

MULTIPLE IMPROVEMENTS OF MULTIPLE IMPUTATION LIKELIHOOD RATIO TESTS

BY KIN WAI CHAN AND XIAO-LI MENG

The Chinese University of Hong Kong and Harvard University

Multiple imputation (MI) inference handles missing data by first properly imputing the missing values m times, and then combining the results from the m complete-data analyses. However, the existing method for combining likelihood ratio tests has multiple defects: (i) the combined test statistic can be negative in practice but its null distribution is approximated by a standard F distribution; (ii) it is not invariant to re-parametrization; (iii) it fails to ensure monotonic power due to its use of an inconsistent estimator of the fraction of missing information (FMI) under the alternative hypothesis; and (iv) it requires non-trivial access to the likelihood ratio test statistic as a function of estimated parameters instead of datasets. This paper shows, via both theoretical derivations and empirical investigations, that essentially all of these problems can be straightforwardly addressed if we are willing to perform an additional likelihood ratio test by stacking the m completed datasets as one big completed dataset. A particularly intriguing finding is that the FMI itself can be estimated consistently by a likelihood ratio statistic for testing whether the m completed datasets produced by MI can be regarded effectively as samples coming from a common model. Practical guidelines are provided based on an extensive comparison of existing MI tests. Intrigued issues regarding nuisance parameters are also discussed.

1. Historical Successes and Failures. Missing-data problems are ubiquitous in practice, to the extent that the absence of any missingness often is a strong indication that the data have been pre-processed or manipulated in some way (e.g., [Blocker and Meng, 2013](#)). Multiple imputation (MI) ([Rubin, 1978, 2004](#)) has been a preferred method by many practitioners, especially those who are ill-equipped to handle missingness on their own, due to lack of information or skills or resources. MI relies on the data collector (e.g., a census bureau) to build a reliable imputation model to fill in the missing data $m(\geq 2)$ times, so the users of the data can apply their favorite software or procedures that are designed to handle complete data,

AMS 2000 subject classifications: Primary 62D05; Secondary 62F03, 62E20.

Keywords and phrases: Fraction of missing information; missing data; invariant test; monotonic power; robust estimation.

Meng thanks NSF and JTF for partial financial support. He is also embarrassed by the research immaturity displayed in his thesis, and is very thankful to Keith's (Kin Wai) creativity and diligence which led to the beautiful remedies presented here.

and do so m times. MI inference, e.g., hypothesis testing, is then performed by appropriately combining these m complete-data results.

Although MI was designed initially for public-use datasets, over the past 30 years or so, it has become a method of choice for handling missing data in general, because it separates the handling missingness from conducting analysis (e.g., [Tu *et al.*, 1993](#); [Rubin, 1996, 2004](#); [Schafer, 1999](#); [King *et al.*, 2001](#); [Peugh and Enders, 2004](#); [Kenward and Carpenter, 2007](#); [Rose and Fraser, 2008](#); [Holan *et al.*, 2010](#); [Kim and Yang, 2017](#)). Software routines for performing MI are now available in R ([van Buuren and Groothuis-Oudshoorn, 2011](#); [Su *et al.*, 2011](#)), Stata ([Royston and White, 2011](#)), SAS ([Berglund and Heeringa, 2014](#)) and SPSS; also see [Harel and Zhou \(2007\)](#) and [Horton and Kleinman \(2007\)](#) for summaries on software that utilize MI.

This convenient separation, however, creates the thorny issue of uncongeniality, i.e., the incompatibility between the imputation model and the subsequent analysis procedures ([Meng, 1994a](#)). This issue is examined in detail by [Xie and Meng \(2017\)](#), which shows that uncongeniality is easiest to deal with when the imputer’s model is more saturated than the user’s model/procedure, and when the user is conducting efficient analysis, such as likelihood inference. The current paper, therefore, focuses on conducting MI likelihood ratio tests (LRTs), assuming the imputation model is sufficiently saturated to render the validity of the common assumptions, which we shall review, made in the literature about conducting LRTs with MI.

