proposed previously but has never been demonstrated in a comprehensive manner over many data sets and thousands of models as was done in MAQC-II<sup>36</sup>.

### The effect of modeling factors is modest

To rigorously identify potential sources of variance that explain the variability in external-validation performance (Fig. 2c), we applied random effect modeling (Fig. 5a). We observed that the endpoint itself is by far the dominant source of variability, explaining >65% of the variability in the external validation performance. All other factors explain <8% of the total variance, and the residual variance is ~6%. Among the factors tested, those involving interactions with endpoint have a relatively large effect, in particular the interaction between endpoint with organization and classification algorithm, highlighting variations in proficiency between analysis teams.

To further investigate the impact of individual levels within each modeling factor, we estimated the empirical best linear unbiased predictors (BLUPs)<sup>37</sup>. Figure 5b shows the plots of BLUPs of the corresponding factors in Figure 5a with proportion of variation >1%. The BLUPs reveal the effect of each level of the factor to the corresponding MCC value. The BLUPs of the main endpoint effect show that rat liver necrosis, breast cancer estrogen receptor status and the sex of the patient (endpoints C, E, H and L) are relatively easier to be predicted with ~0.2–0.4 advantage contributed on the corresponding MCC values. The rest of the endpoints are relatively harder to be predicted with about –0.1 to –0.2 disadvantage contributed to the corresponding MCC values. The main factors of normalization, classification algorithm, the number of selected features and the feature selection method have an impact of –0.1 to 0.1 on the corresponding MCC values. Loess normalization was applied to the endpoints (J, K and L) for the neuroblastoma data set with the two-color Agilent platform and has 0.1 advantage to MCC values. Among the Microarray Analysis Suite version 5 (MAS5), Robust Multichip Analysis (RMA) and dChip normalization methods that were applied to all endpoints (A, C, D, E, F, G and H) for Affymetrix data, the dChip method has a lower BLUP than the others. Because normalization methods are partially confounded with endpoints, it may not be suitable to compare methods between different confounded groups. Among classification methods, discriminant analysis has the largest positive impact of 0.056 on the MCC values. Regarding the number of selected features, larger bin number has better impact on the average across endpoints. The bin number is assigned by applying the ceiling function to the log base 10 of the number of selected features. All the feature selection methods have a slight impact of –0.025 to 0.025 on MCC values except for recursive feature elimination (RFE) that has an impact of –0.006. In the plots of the four selected interactions, the estimated BLUPs vary across endpoints. The large variation across endpoints implies the impact of the corresponding modeling factor on different endpoints can be very different. Among the four interaction plots (see Supplementary Fig. 6 for a clear labeling of each interaction term), the corresponding BLUPs of the three-way interaction of organization, classification algorithm and endpoint show the highest variation. This may be due to different tuning parameters applied to individual algorithms for different organizations, as was the case for KNN<sup>32</sup>.

We also analyzed the relative importance of modeling factors on external-validation prediction performance using a decision tree model<sup>38</sup>. The analysis results revealed observations (Supplementary Fig. 7) largely consistent with those above. First, the endpoint code was the most influential modeling factor. Second, feature selection method, normalization and summarization method, classification method and organization code also contributed to prediction performance, but their contribution was relatively small.

## Feature list stability is correlated with endpoint predictability

Prediction performance is the most important criterion for evaluating the performance of a predictive model and its modeling process. However, the robustness and mechanistic relevance of the model and the corresponding gene signature is also important (Supplementary Fig. 8). That is, given comparable prediction performance between two modeling processes, the one yielding a more robust and reproducible gene signature across similar data sets (e.g., by swapping the training and validation sets), which is therefore less susceptible to sporadic fluctuations in the data, or the one that provides new insights to the underlying biology is preferable. Reproducibility or stability of feature sets is best studied by running the same model selection protocol on two distinct collections of samples, a scenario only possible, in this case, after the blind validation data were distributed to the data analysis teams that were asked to perform their analysis after swapping their original training and test sets. Supplementary Figures 9 and 10 show that, although the feature space is extremely large for microarray data, different teams and protocols were able to consistently select the best-performing features. Analysis of the lists of features indicated that for endpoints relatively easy to predict, various data analysis teams arrived at models that used more common features and the overlap of the lists from the original and swap analyses is greater than those for more difficult endpoints (Supplementary Figs. 9–11). Therefore, the level of stability of feature lists can be associated to the level of difficulty of the prediction problem (Supplementary Fig. 11), although multiple models with different feature lists and comparable performance can be found from the same data set<sup>39</sup>. Functional analysis of the most frequently selected genes by all data analysis protocols shows that many of these genes represent biological processes that are highly relevant to the clinical outcome that is being predicted<sup>36</sup>. The sex-based endpoints have the best overlap, whereas more difficult survival endpoints (in which disease processes are confounded by many other factors) have only marginally better overlap with biological processes relevant to the disease than that expected by random chance.

## Summary of MAQC-II observations and recommendations

The MAQC-II data analysis teams comprised a diverse group, some of whom were experienced microarray analysts whereas others were graduate students with little experience. In aggregate, the group's composition likely mimicked the broad scientific community engaged in building and publishing models derived from microarray data. The more than 30,000 models developed by 36 data analysis teams for 13 endpoints from six diverse clinical and preclinical data sets are a rich source from which to highlight several important observations.

First, model prediction performance was largely endpoint (biology) dependent (Figs. 2c and 3). The incorporation of multiple data sets and endpoints (including positive and negative controls) in the MAQC-II study design made this observation possible. Some endpoints are highly predictive based on the nature of the data, which makes it possible to build good models, provided that sound modeling procedures are used. Other endpoints are inherently difficult to predict regardless of the model development protocol.

Second, there are clear differences in proficiency between data analysis teams (organizations) and such differences are correlated with the level of experience of the team. For example, the top-performing teams shown in Figure 3 were mainly industrial participants with many years of experience in microarray data analysis, whereas bottom-performing teams were mainly less-experienced graduate students or researchers. Based on results from the positive and negative endpoints, we noticed that simple errors were sometimes made, suggesting rushed efforts due to lack of time or unnoticed implementation flaws. This observation strongly suggests that mechanisms are needed to ensure the

reliability of results presented to the regulatory agencies, journal editors and the research community. By examining the practices of teams whose models did not perform well, future studies might be able to identify pitfalls to be avoided. Likewise, practices adopted by top-performing teams can provide the basis for developing good modeling practices.

Third, the internal validation performance from well-implemented, unbiased cross-validation shows a high degree of concordance with the external validation performance in a strict blinding process (Fig. 2). This observation was not possible from previously published studies owing to the small number of available endpoints tested in them.

Fourth, many models with similar performance can be developed from a given data set (Fig. 2). Similar prediction performance is attainable when using different modeling algorithms and parameters, and simple data analysis methods often perform as well as more complicated approaches<sup>32,40</sup>. Although it is not essential to include the same features in these models to achieve comparable prediction performance, endpoints that were easier to predict generally yielded models with more common features, when analyzed by different teams (Supplementary Fig. 11).

Finally, applying good modeling practices appeared to be more important than the actual choice of a particular algorithm over the others within the same step in the modeling process. This can be seen in the diverse choices of the modeling factors used by teams that produced models that performed well in the blinded validation (Table 2) where modeling factors did not universally contribute to variations in model performance among good performing teams (Fig. 5).

Summarized below are the model building steps recommended to the MAQC-II data analysis teams. These may be applicable to model building practitioners in the general scientific community.

Step one (design). There is no exclusive set of steps and procedures, in the form of a checklist, to be followed by any practitioner for all problems. However, normal good practice on the study design and the ratio of sample size to classifier complexity should be followed. The frequently used options for normalization, feature selection and classification are good starting points (Table 2).

Step two (pilot study or internal validation). This can be accomplished by bootstrap or cross-validation such as the ten repeats of a fivefold cross-validation procedure adopted by most MAQC-II teams. The samples from the pilot study are not replaced for the pivotal study; rather they are augmented to achieve 'appropriate' target size.

Step three (pivotal study or external validation). Many investigators assume that the most conservative approach to a pivotal study is to simply obtain a test set completely independent of the training set(s). However, it is good to keep in mind the exchange<sup>34,35</sup> regarding the fragility of results when the training and validation sets are swapped. Results from further resampling (including simple swapping as in MAQC-II) across the training and validation sets can provide important information about the reliability of the models and the modeling procedures, but the complete separation of the training and validation sets should be maintained<sup>41</sup>.

Finally, a perennial issue concerns reuse of the independent validation set after modifications to an originally designed and validated data analysis algorithm or protocol. Such a process turns the validation set into part of the design or training set<sup>42</sup>. Ground rules must be developed for avoiding this approach and penalizing it when it occurs; and practitioners should guard against using it before such ground rules are well established.