Like many hypothesis testing procedures in common practice, the exact null distributions of various MI test statistics, LRTs or not, are intractable. This intractability is not computational, but rather statistical due to the well-known issue of nuisance parameter, that is, the lack of pivotal quantity, as highlighted, historically, by the Behrens-Fisher problem ([Wallace, 1980](#)). Indeed, the nuisance parameter in the MI context is the so-called “the fraction of missing information” (FMI), which is determined by the ratio of the between-imputation variance to within-imputation variance (and its multivariate counterparts), and hence the challenge we face is almost identical to the one faced by the Behrens-Fisher problem, as shown in [Meng \(1994b\)](#).

An added challenge in the MI context is that the user’s complete-data procedures can be very restrictive. What is available to the user could vary from the entire likelihood function, to point estimators such as MLE and Fisher information, to a single p -value. Therefore, there have been a variety of procedures proposed in the literature, depending on what quantities we assume the user has access to, as we shall review shortly.

Among them, a promising idea was to directly combine LRT statistics. However, the execution of this idea as presented in [Meng and Rubin \(1992\)](#)

relied too heavily on the usual asymptotic equivalence (in terms of the data size, not the number of imputations, m) between the LRT and Wald test *under the null*. Its asymptotic validity, unfortunately, does not protect it from quick deterioration for small data sizes, such as delivering negative “ F test statistic” or FMI. Worst of all, the test can have essentially zero power because the estimator of FMI can be badly inconsistent under some alternative hypotheses. In addition, the combining rule of [Meng and Rubin \(1992\)](#) requires the user to have access to the LRT as a function of parameter values, not just as a function of the data. The former one is often unavailable from standard software packages. This defective MI LRT, however, has been adopted by textbooks (e.g., [van Buuren S, 2012](#); [Kim and Shao, 2013](#)) and popular software, e.g., the function `pool.compare` in the R package `mice` ([van Buuren and Groothuis-Oudshoorn, 2011](#)), the function `testModels` in the R package `mitml` ([Grund et al., 2017](#)), the function `milrtest` ([Medeiros, 2008](#)) in the Stata module `mim` ([Carlin et al., 2008](#)).

To minimize the negative impact of this defective LRT test, this paper derives MI LRTs that are free of the defects as outlined in the abstract and detailed in § 1.3 below. We achieve this mainly by switching the order of two main operators in the combining rule of [Meng and Rubin \(1992\)](#): Maximizing the average of the m log-likelihoods instead of averaging the maximizers of them. This switching, guided by the likelihood principle, automatically renders positivity, invariance and monotonic power. Other judicious uses of the likelihood functions permit us to overcome the remaining defects.

The remainder of Section 1 provides background and notation. Section 2 then discusses the defects of the existing MI LRT and our remedies. Section 3 investigates computational requirements for our proposals, including theoretical considerations and comparisons. In particular, Algorithm 1 of Section 3.1 computes our most recommended test. Section 4 provides empirical evidence with simulated and real data. Section 5 calls for future work. Appendices A–C supplement with proofs, and additional investigations.

1.1. Notation and Complete-data Tests. Let X_{obs} and X_{mis} be, respectively, the observed and missing parts of an intended complete dataset $X = X_{\text{com}} = \{X_{\text{obs}}, X_{\text{mis}}\}$, which consists of n observations. Denote the sampling model — probability or density, depending on the data type — of X by $f(\cdot \mid \psi)$, where $\psi \in \Psi \subset \mathbb{R}^h$ is a vector of parameters. The goal is to test $H_0 : \theta = \theta_0$ when only X_{obs} is available, where $\theta = \theta(\psi) \in \Theta \subset \mathbb{R}^k$ is a function of ψ , and θ_0 is a specified vector. For simplicity, we will focus on the standard two-sided alternative, but our approach adapts to general complete-data LRTs. Denote the true values of ψ and θ by ψ^* and θ^* . Here

we assume X_{obs} is rich enough that the missing data mechanism is ignorable (Rubin, 1976), or it has been properly incorporated into the imputation model by the imputer, who may have access to additional confidential data.