## DISCUSSION

MAQC-II conducted a broad observational study of the current community landscape of gene-expression profile-based predictive model development. Microarray gene expression profiling is among the most commonly used analytical tools in biomedical research. Analysis of the high-dimensional data generated by these experiments involves multiple steps and several critical decision points that can profoundly influence the soundness of the results<sup>43</sup>. An important requirement of a sound internal validation is that it must include feature selection and parameter optimization within each iteration to avoid overly optimistic estimations of prediction performance<sup>28,29,44</sup>. To what extent this information has been disseminated and followed by the scientific community in current microarray analysis remains unknown<sup>33</sup>. Concerns have been raised that results published by one group of investigators often cannot be confirmed by others even if the same data set is used<sup>26</sup>. An inability to confirm results may stem from any of several reasons: (i) insufficient information is provided about the methodology that describes which analysis has actually been done; (ii) data preprocessing (normalization, gene filtering and feature selection) is too complicated and insufficiently documented to be reproduced; or (iii) incorrect or biased complex analytical methods<sup>26</sup> are performed. A distinct but related concern is that genomic data may yield prediction models that, even if reproducible on the discovery data set, cannot be extrapolated well in independent validation. The MAQC-II project provided a unique opportunity to address some of these concerns.

Notably, we did not place restrictions on the model building methods used by the data analysis teams. Accordingly, they adopted numerous different modeling approaches (Table 2 and Supplementary Table 4). For example, feature selection methods varied widely, from statistical significance tests, to machine learning algorithms, to those more reliant on differences in expression amplitude, to those employing knowledge of putative biological mechanisms associated with the endpoint. Prediction algorithms also varied widely. To make internal validation performance results comparable across teams for different models, we recommended that a model's internal performance was estimated using a ten times repeated fivefold cross-validation, but this recommendation was not strictly followed by all teams, which also allows us to survey internal validation approaches. The diversity of analysis protocols used by the teams is likely to closely resemble that of current research going forward, and in this context mimics reality. In terms of the space of modeling factors explored, MAQC-II is a survey of current practices rather than a randomized, controlled experiment; therefore, care should be taken in interpreting the results. For example, some teams did not analyze all endpoints, causing missing data (models) that may be confounded with other modeling factors.

Overall, the procedure followed to nominate MAQC-II candidate models was quite effective in selecting models that performed reasonably well during validation using independent data sets, although generally the selected models did not do as well in validation as in training. The drop in performance associated with the validation highlights the importance of not relying solely on internal validation performance, and points to the need to subject every classifier to at least one external validation. The selection of the 13 candidate models from many nominated models was achieved through a peer-review collaborative effort of many experts and could be described as slow, tedious and sometimes subjective (e.g., a data analysis team could only contribute one of the 13 candidate models). Even though they were still subject to over-optimism, the internal and external performance estimates of the candidate models were more concordant than those of the overall set of models. Thus the review was productive in identifying characteristics of reliable models.

An important lesson learned through MAQC-II is that it is almost impossible to retrospectively retrieve and document decisions that were made at every step during the feature selection and model development stage. This lack of complete description of the model building process is likely to be a common reason for the inability of different data analysis teams to fully reproduce each other's results<sup>32</sup>. Therefore, although meticulously documenting the classifier building procedure can be cumbersome, we recommend that all genomic publications include supplementary materials describing the model building and evaluation process in an electronic format. MAQC-II is making available six data sets with 13 endpoints that can be used in the future as a benchmark to verify that software used to implement new approaches performs as expected. Subjecting new software to benchmarks against these data sets could reassure potential users that the software is mature enough to be used for the development of predictive models in new data sets. It would seem advantageous to develop alternative ways to help determine whether specific implementations of modeling approaches and performance evaluation procedures are sound, and to identify procedures to capture this information in public databases.

The findings of the MAQC-II project suggest that when the same data sets are provided to a large number of data analysis teams, many groups can generate similar results even when different model building approaches are followed. This is concordant with studies<sup>29,33</sup> that found that given good quality data and an adequate number of informative features, most classification methods, if properly used, will yield similar predictive performance. This also confirms reports<sup>6,7,39</sup> on small data sets by individual groups that have suggested that several different feature selection methods and prediction algorithms can yield many models that are distinct, but have statistically similar performance. Taken together, these results provide perspective on the large number of publications in the bioinformatics literature that have examined the various steps of the multivariate prediction model building process and identified elements that are critical for achieving reliable results.

An important and previously underappreciated observation from MAQC-II is that different clinical endpoints represent very different levels of classification difficulty. For some endpoints the currently available data are sufficient to generate robust models, whereas for other endpoints currently available data do not seem to be sufficient to yield highly predictive models. An analysis done as part of the MAQC-II project and that focused on the breast cancer data demonstrates these points in more detail<sup>40</sup>. It is also important to point out that for some clinically meaningful endpoints studied in the MAQC-II project, gene expression data did not seem to significantly outperform models based on clinical covariates alone, highlighting the challenges in predicting the outcome of patients in a heterogeneous population and the potential need to combine gene expression data with clinical covariates (unpublished data).

The accuracy of the clinical sample annotation information may also play a role in the difficulty to obtain accurate prediction results on validation samples. For example, some samples were misclassified by almost all models (Supplementary Fig. 12). It is true even for some samples within the positive control endpoints H and L, as shown in Supplementary Table 8. Clinical information of neuroblastoma patients for whom the positive control endpoint L was uniformly misclassified were rechecked and the sex of three out of eight cases (NB412, NB504 and NB522) was found to be incorrectly annotated.

The companion MAQC-II papers published elsewhere give more in-depth analyses of specific issues such as the clinical benefits of genomic classifiers (unpublished data), the impact of different modeling factors on prediction performance<sup>45</sup>, the objective assessment of microarray cross-platform prediction<sup>46</sup>, cross-tissue prediction<sup>47</sup>, one-color versus two-color prediction comparison<sup>48</sup>, functional analysis of gene signatures<sup>36</sup> and recommendation

of a simple yet robust data analysis protocol based on the KNN<sup>32</sup>. For example, we systematically compared the classification performance resulting from one- and two-color gene-expression profiles of 478 neuroblastoma samples and found that analyses based on either platform yielded similar classification performance<sup>48</sup>. This newly generated one-color data set has been used to evaluate the applicability of the KNN-based simple data analysis protocol to future data sets<sup>32</sup>. In addition, the MAQC-II Genome-Wide Association Working Group assessed the variabilities in genotype calling due to experimental or algorithmic factors<sup>49</sup>.

In summary, MAQC-II has demonstrated that current methods commonly used to develop and assess multivariate gene-expression based predictors of clinical outcome were used appropriately by most of the analysis teams in this consortium. However, differences in proficiency emerged and this underscores the importance of proper implementation of otherwise robust analytical methods. Observations based on analysis of the MAQC-II data sets may be applicable to other diseases. The MAQC-II data sets are publicly available and are expected to be used by the scientific community as benchmarks to ensure proper modeling practices. The experience with the MAQC-II clinical data sets also reinforces the notion that clinical classification problems represent several different degrees of prediction difficulty that are likely to be associated with whether mRNA abundances measured in a specific data set are informative for the specific prediction problem. We anticipate that including other types of biological data at the DNA, microRNA, protein or metabolite levels will enhance our capability to more accurately predict the clinically relevant endpoints. The good modeling practice guidelines established by MAQC-II and lessons learned from this unprecedented collaboration provide a solid foundation from which other high-dimensional biological data could be more reliably used for the purpose of predictive and personalized medicine.

## METHODS

Methods and any associated references are available here.

## Supplementary Material

Refer to Web version on PubMed Central for supplementary material.

## Acknowledgments

The MAQC-II project was funded in part by the FDA's Office of Critical Path Programs (to L.S.). Participants from the National Institutes of Health (NIH) were supported by the Intramural Research Program of NIH, Bethesda, Maryland or the Intramural Research Program of the NIH, National Institute of Environmental Health Sciences (NIEHS), Research Triangle Park, North Carolina. J.F. was supported by the Division of Intramural Research of the NIEHS under contract HHSN273200700046U. Participants from the Johns Hopkins University were supported by grants from the NIH (1R01GM083084-01 and 1R01RR021967-01A2 to R.A.I. and T32GM074906 to M.M.). Participants from the Weill Medical College of Cornell University were partially supported by the Biomedical Informatics Core of the Institutional Clinical and Translational Science Award RFA-RM-07-002. F.C. acknowledges resources from The HRH Prince Alwaleed Bin Talal Bin Abdulaziz Alsaud Institute for Computational Biomedicine and from the David A. Cofrin Center for Biomedical Information at Weill Cornell. The data set from The Hamner Institutes for Health Sciences was supported by a grant from the American Chemistry Council's Long Range Research Initiative. The breast cancer data set was generated with support of grants from NIH (R-01 to L.P.), The Breast Cancer Research Foundation (to L.P. and W.F.S.) and the Faculty Incentive Funds of the University of Texas MD Anderson Cancer Center (to W.F.S.). The data set from the University of Arkansas for Medical Sciences was supported by National Cancer Institute (NCI) PO1 grant CA55819-01A1, NCI R33 Grant CA97513-01, Donna D. and Donald M. Lambert Lebow Fund to Cure Myeloma and Nancy and Steven Grand Foundation. We are grateful to the individuals whose gene expression data were used in this study. All MAQC-II participants freely donated their time and reagents for the completion and analyses of the MAQC-II project. The MAQC-II consortium also thanks R. O'Neill for his encouragement and coordination among FDA Centers on the formation of the RBWG. The MAQC-II consortium gratefully dedicates this work in memory of R.F. Wagner who

enthusiastically worked on the MAQC-II project and inspired many of us until he unexpectedly passed away in June 2008.

## References

1. Marshall E. Getting the noise out of gene arrays. *Science*. 2004; 306:630–631. [PubMed: 15499004]
2. Frantz S. An array of problems. *Nat Rev Drug Discov*. 2005; 4:362–363. [PubMed: 15902768]
3. Michiels S, Koscielny S, Hill C. Prediction of cancer outcome with microarrays: a multiple random validation strategy. *Lancet*. 2005; 365:488–492. [PubMed: 15705458]
4. Ntzani EE, Ioannidis JP. Predictive ability of DNA microarrays for cancer outcomes and correlates: an empirical assessment. *Lancet*. 2003; 362:1439–1444. [PubMed: 14602436]
5. Ioannidis JP. Microarrays and molecular research: noise discovery? *Lancet*. 2005; 365:454–455. [PubMed: 15705441]
6. Ein-Dor L, Kela I, Getz G, Givol D, Domany E. Outcome signature genes in breast cancer: is there a unique set? *Bioinformatics*. 2005; 21:171–178. [PubMed: 15308542]
7. Ein-Dor L, Zuk O, Domany E. Thousands of samples are needed to generate a robust gene list for predicting outcome in cancer. *Proc Natl Acad Sci USA*. 2006; 103:5923–5928. [PubMed: 16585533]
8. Shi L, et al. QA/QC: challenges and pitfalls facing the microarray community and regulatory agencies. *Expert Rev Mol Diagn*. 2004; 4:761–777. [PubMed: 15525219]
9. Shi L, et al. Cross-platform comparability of microarray technology: intra-platform consistency and appropriate data analysis procedures are essential. *BMC Bioinformatics*. 2005; 6 (Suppl 2):S12. [PubMed: 16026597]
10. Shi L, et al. The MicroArray Quality Control (MAQC) project shows inter- and intraplatform reproducibility of gene expression measurements. *Nat Biotechnol*. 2006; 24:1151–1161. [PubMed: 16964229]
11. Guo L, et al. Rat toxicogenomic study reveals analytical consistency across microarray platforms. *Nat Biotechnol*. 2006; 24:1162–1169. [PubMed: 17061323]
12. Canales RD, et al. Evaluation of DNA microarray results with quantitative gene expression platforms. *Nat Biotechnol*. 2006; 24:1115–1122. [PubMed: 16964225]
13. Patterson TA, et al. Performance comparison of one-color and two-color platforms within the MicroArray Quality Control (MAQC) project. *Nat Biotechnol*. 2006; 24:1140–1150. [PubMed: 16964228]
14. Shippy R, et al. Using RNA sample titrations to assess microarray platform performance and normalization techniques. *Nat Biotechnol*. 2006; 24:1123–1131. [PubMed: 16964226]
15. Tong W, et al. Evaluation of external RNA controls for the assessment of microarray performance. *Nat Biotechnol*. 2006; 24:1132–1139. [PubMed: 16964227]
16. Irizarry RA, et al. Multiple-laboratory comparison of microarray platforms. *Nat Methods*. 2005; 2:345–350. [PubMed: 15846361]
17. Strauss E. Arrays of hope. *Cell*. 2006; 127:657–659. [PubMed: 17110319]
18. Shi L, Perkins RG, Fang H, Tong W. Reproducible and reliable microarray results through quality control: good laboratory proficiency and appropriate data analysis practices are essential. *Curr Opin Biotechnol*. 2008; 19:10–18. [PubMed: 18155896]
19. Dudoit S, Fridlyand J, Speed TP. Comparison of discrimination methods for the classification of tumors using gene expression data. *J Am Stat Assoc*. 2002; 97:77–87.
20. Goodsaid FM, et al. Voluntary exploratory data submissions to the US FDA and the EMA: experience and impact. *Nat Rev Drug Discov*. 2010; 9:435–445. [PubMed: 20514070]
21. van 't Veer LJ, et al. Gene expression profiling predicts clinical outcome of breast cancer. *Nature*. 2002; 415:530–536. [PubMed: 11823860]
22. Buyse M, et al. Validation and clinical utility of a 70-gene prognostic signature for women with node-negative breast cancer. *J Natl Cancer Inst*. 2006; 98:1183–1192. [PubMed: 16954471]
23. Dumur CI, et al. Interlaboratory performance of a microarray-based gene expression test to determine tissue of origin in poorly differentiated and undifferentiated cancers. *J Mol Diagn*. 2008; 10:67–77. [PubMed: 18083688]

24. Deng MC, et al. Noninvasive discrimination of rejection in cardiac allograft recipients using gene expression profiling. *Am J Transplant*. 2006; 6:150–160. [PubMed: 16433769]
25. Coombes KR, Wang J, Baggerly KA. Microarrays: retracing steps. *Nat Med*. 2007; 13:1276–1277. author reply 1277–1278. [PubMed: 17987014]
26. Ioannidis JPA, et al. Repeatability of published microarray gene expression analyses. *Nat Genet*. 2009; 41:149–155. [PubMed: 19174838]
27. Baggerly KA, Edmonson SR, Morris JS, Coombes KR. High-resolution serum proteomic patterns for ovarian cancer detection. *Endocr Relat Cancer*. 2004; 11:583–584. author reply 585–587. [PubMed: 15613439]
28. Ambroise C, McLachlan GJ. Selection bias in gene extraction on the basis of microarray gene-expression data. *Proc Natl Acad Sci USA*. 2002; 99:6562–6566. [PubMed: 11983868]
29. Simon R. Using DNA microarrays for diagnostic and prognostic prediction. *Expert Rev Mol Diagn*. 2003; 3:587–595. [PubMed: 14510179]
30. Dobbin KK, et al. Interlaboratory comparability study of cancer gene expression analysis using oligonucleotide microarrays. *Clin Cancer Res*. 2005; 11:565–572. [PubMed: 15701842]
31. Shedden K, et al. Gene expression-based survival prediction in lung adenocarcinoma: a multi-site, blinded validation study. *Nat Med*. 2008; 14:822–827. [PubMed: 18641660]
32. Parry RM, et al. K-nearest neighbors (KNN) models for microarray gene-expression analysis and reliable clinical outcome prediction. *Pharmacogenomics J*. 2010; 10:292–309. [PubMed: 20676068]
33. Dupuy A, Simon RM. Critical review of published microarray studies for cancer outcome and guidelines on statistical analysis and reporting. *J Natl Cancer Inst*. 2007; 99:147–157. [PubMed: 17227998]
34. Dave SS, et al. Prediction of survival in follicular lymphoma based on molecular features of tumor-infiltrating immune cells. *N Engl J Med*. 2004; 351:2159–2169. [PubMed: 15548776]
35. Tibshirani R. Immune signatures in follicular lymphoma. *N Engl J Med*. 2005; 352:1496–1497. author reply 1496–1497. [PubMed: 15814892]
36. Shi W, et al. Functional analysis of multiple genomic signatures demonstrates that classification algorithms choose phenotype-related genes. *Pharmacogenomics J*. 2010; 10:310–323. [PubMed: 20676069]
37. Robinson GK. That BLUP is a good thing: the estimation of random effects. *Stat Sci*. 1991; 6:15–32.
38. Hothorn T, Hornik K, Zeileis A. Unbiased recursive partitioning: a conditional inference framework. *J Comput Graph Statist*. 2006; 15:651–674.
39. Boutros PC, et al. Prognostic gene signatures for non-small-cell lung cancer. *Proc Natl Acad Sci USA*. 2009; 106:2824–2828. [PubMed: 19196983]
40. Popovici V, et al. Effect of training sample size and classification difficulty on the accuracy of genomic predictors. *Breast Cancer Res*. 2010; 12:R5. [PubMed: 20064235]
41. Yousef WA, Wagner RF, Loew MH. Assessing classifiers from two independent data sets using ROC analysis: a nonparametric approach. *IEEE Trans Pattern Anal Mach Intell*. 2006; 28:1809–1817. [PubMed: 17063685]
42. Gur D, Wagner RF, Chan HP. On the repeated use of databases for testing incremental improvement of computer-aided detection schemes. *Acad Radiol*. 2004; 11:103–105. [PubMed: 14746409]
43. Allison DB, Cui X, Page GP, Sabripour M. Microarray data analysis: from disarray to consolidation and consensus. *Nat Rev Genet*. 2006; 7:55–65. [PubMed: 16369572]
44. Wood IA, Visscher PM, Mengersen KL. Classification based upon gene expression data: bias and precision of error rates. *Bioinformatics*. 2007; 23:1363–1370. [PubMed: 17392326]
45. Luo J, et al. A comparison of batch effect removal methods for enhancement of prediction performance using MAQC-II microarray gene expression data. *Pharmacogenomics J*. 2010; 10:278–291. [PubMed: 20676067]
46. Fan X, et al. Consistency of predictive signature genes and classifiers generated using different microarray platforms. *Pharmacogenomics J*. 2010; 10:247–257. [PubMed: 20676064]