Let $\hat{\theta} = \hat{\theta}(X)$ and $U = U_{\theta} = U_{\theta}(X)$ be respectively the complete-data MLE of θ and an efficient estimator of $\text{Var}(\hat{\theta})$ (e.g., the inverse of the observed Fisher information). Also, let $\hat{\psi}_0 = \hat{\psi}_0(X)$ and $\hat{\psi} = \hat{\psi}(X)$ be respectively the H_0 -constrained and unconstrained complete-data MLEs of ψ , and $U_{\psi} = U_{\psi}(X)$ be an efficient estimator of $\text{Var}(\hat{\psi})$. For testing H_0 against H_1 , the common choices include the Wald statistic $D_W = d_W(\hat{\theta}, U)/k$ and the LRT statistic $D_L = d_L(\hat{\psi}_0, \hat{\psi} \mid X)/k$, where

$$d_W(\hat{\theta}, U) = (\hat{\theta} - \theta_0)^{\top} U^{-1} (\hat{\theta} - \theta_0), \quad d_L(\hat{\psi}_0, \hat{\psi} \mid X) = 2 \log \frac{f(X \mid \hat{\psi})}{f(X \mid \hat{\psi}_0)}.$$

Under regularity conditions (RCs), such as those in § 4.2.2 and § 4.4.2 of Serfling (2001) when the rows of X are independent and identically distributed, we have the following classical results.

PROPERTY 1.1. *Under H_0 , (i) $D_W \Rightarrow \chi_k^2/k$ and $D_L \Rightarrow \chi_k^2/k$; and (ii) $n(D_W - D_L) \xrightarrow{\text{Pr}} 0$ as $n \rightarrow \infty$ where “ $\Rightarrow$ ” and “ $\xrightarrow{\text{Pr}}$ ” denote convergence in distribution and in probability, respectively.*

Testing $\theta = \theta_0$ based on X_{obs} is more involved. For MI, let $X^{\ell} = \{X_{\text{obs}}, X_{\text{mis}}^{\ell}\}$, $\ell = 1, \dots, m$, be the m completed datasets, where X_{mis}^{ℓ} are drawn conditionally independently from a proper imputation model given X_{obs} ; see Rubin (2004). We then carry out a complete-data estimation or testing procedure on X^{ℓ} , $\ell = 1, \dots, m$, resulting in a set of m quantities. The so-called MI inference is to appropriately combine them to obtain a single answer.

1.2. MI Wald Test and Fraction of Missing Information. Let $d_W^{\ell} = d_W(\hat{\theta}^{\ell}, U^{\ell})$, $\hat{\theta}^{\ell} = \hat{\theta}(X^{\ell})$ and $U^{\ell} = U(X^{\ell})$ be the imputed counterparts of $d_W(\hat{\theta}, U)$, $\hat{\theta}$ and U , respectively, for each ℓ . Also, write their averages as

$$(1.1) \quad \bar{d}_W = \frac{1}{m} \sum_{\ell=1}^m d_W^{\ell}, \quad \bar{\theta} = \frac{1}{m} \sum_{\ell=1}^m \hat{\theta}^{\ell}, \quad \bar{U} = \frac{1}{m} \sum_{\ell=1}^m U^{\ell}.$$

Under congeniality (Meng, 1994a), one can show that asymptotically (Rubin and Schenker, 1986) $\text{Var}(\bar{\theta})$ can be consistently estimated by

$$(1.2) \quad T = \bar{U} + (1 + 1/m)B, \quad \text{where} \quad B = \frac{1}{m-1} \sum_{\ell=1}^m (\hat{\theta}^{\ell} - \bar{\theta})(\hat{\theta}^{\ell} - \bar{\theta})^{\top}$$

is known as the *between-imputation variance*, in contrast to $\bar{U}$ in (1.1), which measures *within-imputation variance*. Intriguingly, $2T$ serves as a universal (estimated) upper bound of $\text{Var}(\bar{\theta})$ under uncongeniality (Xie and Meng, 2017). Under RCs, we have that, as $m, n \rightarrow \infty$,

$$n(\bar{U} - \mathcal{U}_\theta) \xrightarrow{\text{pr}} \mathbf{0}, \quad n(T - \mathcal{T}_\theta) \xrightarrow{\text{pr}} \mathbf{0}, \quad n(B - \mathcal{B}_\theta) \xrightarrow{\text{pr}} \mathbf{0}$$

for some deterministic matrices $\mathcal{U}_\theta$, $\mathcal{T}_\theta$ and $\mathcal{B}_\theta = \mathcal{T}_\theta - \mathcal{U}_\theta$, where $\mathbf{0}$ denotes a matrix of zeros, and the subscript θ highlights that these matrices are for estimating θ , because there are also corresponding $\mathcal{T}_\psi, \mathcal{B}_\psi, \mathcal{U}_\psi$ for the entire parameter ψ . Similar to $\bar{U}$, T and B , we define $\bar{U}_\psi, T_\psi$ and B_ψ for the component ψ . Note that if $\hat{\theta}_{\text{com}}$ and $\hat{\theta}_{\text{obs}}$ are the MLEs of θ based on X_{com} and X_{obs} (under congeniality), respectively, then $\mathcal{U}_\theta \simeq \text{Var}(\hat{\theta}_{\text{com}})$ and $\mathcal{T}_\theta \simeq \text{Var}(\hat{\theta}_{\text{obs}})$ as $n \rightarrow \infty$, where $A_n \simeq B_n$ means that $A_n - B_n = o_p(A_n + B_n)$.

The straightforward MI Wald test $D_W(T) = d_W(\bar{\theta}, T)/k$ is not practical because T is singular when $m < k$ (usually $3 \leq m \leq 10$). Even when it is not singular, it is usually not a very stable estimator of $\mathcal{T}_\theta$ because m is small. To circumvent this problem, Rubin (1978) adopted the following assumption of equal fraction of missing information (EFMI).

ASSUMPTION 1 (EFMI of θ). *There is $\nu \geq 0$ such that $\mathcal{T}_\theta = (1 + \nu)\mathcal{U}_\theta$.*

EFMI clearly is a very strong assumption, implying that the missing data have caused an equal amount of loss of information for estimating every component of θ . However, as we shall see shortly, the adoption of this assumption, *for the purpose of hypothesis testing*, is essentially the same as to summarize the impact of (at least) k nuisance parameters due to FMI by a single nuisance parameter, i.e., the average FMI across different components. How well this reduction strategy works therefore will affect more the power of the test than its validity, as long as we can construct an approximate null distribution that is more robust to the EFMI assumption. The issue of power turns out to be a rather tricky one, because without the reduction strategy we would also lose power when m/k is small or even modest. It is because we simply do not have enough degrees of freedom to estimate all the nuisance parameters well or at all. We will illustrate this point in § 4.2. (To clarify some confusions in literature, ν in Assumption 1 is the *odds of the missing information*, not the FMI, which is $\ell = \nu/(1 + \nu)$.)

Under EFMI, Rubin (2004) replaced T by $(1 + \tilde{r}'_W)\bar{U}$, where

$$(1.3) \quad \tilde{r}'_W = \frac{(m+1)}{k(m-1)}(\bar{d}'_W - \tilde{d}'_W); \quad \bar{d}'_W = \frac{1}{m} \sum_{\ell=1}^m d_W(\hat{\theta}^\ell, \bar{U}),$$

$\tilde{d}'_W = d_W(\bar{\theta}, \bar{U})$, and the prime “ r ” indicates that $\bar{U}$ is used instead of individual $\{U^\ell\}_{\ell=1}^m$. Then, [Rubin \(2004\)](#) proposed a simple MI Wald test statistic:

$$(1.4) \quad \tilde{D}'_W = \frac{\tilde{d}'_W}{k(1 + \tilde{r}'_W)}.$$

The intuition behind (1.3)–(1.4) is important because the forms here are the building blocks for virtually all the subsequent developments. The “obvious” Wald test statistic $\tilde{d}'_W/k$ is too large (compared to the usual χ^2_k/k) because it fails to take into account of missing information. The $(1 + \tilde{r}'_W)$ factor attempts to correct this, with the amount of correction determined by the (average) amount of between-imputation variance relative to the within-imputation variance. Expression (1.3) shows that this relative amount can be estimated by contrasting the average of individual Wald statistics and the Wald statistic based on an average of individual estimates. Using the difference between “average of functions” and “function of average”, namely,

$$(1.5) \quad \text{Ave}\{G(x)\} - G(\text{Ave}\{x\})$$

is a common practice, e.g., $G(x) = x^2$ for variance; see [Meng \(2002\)](#).