47. Huang J, et al. Genomic indicators in the blood predict drug-induced liver injury. *Pharmacogenomics J.* 2010; 10:267–277. [PubMed: 20676066]
48. Oberthuer A, et al. Comparison of performance of one-color and two-color geneexpression analyses in predicting clinical endpoints of neuroblastoma patients. *Pharmacogenomics J.* 2010; 10:258–266. [PubMed: 20676065]
49. Hong H, et al. Assessing sources of inconsistencies in genotypes and their effects on genome-wide association studies with HapMap samples. *Pharmacogenomics J.* 2010; 10:364–374. [PubMed: 20368714]

## Appendix

Leming Shi<sup>1</sup>, Gregory Campbell<sup>2</sup>, Wendell D Jones<sup>3</sup>, Fabien Campagne<sup>4</sup>, Zhining Wen<sup>1</sup>, Stephen J Walker<sup>5</sup>, Zhenqiang Su<sup>6</sup>, Tzu-Ming Chu<sup>7</sup>, Federico M Goodsaid<sup>8</sup>, Lajos Pusztai<sup>9</sup>, John D Shaughnessy Jr<sup>10</sup>, André Oberthuer<sup>11</sup>, Russell S Thomas<sup>12</sup>, Richard S Paules<sup>13</sup>, Mark Fielden<sup>14</sup>, Bart Barlogie<sup>10</sup>, Weijie Chen<sup>2</sup>, Pan Du<sup>15</sup>, Matthias Fischer<sup>11</sup>, Cesare Furlanello<sup>16</sup>, Brandon D Gallas<sup>2</sup>, Xijin Ge<sup>17</sup>, Dalila B Megherbi<sup>18</sup>, W Fraser Symmans<sup>19</sup>, May D Wang<sup>20</sup>, John Zhang<sup>21</sup>, Hans Bitter<sup>22</sup>, Benedikt Brors<sup>23</sup>, Pierre R Bushel<sup>13</sup>, Max Bylesjo<sup>24</sup>, Minjun Chen<sup>1</sup>, Jie Cheng<sup>25</sup>, Jing Cheng<sup>26</sup>, Jeff Chou<sup>13</sup>, Timothy S Davison<sup>27</sup>, Mauro Delorenzi<sup>28</sup>, Youping Deng<sup>29</sup>, Viswanath Devanarayan<sup>30</sup>, David J Dix<sup>31</sup>, Joaquin Dopazo<sup>32</sup>, Kevin C Dorff<sup>33</sup>, Fathi Elloumi<sup>31</sup>, Jianqing Fan<sup>34</sup>, Shicai Fan<sup>35</sup>, Xiaohui Fan<sup>36</sup>,

<sup>1</sup>National Center for Toxicological Research, US Food and Drug Administration, Jefferson, Arkansas, USA

<sup>2</sup>Center for Devices and Radiological Health, US Food and Drug Administration, Silver Spring, Maryland, USA

<sup>3</sup>Expression Analysis Inc., Durham, North Carolina, USA

<sup>4</sup>Department of Physiology and Biophysics and HRH Prince Alwaleed Bin Talal Bin Abdulaziz Alsaud Institute for Computational Biomedicine, Weill Medical College of Cornell University, New York, New York, USA

<sup>5</sup>Wake Forest Institute for Regenerative Medicine, Wake Forest University, Winston-Salem, North Carolina, USA

<sup>6</sup>Z-Tech, an ICF International Company at NCTR/FDA, Jefferson, Arkansas, USA

<sup>7</sup>SAS Institute Inc., Cary, North Carolina, USA

<sup>8</sup>Center for Drug Evaluation and Research, US Food and Drug Administration, Silver Spring, Maryland, USA

<sup>9</sup>Breast Medical Oncology Department, University of Texas (UT) M.D. Anderson Cancer Center, Houston, Texas, USA

<sup>10</sup>Myeloma Institute for Research and Therapy, University of Arkansas for Medical Sciences, Little Rock, Arkansas, USA

<sup>11</sup>Department of Pediatric Oncology and Hematology and Center for Molecular Medicine (CMMC), University of Cologne, Cologne, Germany

<sup>12</sup>The Hamner Institutes for Health Sciences, Research Triangle Park, North Carolina, USA

<sup>13</sup>National Institute of Environmental Health Sciences, National Institutes of Health, Research Triangle Park, North Carolina, USA

<sup>14</sup>Roche Palo Alto LLC, South San Francisco, California, USA

<sup>15</sup>Biomedical Informatics Center, Northwestern University, Chicago, Illinois, USA

<sup>16</sup>Fondazione Bruno Kessler, Povo-Trento, Italy

<sup>17</sup>Department of Mathematics & Statistics, South Dakota State University, Brookings, South Dakota, USA

<sup>18</sup>CMINDS Research Center, Department of Electrical and Computer Engineering, University of Massachusetts Lowell, Lowell, Massachusetts, USA

<sup>19</sup>Department of Pathology, UT M.D. Anderson Cancer Center, Houston, Texas, USA

<sup>20</sup>Department of Biomedical Engineering, Georgia Institute of Technology and Emory University, Atlanta, Georgia, USA

<sup>21</sup>Systems Analytics Inc., Waltham, Massachusetts, USA

<sup>22</sup>Hoffmann-LaRoche, Nutley, New Jersey, USA

<sup>23</sup>Department of Theoretical Bioinformatics, German Cancer Research Center (DKFZ), Heidelberg, Germany

<sup>24</sup>Computational Life Science Cluster (CLiC), Chemical Biology Center (KBC), Umeå University, Umeå, Sweden

<sup>25</sup>GlaxoSmithKline, Collegeville, Pennsylvania, USA

<sup>26</sup>Medical Systems Biology Research Center, School of Medicine, Tsinghua University, Beijing, China

<sup>27</sup>Almac Diagnostics Ltd., Craigavon, UK

<sup>28</sup>Swiss Institute of Bioinformatics, Lausanne, Switzerland

<sup>29</sup>Department of Biological Sciences, University of Southern Mississippi, Hattiesburg, Mississippi, USA

<sup>30</sup>Global Pharmaceutical R&D, Abbott Laboratories, Souderton, Pennsylvania, USA

<sup>31</sup>National Center for Computational Toxicology, US Environmental Protection Agency, Research Triangle Park, North Carolina, USA

<sup>32</sup>Department of Bioinformatics and Genomics, Centro de Investigación Príncipe Felipe (CIPF), Valencia, Spain

<sup>33</sup>HRH Prince Alwaleed Bin Talal Bin Abdulaziz Alsaud Institute for Computational Biomedicine, Weill Medical College of Cornell University, New York, New York, USA

<sup>34</sup>Department of Operation Research and Financial Engineering, Princeton University, Princeton, New Jersey, USA

<sup>35</sup>MOE Key Laboratory of Bioinformatics and Bioinformatics Division, TNLIST/Department of Automation, Tsinghua University, Beijing, China

<sup>36</sup>Institute of Pharmaceutical Informatics, College of Pharmaceutical Sciences, Zhejiang University, Hangzhou, Zhejiang, China