Since the exact null distribution of $\tilde{D}'_W$ is intractable, [Li et al. \(1991b\)](#) proposed to approximate it by $F_{k, \tilde{\text{df}}(\tilde{r}'_W, k)}$, the F distribution with degrees of freedom k and $\tilde{\text{df}}(\tilde{r}'_W, k)$, where, denoting $K_m = k(m-1)$,

$$(1.6) \quad \tilde{\text{df}}(\tilde{r}'_W, k) = \begin{cases} 4 + (K_m - 4)\{1 + (1 - 2/K_m)/\tilde{r}'_W\}^2, & \text{if } K_m > 4; \\ (m-1)(1 + 1/\tilde{r}'_W)^2(k+1)/2, & \text{otherwise.} \end{cases}$$

This approximation assumes n is sufficiently large so that the standard asymptotic χ^2 distribution in Property 1.1 can be used. If n is small, the small sample degree of freedom in [Barnard and Rubin \(1999\)](#) should be used.

1.3. The Current MI Likelihood Ratio Test and Its Defect. Let $d_L^\ell = d_L(\hat{\psi}_0^\ell, \hat{\psi}^\ell \mid X^\ell)$, $\hat{\psi}_0^\ell = \hat{\psi}_0(X^\ell)$ and $\hat{\psi}^\ell = \hat{\psi}(X^\ell)$ be the imputed counterparts of $d_L(\hat{\psi}_0, \hat{\psi} \mid X)$, $\hat{\psi}_0$ and $\hat{\psi}$, respectively, for each ℓ . Let their averages be

$$(1.7) \quad \bar{d}_L = \frac{1}{m} \sum_{\ell=1}^m d_L^\ell, \quad \bar{\psi}_0 = \frac{1}{m} \sum_{\ell=1}^m \hat{\psi}_0^\ell, \quad \bar{\psi} = \frac{1}{m} \sum_{\ell=1}^m \hat{\psi}^\ell.$$

Similar to $\tilde{r}'_W$, [Meng and Rubin \(1992\)](#) proposed to estimate $\tilde{r}_m$ by

$$(1.8) \quad \tilde{r}_L = \frac{m+1}{k(m-1)}(\bar{d}_L - \tilde{d}_L), \quad \text{where} \quad \tilde{d}_L = \frac{1}{m} \sum_{\ell=1}^m d_L(\bar{\psi}_0, \bar{\psi} \mid X^\ell),$$

and hence it is again in the form of (1.5). Computation of $\tilde{r}_L$ requires users to have access to (i) a subroutine for $(X, \psi_0, \psi) \mapsto d_L(\psi_0, \psi \mid X)$, and (ii) the estimates $\hat{\psi}_0^\ell$ and $\hat{\psi}^\ell$, rather than the matrices $\bar{U}$ and B . Therefore computing $\tilde{r}_L$ is easier than computing $\tilde{r}_W$. The resulting MI LRT is

$$(1.9) \quad \tilde{D}_L = \frac{\tilde{d}_L}{k(1 + \tilde{r}_L)},$$

whose null distribution can be approximated by $F_{k, \text{df}(\tilde{r}_L, k)}$.

The main theoretical justification (and motivation) was the asymptotic equivalence between the complete-data Wald test statistic and LRT statistic *under the null*, as stated in Property 1.1. This equivalence permitted the replacement of $\bar{d}_W$ and $\tilde{d}_W$ in (1.3) respectively by $\bar{d}_L$ and $\tilde{d}_L$ in (1.8). However, this is also where the problems lie.

First, with finite samples, $0 \leq \tilde{d}_L \leq \bar{d}_L$ is not guaranteed, consequently nor is $\tilde{D}_L \geq 0$ or $\tilde{r}_L \geq 0$. Since $\tilde{D}_L$ is referred to an F distribution and $\tilde{r}_L$ estimates $\kappa_m \geq 0$, clearly negative values of $\tilde{D}_L$ or $\tilde{r}_L$ will cause trouble.