Hong Fang<sup>6</sup>, Nina Gonzaludo<sup>37</sup>, Kenneth R Hess<sup>38</sup>, Huixiao Hong<sup>1</sup>, Jun Huan<sup>39</sup>, Rafael A Irizarry<sup>40</sup>, Richard Judson<sup>31</sup>, Dilafruz Juraeva<sup>23</sup>, Samir Lababidi<sup>41</sup>, Christophe G Lambert<sup>42</sup>, Li Li<sup>7</sup>, Yanen Li<sup>43</sup>, Zhen Li<sup>31</sup>, Simon M Lin<sup>15</sup>, Guozhen Liu<sup>44</sup>, Edward K Lobenhofer<sup>45</sup>, Jun Luo<sup>21</sup>, Wen Luo<sup>46</sup>, Matthew N McCall<sup>40</sup>, Yuri Nikolsky<sup>47</sup>, Gene A Pennello<sup>2</sup>, Roger G Perkins<sup>1</sup>, Reena Philip<sup>2</sup>, Vlad Popovici<sup>28</sup>, Nathan D Price<sup>48</sup>, Feng Qian<sup>6</sup>, Andreas Scherer<sup>49</sup>, Tielu Shi<sup>50</sup>, Weiwei Shi<sup>47</sup>, Jaeyun Sung<sup>48</sup>, Danielle Thierry-Mieg<sup>51</sup>, Jean Thierry-Mieg<sup>51</sup>, Venkata Thodima<sup>52</sup>, Johan Trygg<sup>24</sup>, Lakshmi Vishnuvajjala<sup>2</sup>, Sue Jane Wang<sup>8</sup>, Jianping Wu<sup>53</sup>, Yichao Wu<sup>54</sup>, Qian Xie<sup>55</sup>, Waleed A Yousef<sup>56</sup>, Liang Zhang<sup>53</sup>, Xuegong Zhang<sup>35</sup>, Sheng Zhong<sup>57</sup>, Yiming Zhou<sup>10</sup>, Sheng Zhu<sup>53</sup>, Dhivya Arasappan<sup>6</sup>, Wenjun Bao<sup>7</sup>, Anne Bergstrom Lucas<sup>58</sup>, Frank Berthold<sup>11</sup>, Richard J Brennan<sup>47</sup>, Andreas Bunes<sup>59</sup>, Jennifer G Catalano<sup>41</sup>, Chang Chang<sup>50</sup>, Rong Chen<sup>60</sup>, Yiyu Cheng<sup>36</sup>, Jian Cui<sup>50</sup>, Wendy Czika<sup>7</sup>, Francesca Demichelis<sup>61</sup>, Xutao Deng<sup>62</sup>, Damir Dosymbekov<sup>63</sup>, Roland Eils<sup>23</sup>, Yang Feng<sup>34</sup>, Jennifer Foster<sup>13</sup>, Stephanie Fulmer-Smentek<sup>58</sup>, James C Fuscoe<sup>1</sup>, Laurent Gatto<sup>64</sup>, Weigong Ge<sup>1</sup>, Darlene R Goldstein<sup>65</sup>, Li Guo<sup>66</sup>, Donald N Halbert<sup>67</sup>, Jing Han<sup>41</sup>, Stephen C Harris<sup>1</sup>, Christos Hatzis<sup>68</sup>, Damir Herman<sup>69</sup>, Jianping Huang<sup>36</sup>, Roderick V Jensen<sup>70</sup>, Rui Jiang<sup>35</sup>, Charles D Johnson<sup>71</sup>, Giuseppe Jurman<sup>16</sup>, Yvonne Kahlert<sup>11</sup>, Sadik A Khuder<sup>72</sup>, Matthias Kohl<sup>73</sup>, Jianying Li<sup>74</sup>, Li Li<sup>75</sup>, Menglong Li<sup>76</sup>, Quan-Zhen Li<sup>77</sup>, Shao Li<sup>36</sup>, Zhiguang Li<sup>1</sup>, Jie Liu<sup>1</sup>, Ying Liu<sup>35</sup>,

<sup>37</sup>Roche Palo Alto LLC, Palo Alto, California, USA

<sup>38</sup>Department of Biostatistics, UT M.D. Anderson Cancer Center, Houston, Texas, USA

<sup>39</sup>Department of Electrical Engineering & Computer Science, University of Kansas, Lawrence, Kansas, USA

<sup>40</sup>Department of Biostatistics, Johns Hopkins University, Baltimore, Maryland, USA

<sup>41</sup>Center for Biologics Evaluation and Research, US Food and Drug Administration, Bethesda, Maryland, USA

<sup>42</sup>Golden Helix Inc., Bozeman, Montana, USA

<sup>43</sup>Department of Computer Science, University of Illinois at Urbana-Champaign, Urbana, Illinois, USA

<sup>44</sup>SABiosciences Corp., a Qiagen Company, Frederick, Maryland, USA

<sup>45</sup>Cogenics, a Division of Clinical Data Inc., Morrisville, North Carolina, USA

<sup>46</sup>Ligand Pharmaceuticals Inc., La Jolla, California, USA

<sup>47</sup>GeneGo Inc., Encinitas, California, USA

<sup>48</sup>Department of Chemical and Biomolecular Engineering, University of Illinois at Urbana-Champaign, Urbana, Illinois, USA

<sup>49</sup>Spheromics, Kontiolahti, Finland

<sup>50</sup>The Center for Bioinformatics and The Institute of Biomedical Sciences, School of Life Science, East China Normal University, Shanghai, China

<sup>51</sup>National Center for Biotechnology Information, National Institutes of Health, Bethesda, Maryland, USA

<sup>52</sup>Rockefeller Research Laboratories, Memorial Sloan-Kettering Cancer Center, New York, New York, USA

<sup>53</sup>CapitalBio Corporation, Beijing, China

<sup>54</sup>Department of Statistics, North Carolina State University, Raleigh, North Carolina, USA

<sup>55</sup>SRA International (EMMES), Rockville, Maryland, USA

<sup>56</sup>Helwan University, Helwan, Egypt

<sup>57</sup>Department of Bioengineering, University of Illinois at Urbana-Champaign, Urbana, Illinois, USA

<sup>58</sup>Agilent Technologies Inc., Santa Clara, California, USA

<sup>59</sup>F. Hoffmann-La Roche Ltd., Basel, Switzerland

<sup>60</sup>Stanford Center for Biomedical Informatics Research, Stanford University, Stanford, California, USA

<sup>61</sup>Department of Pathology and Laboratory Medicine and HRH Prince Alwaleed Bin Talal Bin Abdulaziz Alsaud Institute for Computational Biomedicine, Weill Medical College of Cornell University, New York, New York, USA

<sup>62</sup>Cedars-Sinai Medical Center, UCLA David Geffen School of Medicine, Los Angeles, California, USA

<sup>63</sup>Vavilov Institute for General Genetics, Russian Academy of Sciences, Moscow, Russia

<sup>64</sup>DNAVision SA, Gosselies, Belgium

<sup>65</sup>École Polytechnique Fédérale de Lausanne (EPFL), Lausanne, Switzerland

<sup>66</sup>State Key Laboratory of Multi-phase Complex Systems, Institute of Process Engineering, Chinese Academy of Sciences, Beijing, China

<sup>67</sup>Abbott Laboratories, Abbott Park, Illinois, USA

<sup>68</sup>Nuvera Biosciences Inc., Woburn, Massachusetts, USA

<sup>69</sup>Winthrop P. Rockefeller Cancer Institute, University of Arkansas for Medical Sciences, Little Rock, Arkansas, USA

<sup>70</sup>VirginiaTech, Blacksburg, Virginia, USA

<sup>71</sup>BioMath Solutions, LLC, Austin, Texas, USA

<sup>72</sup>Bioinformatic Program, University of Toledo, Toledo, Ohio, USA

<sup>73</sup>Department of Mathematics, University of Bayreuth, Bayreuth, Germany

<sup>74</sup>Lineberger Comprehensive Cancer Center, University of North Carolina, Chapel Hill, North Carolina, USA

<sup>75</sup>Pediatric Department, Stanford University, Stanford, California, USA

<sup>76</sup>College of Chemistry, Sichuan University, Chengdu, Sichuan, China

<sup>77</sup>University of Texas Southwestern Medical Center (UTSW), Dallas, Texas, USA