Second, $\tilde{D}_L$ is not invariant to re-parameterization of ψ . For each individual LRT statistic d_L^ℓ and bijective map g such that $\varphi = g(\psi)$, we have

$$(1.10) \quad d_L^\ell = d_L(\hat{\psi}_0^\ell, \hat{\psi}^\ell \mid X^\ell) = d_L(g^{-1}(\hat{\varphi}_0^\ell), g^{-1}(\hat{\varphi}^\ell) \mid X^\ell),$$

where $\hat{\varphi}_0^\ell$ and $\hat{\varphi}^\ell$ are the constrained and unconstrained MLEs of φ based on X^ℓ . However, the MI LRT statistic $\tilde{d}_L$ no longer has this property because

$$\tilde{d}_L = \frac{1}{m} \sum_{\ell=1}^m d_L(\bar{\psi}_0, \bar{\psi} \mid X^\ell) \neq \frac{1}{m} \sum_{\ell=1}^m d_L(g^{-1}(\bar{\varphi}_0), g^{-1}(\bar{\varphi}) \mid X^\ell)$$

for most bijective maps g , where $\bar{\varphi}_0 = m^{-1} \sum_{\ell=1}^m \hat{\varphi}_0^\ell$ and $\bar{\varphi} = m^{-1} \sum_{\ell=1}^m \hat{\varphi}^\ell$. See § 4 how $\tilde{D}_L$ vary dramatically with parametrizations in finite samples.

Third, the estimator $\tilde{r}_L$ involves the estimators of ψ under H_0 , i.e., $\hat{\psi}_0^\ell$ and $\bar{\psi}_0$. When H_0 fails, they may not be consistent for ψ . As a result, $\tilde{r}_L$ is no longer consistent for κ_m . A serious consequence is that the power of the test statistic $\tilde{D}_L$ is not guaranteed to monotonically increase as H_1 moves away from H_0 . Indeed our simulations (see § 3.2) show that under certain parametrizations, the power may nearly vanish for obviously false H_0 .

Fourth, in order to compute $\tilde{d}_L$ in (1.8), users need to have access to the LRT function $\tilde{\mathcal{D}}_L$, but, in most software, the function is built as $\mathcal{D}_L$, where

$$(1.11) \quad \tilde{\mathcal{D}}_L : (X, \psi_0, \psi) \mapsto d_L(\psi_0, \psi \mid X), \quad \mathcal{D}_L : X \mapsto d_L(\hat{\psi}_0(X), \hat{\psi}(X) \mid X).$$

Hence users would need to write themselves a subroutine for evaluating $\tilde{\mathcal{D}}_L$. This may not be feasible for users because of lack of information or skills.

In short, four problems need to be resolved: (i) lack of non-negativity, (ii) lack of invariance, (iii) lack of consistency and power, and (iv) lack of a computationally feasible algorithm. Problems (i)–(iii) are resolved in § 2 below, where § 2.1 presents an invariant combining rule, which fully resolves (ii). Next, we propose two estimators of $\boldsymbol{\pi}_m$ (or equivalently $\boldsymbol{\pi}$) in § 2.2 and § 2.4. We start with a quick ad hoc fix that requires slightly less assumption but only addresses (i), and then construct a test that fully resolves (i) and (iii). Finally, in § 3, we derive a very handy algorithm, which resolves (iv).

2. Improved MI Likelihood Ratio Tests.

2.1. An Invariant Combining Rule. To derive a MI LRT that is invariant to re-parametrization, we replace $\tilde{d}_L$ by an asymptotically equivalent version that behaves like a standard LRT statistic. Specifically, let

$$(2.1) \quad \bar{L}(\psi) = \frac{1}{m} \sum_{\ell=1}^m L^\ell(\psi), \quad \text{where} \quad L^\ell(\psi) = \log f(X^\ell \mid \psi).$$