Zhichao Liu<sup>1</sup>, Lu Meng<sup>35</sup>, Manuel Madera<sup>18</sup>, Francisco Martinez-Murillo<sup>2</sup>, Ignacio Medina<sup>78</sup>, Joseph Meehan<sup>6</sup>, Kelci Miclaus<sup>7</sup>, Richard A Moffitt<sup>20</sup>, David Montaner<sup>78</sup>, Piali Mukherjee<sup>33</sup>, George J Mulligan<sup>79</sup>, Padraic Neville<sup>7</sup>, Tatiana Nikolskaya<sup>47</sup>, Baitang Ning<sup>1</sup>, Grier P Page<sup>80</sup>, Joel Parker<sup>3</sup>, R Mitchell Parry<sup>20</sup>, Xuejun Peng<sup>81</sup>, Ron L Peterson<sup>82</sup>, John H Phan<sup>20</sup>, Brian Quanz<sup>39</sup>, Yi Ren<sup>83</sup>, Samantha Riccadonna<sup>16</sup>, Alan H Roter<sup>84</sup>, Frank W Samuelson<sup>2</sup>, Martin M Schumacher<sup>85</sup>, Joseph D Shambaugh<sup>86</sup>, Qiang Shi<sup>1</sup>, Richard Shippy<sup>87</sup>, Shengzhu Si<sup>88</sup>, Aaron Smalter<sup>39</sup>, Christos Sotiriou<sup>89</sup>, Mat Soukup<sup>8</sup>, Frank Staedtler<sup>85</sup>, Guido Steiner<sup>90</sup>, Todd H Stokes<sup>20</sup>, Qinglan Sun<sup>53</sup>, Pei-Yi Tan<sup>7</sup>, Rong Tang<sup>2</sup>, Zivana Tezak<sup>2</sup>, Brett Thorn<sup>1</sup>, Marina Tsyganova<sup>63</sup>, Yaron Turpaz<sup>91</sup>, Silvia C Vega<sup>92</sup>, Roberto Visintainer<sup>16</sup>, Juergen von Frese<sup>93</sup>, Charles Wang<sup>62</sup>, Eric Wang<sup>21</sup>, Junwei Wang<sup>50</sup>, Wei Wang<sup>94</sup>, Frank Westermann<sup>23</sup>, James C Willey<sup>95</sup>, Matthew Woods<sup>21</sup>, Shujian Wu<sup>96</sup>, Nianqing Xiao<sup>97</sup>, Joshua Xu<sup>6</sup>, Lei Xu<sup>1</sup>, Lun Yang<sup>1</sup>, Xiao Zeng<sup>44</sup>, Jialu Zhang<sup>8</sup>, Li Zhang<sup>8</sup>, Min Zhang<sup>1</sup>, Chen Zhao<sup>50</sup>, Raj K Puri<sup>41</sup>, Uwe Scherf<sup>2</sup>, Weida Tong<sup>1</sup> & Russell D Wolfinger<sup>7</sup>

<sup>78</sup>Centro de Investigación Príncipe Felipe (CIPF), Valencia, Spain

<sup>79</sup>Millennium Pharmaceuticals Inc., Cambridge, Massachusetts, USA

<sup>80</sup>RTI International, Atlanta, Georgia, USA

<sup>81</sup>Takeda Global R & D Center, Inc., Deerfield, Illinois, USA

<sup>82</sup>Novartis Institutes of Biomedical Research, Cambridge, Massachusetts, USA

<sup>83</sup>W.M. Keck Center for Collaborative Neuroscience, Rutgers, The State University of New Jersey, Piscataway, New Jersey, USA

<sup>84</sup>Entelos Inc., Foster City, California, USA

<sup>85</sup>Biomarker Development, Novartis Institutes of BioMedical Research, Novartis Pharma AG, Basel, Switzerland

<sup>86</sup>Genedata Inc., Lexington, Massachusetts, USA

<sup>87</sup>Affymetrix Inc., Santa Clara, California, USA

<sup>88</sup>Department of Chemistry and Chemical Engineering, Hefei Teachers College, Hefei, Anhui, China

<sup>89</sup>Institut Jules Bordet, Brussels, Belgium

<sup>90</sup>Biostatistics, F. Hoffmann-La Roche Ltd., Basel, Switzerland

<sup>91</sup>Lilly Singapore Centre for Drug Discovery, Immunos, Singapore

<sup>92</sup>Microsoft Corporation, US Health Solutions Group, Redmond, Washington, USA

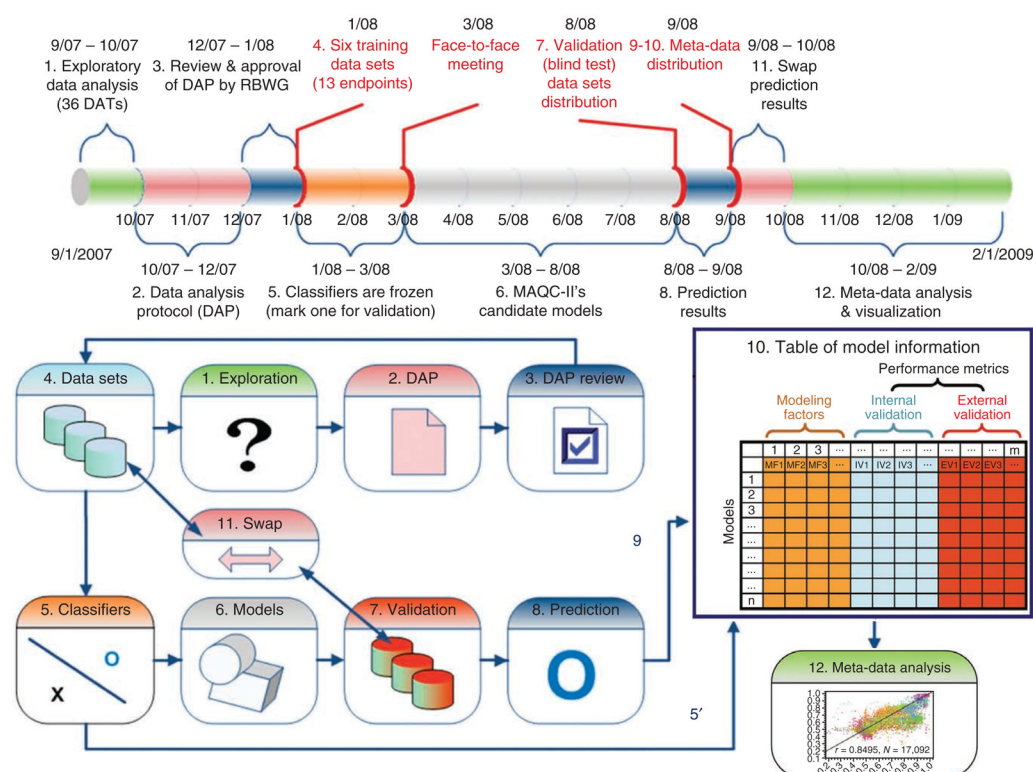
<sup>93</sup>Data Analysis Solutions DA-SOL GmbH, Greifenberg, Germany

<sup>94</sup>Cornell University, Ithaca, New York, USA

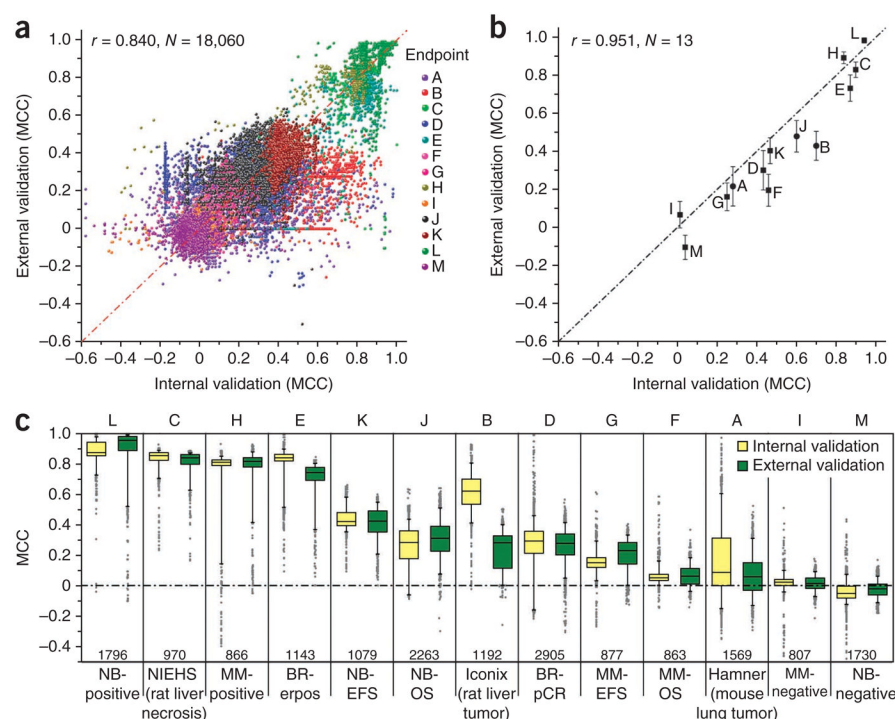
<sup>95</sup>Division of Pulmonary and Critical Care Medicine, Department of Medicine, University of Toledo Health Sciences Campus, Toledo, Ohio, USA

<sup>96</sup>Bristol-Myers Squibb, Pennington, New Jersey, USA

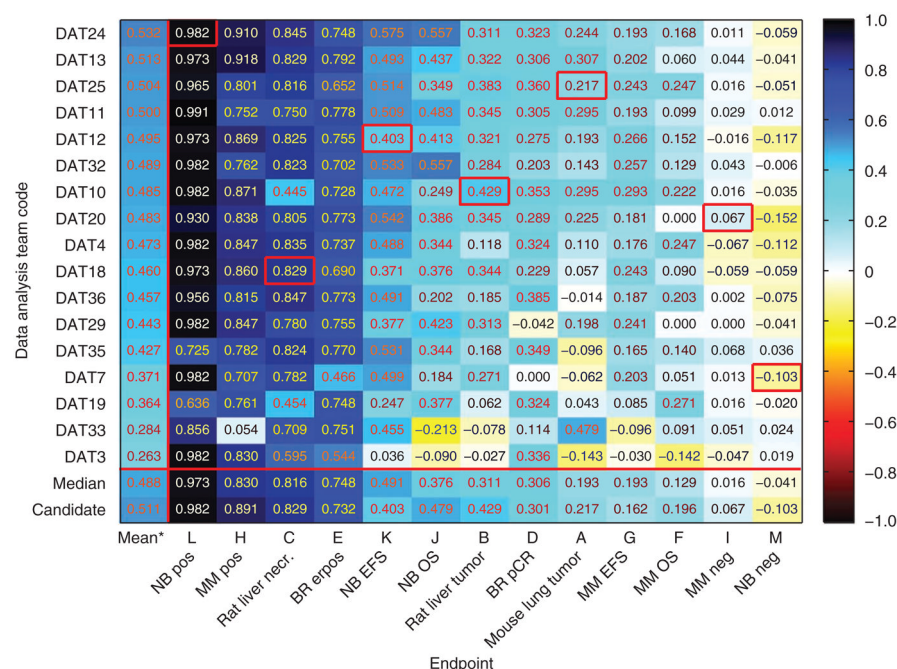
<sup>97</sup>OpGen Inc., Gaithersburg, Maryland, USA

**Figure 1.**

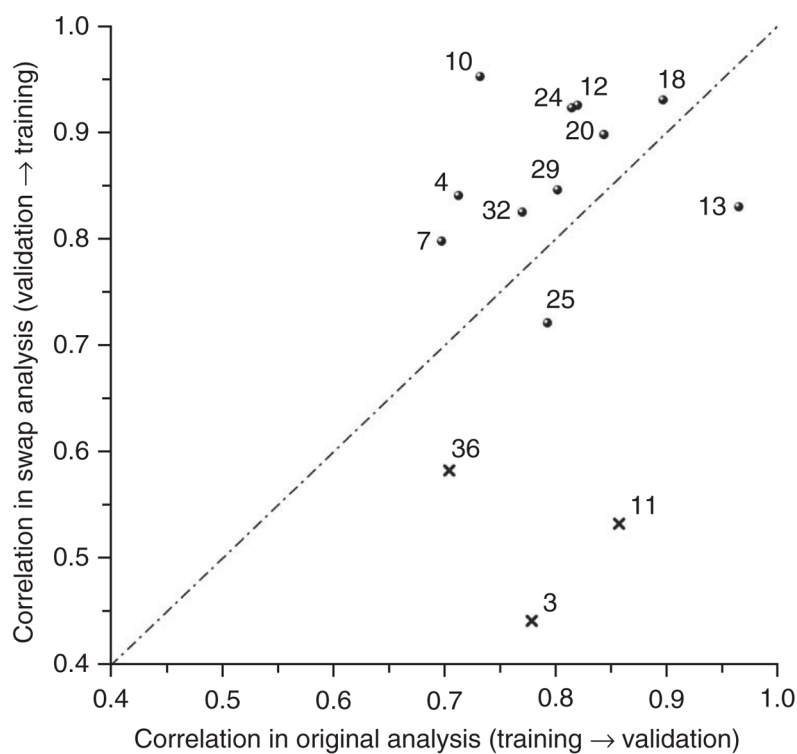
Experimental design and timeline of the MAQC-II project. Numbers (1–11) order the steps of analysis. Step 11 indicates when the original training and validation data sets were swapped to repeat steps 4–10. See main text for description of each step. Every effort was made to ensure the complete independence of the validation data sets from the training sets. Each model is characterized by several modeling factors and seven internal and external validation performance metrics (Supplementary Tables 1 and 2). The modeling factors include: (i) organization code; (ii) data set code; (iii) endpoint code; (iv) summary and normalization; (v) feature selection method; (vi) number of features used; (vii) classification algorithm; (viii) batch-effect removal method; (ix) type of internal validation; and (x) number of iterations of internal validation. The seven performance metrics for internal validation and external validation are: (i) MCC; (ii) accuracy; (iii) sensitivity; (iv) specificity; (v) AUC; (vi) mean of sensitivity and specificity; and (vii) r.m.s.e. s.d. of metrics are also provided for internal validation results.



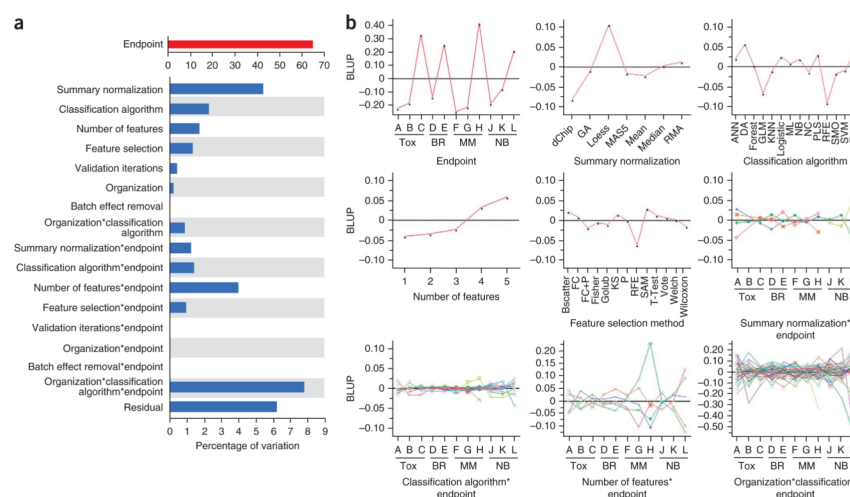
**Figure 2.** Model performance on internal validation compared with external validation. **(a)** Performance of 18,060 models that were validated with blinded validation data. **(b)** Performance of 13 candidate models.  $r$ , Pearson correlation coefficient;  $N$ , number of models. Candidate models with binary and continuous prediction values are marked as circles and squares, respectively, and the standard error estimate was obtained using 500-times resampling with bagging of the prediction results from each model. **(c)** Distribution of MCC values of all models for each endpoint in internal (left, yellow) and external (right, green) validation performance. Endpoints H and L (sex of the patients) are included as positive controls and endpoints I and M (randomly assigned sample class labels) as negative controls. Boxes indicate the 25% and 75% percentiles, and whiskers indicate the 5% and 95% percentiles.



**Figure 3.** Performance, measured using MCC, of the best models nominated by the 17 data analysis teams (DATs) that analyzed all 13 endpoints in the original training-validation experiment. The median MCC value for an endpoint, representative of the level of predicability of the endpoint, was calculated based on values from the 17 data analysis teams. The mean MCC value for a data analysis team, representative of the team's proficiency in developing predictive models, was calculated based on values from the 11 non-random endpoints (excluding negative controls I and M). Red boxes highlight candidate models. Lack of a red box in an endpoint indicates that the candidate model was developed by a data analysis team that did not analyze all 13 endpoints.



**Figure 4.** Correlation between internal and external validation is dependent on data analysis team. Pearson correlation coefficients between internal and external validation performance in terms of MCC are displayed for the 14 teams that submitted models for all 13 endpoints in both the original (x axis) and swap (y axis) analyses. The unusually low correlation in the swap analysis for DAT3, DAT11 and DAT36 is a result of their failure to accurately predict the positive endpoint H, likely due to operator errors (Supplementary Table 6).



**Figure 5.**

Effect of modeling factors on estimates of model performance. **(a)** Random-effect models of external validation performance (MCC) were developed to estimate a distinct variance component for each modeling factor and several selected interactions. The estimated variance components were then divided by their total in order to compare the proportion of variability explained by each modeling factor. The endpoint code contributes the most to the variability in external validation performance. **(b)** The BLUP plots of the corresponding factors having proportion of variation larger than 1% in **a**. Endpoint abbreviations (Tox., preclinical toxicity; BR, breast cancer; MM, multiple myeloma; NB, neuroblastoma). Endpoints H and L are the sex of the patient. Summary normalization abbreviations (GA, genetic algorithm; RMA, robust multichip analysis). Classification algorithm abbreviations (ANN, artificial neural network; DA, discriminant analysis; Forest, random forest; GLM, generalized linear model; KNN, K-nearest neighbors; Logistic, logistic regression; ML, maximum likelihood; NB, Naïve Bayes; NC, nearest centroid; PLS, partial least squares; RFE, recursive feature elimination; SMO, sequential minimal optimization; SVM, support vector machine; Tree, decision tree). Feature selection method abbreviations (Bscatter, between-class scatter; FC, fold change; KS, Kolmogorov-Smirnov algorithm; SAM, significance analysis of microarrays).

Table 1

Microarray data sets used for model development and validation in the MAQC-II project

| Date set code         | Endpoint code | Endpoint description                                                                                                  | Microarray platform                             | Training set <sup>d</sup> |               |               |           | Validation set <sup>d</sup> |               |               |           | Comments and references                                                                                                                                                                                                                                                                                                                                                      |
|-----------------------|---------------|-----------------------------------------------------------------------------------------------------------------------|-------------------------------------------------|---------------------------|---------------|---------------|-----------|-----------------------------|---------------|---------------|-----------|------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
|                       |               |                                                                                                                       |                                                 | Number of samples         | Positives (P) | Negatives (N) | P/N ratio | Number of samples           | Positives (P) | Negatives (N) | P/N ratio |                                                                                                                                                                                                                                                                                                                                                                              |
| Hamner                | A             | Lung tumorigen vs. non-tumorigen (mouse)                                                                              | Affymetrix Mouse 430 2.0                        | 70                        | 26            | 44            | 0.59      | 88                          | 28            | 60            | 0.47      | The training set was first published in 2007 (ref. 50) and the validation set was generated for MAQC-II                                                                                                                                                                                                                                                                      |
|                       | B             | Non-genotoxic liver carcinogens vs. non-carcinogens (rat)                                                             | Amersham Uniset Rat 1 Bioarray                  | 216                       | 73            | 143           | 0.51      | 201                         | 57            | 144           | 0.40      |                                                                                                                                                                                                                                                                                                                                                                              |
|                       | C             | Liver toxicants vs. non-toxicants based on overall necrosis score (rat)                                               | Affymetrix Rat 230 2.0                          | 214                       | 79            | 135           | 0.58      | 204                         | 78            | 126           | 0.62      |                                                                                                                                                                                                                                                                                                                                                                              |
| Breast cancer (BR)    | D             | Pre-operative treatment response (pCR, pathologic complete response)                                                  | Affymetrix Human U133A                          | 130                       | 33            | 97            | 0.34      | 100                         | 15            | 85            | 0.18      | The training set was first published in 2006 (ref. 56) and the validation set was specifically generated for MAQC-II. In addition, two distinct endpoints (D and E) were analyzed in MAQC-II                                                                                                                                                                                 |
|                       | E             | Estrogen receptor status (erpos)                                                                                      |                                                 | 130                       | 80            | 50            | 1.6       | 100                         | 61            | 39            | 1.56      |                                                                                                                                                                                                                                                                                                                                                                              |
|                       | F             | Overall survival milestone outcome (OS, 730-d cutoff)                                                                 | Affymetrix Human U133Plus 2.0                   | 340                       | 51            | 289           | 0.18      | 214                         | 27            | 187           | 0.14      |                                                                                                                                                                                                                                                                                                                                                                              |
| Multiple myeloma (MM) | G             | Event-free survival milestone outcome (EFS, 730-d cutoff)                                                             |                                                 | 340                       | 84            | 256           | 0.33      | 214                         | 34            | 180           | 0.19      | The data set was first published in 2006 (ref. 57) and 2007 (ref. 58). However, patient survival data were updated and the raw microarray data (CEL files) were provided specifically for MAQC-II data analysis. In addition, endpoints H and I were designed and analyzed specifically in MAQC-II                                                                           |
|                       | H             | Clinical parameter S1 (CPS1). The actual class label is the sex of the patient. Used as a “positive” control endpoint |                                                 | 340                       | 194           | 146           | 1.33      | 214                         | 140           | 74            | 1.89      |                                                                                                                                                                                                                                                                                                                                                                              |
|                       | I             | Clinical parameter R1 (CPR1). The actual class label is randomly assigned. Used as a “negative” control endpoint      |                                                 | 340                       | 200           | 140           | 1.43      | 214                         | 122           | 92            | 1.33      |                                                                                                                                                                                                                                                                                                                                                                              |
| Neuro-blastoma (NB)   | J             | Overall survival milestone outcome (OS, 900-d cutoff)                                                                 | Different versions of Agilent human microarrays | 238                       | 22            | 216           | 0.10      | 177                         | 39            | 138           | 0.28      | The training data set was first published in 2006 (ref. 63). The validation set (two-color Agilent platform) was generated specifically for MAQC-II. In addition, one-color Agilent platform data were also generated for most samples used in the training and validation sets specifically for MAQC-II to compare the prediction performance of two-color versus one-color |

| Data set code | Endpoint code | Endpoint description                                                                                                           | Microarray platform | Training set <sup>d</sup> |               |               |           | Validation set <sup>d</sup> |               |               |           | Comments and references                                                                                                               |
|---------------|---------------|--------------------------------------------------------------------------------------------------------------------------------|---------------------|---------------------------|---------------|---------------|-----------|-----------------------------|---------------|---------------|-----------|---------------------------------------------------------------------------------------------------------------------------------------|
|               |               |                                                                                                                                |                     | Number of samples         | Positives (P) | Negatives (N) | P/N ratio | Number of samples           | Positives (P) | Negatives (N) | P/N ratio |                                                                                                                                       |
| K             |               | Event-free survival milestone outcome (EFS, 900-d cutoff)                                                                      |                     | 239                       | 49            | 190           | 0.26      | 193                         | 83            | 110           | 0.75      | platforms. Patient survival data were also updated. In addition, endpoints L and M were designed and analyzed specifically in MAQC-II |
|               |               |                                                                                                                                |                     |                           |               |               |           |                             |               |               |           |                                                                                                                                       |
| L             |               | Newly established parameter S (NEP_S). The actual class label is the sex of the patient. Used as a “positive” control endpoint |                     | 246                       | 145           | 101           | 1.44      | 231                         | 133           | 98            | 1.36      |                                                                                                                                       |
| M             |               | Newly established parameter R (NEP_R). The actual class label is randomly assigned. Used as a “negative” control endpoint      |                     | 246                       | 145           | 101           | 1.44      | 253                         | 143           | 110           | 1.30      |                                                                                                                                       |
|               |               |                                                                                                                                |                     |                           |               |               |           |                             |               |               |           |                                                                                                                                       |
|               |               |                                                                                                                                |                     |                           |               |               |           |                             |               |               |           |                                                                                                                                       |

The first three data sets (Hamner, Iconix and NIEHS) are from preclinical toxicogenomics studies, whereas the other three data sets are from clinical studies. Endpoints H and L are positive controls (sex of patient) and endpoints I and M are negative controls (randomly assigned class labels). The nature of H, I, L and M was unknown to MAQC-II participants except for the project leader until all calculations were completed.

<sup>d</sup>Numbers shown are the actual number of samples used for model development or validation.

**Table 2**

Modeling factor options frequently adopted by MAQC-II data analysis teams

| Modeling factor           | Option     | Original analysis (training => validation) |                     |                  |
|---------------------------|------------|--------------------------------------------|---------------------|------------------|
|                           |            | Number of teams                            | Number of endpoints | Number of models |
| Summary and normalization | Loess      | 12                                         | 3                   | 2,563            |
|                           | RMA        | 3                                          | 7                   | 46               |
|                           | MAS5       | 11                                         | 7                   | 4,947            |
| Batch-effect removal      | None       | 10                                         | 11                  | 2,281            |
|                           | Mean shift | 3                                          | 11                  | 7,279            |
| Feature selection         | SAM        | 4                                          | 11                  | 3,771            |
|                           | FC+P       | 8                                          | 11                  | 4,711            |
|                           | T-Test     | 5                                          | 11                  | 400              |
|                           | RFE        | 2                                          | 11                  | 647              |
| Number of features        | 0~9        | 10                                         | 11                  | 393              |
|                           | 10~99      | 13                                         | 11                  | 4,445            |
|                           | ≥1,000     | 3                                          | 11                  | 474              |
|                           | 100~999    | 10                                         | 11                  | 4,298            |
| Classification algorithm  | DA         | 4                                          | 11                  | 103              |
|                           | Tree       | 5                                          | 11                  | 358              |
|                           | NB         | 4                                          | 11                  | 924              |
|                           | KNN        | 8                                          | 11                  | 6,904            |
|                           | SVM        | 9                                          | 11                  | 986              |

Analytic options used by two or more of the 14 teams that submitted models for all endpoints in both the original and swap experiments. RMA, robust multichip analysis; SAM, significance analysis of microarrays; FC, fold change; RFE, recursive feature elimination; DA, discriminant analysis; Tree, decision tree; NB, naive Bayes; KNN, K-nearest neighbors; SVM, support vector machine.