We emphasize that $\bar{L}(\psi)$ is not a real log-likelihood (even if we drop the divider m), because it does not properly model the completed datasets: $\mathbb{X} = \{X^1, \dots, X^m\}$ (e.g., addressing the issue that all X^ℓ s share the same X_{obs}). Nevertheless, $\bar{L}(\psi)$ can be treated as a log-likelihood for computational purposes. In particular, we can maximize it to obtain

$$(2.2) \quad \hat{\psi}_0^* = \hat{\psi}_0^*(\mathbb{X}) = \arg \max_{\psi \in \Psi : \theta(\psi) = \theta_0} \bar{L}(\psi), \quad \hat{\psi}^* = \hat{\psi}^*(\mathbb{X}) = \arg \max_{\psi \in \Psi} \bar{L}(\psi).$$

The corresponding log-likelihood ratio test statistic is given by

$$(2.3) \quad \hat{d}_L = 2 \left\{ \bar{L}(\hat{\psi}^*) - \bar{L}(\hat{\psi}_0^*) \right\} = \frac{1}{m} \sum_{\ell=1}^m d_L(\hat{\psi}_0^*, \hat{\psi}^* \mid X^\ell).$$

Thus, in contrast to $\tilde{d}_L$ of (1.8), $\hat{d}_L$ aggregates MI datasets through averaging MI LRT functions as in (2.1), rather than averaging MI test statistics and moments, as in (1.7). Although $\sqrt{n}(\hat{\psi}_0^* - \bar{\psi}_0) \xrightarrow{\text{pr}} \mathbf{0}$ and $\sqrt{n}(\hat{\psi}^* - \bar{\psi}) \xrightarrow{\text{pr}} \mathbf{0}$ as $n \rightarrow \infty$ for each m , only $\hat{d}_L$, not $\tilde{d}_L$, is guaranteed to be non-negative and invariant to parametrization of ψ for all m, n . Indeed, the likelihood principle guides us to consider averaging individual log-likelihoods than individual MLEs, since the former has a much better chance to capture functional features of the real log-likelihood than any of their (local) maximizers can.

To derive properties of $\hat{d}_L$, we need the usual RCs on MLE and MI.

ASSUMPTION 2. *The sampling model $f(X | \psi)$ satisfies the following:*

- (a) $\psi \mapsto \underline{L}(\psi) = n^{-1} \log f(X | \psi)$ is twice continuously differentiable;
- (b) the complete-data MLE $\hat{\psi}(X)$ is the unique solution of $\partial \underline{L}(\psi) / \partial \psi = \mathbf{0}$;
- (c) if $\underline{I}(\psi) = -\partial^2 \underline{L}(\psi) / \partial \psi \partial \psi^\top$, then for each ψ , there exists a positive definite matrix $\underline{\mathcal{J}}(\psi) = \mathcal{U}_\psi^{-1}$ such that $\underline{I}(\psi) \xrightarrow{\text{pr}} \underline{\mathcal{J}}(\psi)$ as $n \rightarrow \infty$; and
- (d) the observed-data MLE $\hat{\psi}_{\text{obs}}$ of ψ obeys

$$(2.4) \quad \left[\mathcal{J}_\psi^{-1/2} \left(\hat{\psi}_{\text{obs}} - \psi^\star \right) \mid \psi^\star \right] \Rightarrow \mathcal{N}_h(\mathbf{0}, I_h)$$

as $n \rightarrow \infty$, where I_h is the $h \times h$ identity matrix.

ASSUMPTION 3. *The imputation model is proper (Rubin, 2004):*

$$(2.5) \quad \left[\mathcal{B}_\psi^{-1/2} \left(\hat{\psi}^\ell - \hat{\psi}_{\text{obs}} \right) \mid X_{\text{obs}} \right] \Rightarrow \mathcal{N}_h(\mathbf{0}, I_h),$$

$$(2.6) \quad \left[\mathcal{T}_\psi^{-1} \left(U_\psi^\ell - \mathcal{U}_\psi \right) \mid X_{\text{obs}} \right] \xrightarrow{\text{pr}} \mathbf{0}, \quad \left[\mathcal{T}_\psi^{-1} (B_\psi - \mathcal{B}_\psi) \mid X_{\text{obs}} \right] \xrightarrow{\text{pr}} \mathbf{0}$$

independently for $\ell = 1, \dots, m$, as $n \rightarrow \infty$, provided that $\mathcal{B}_\psi^{-1}$ is well-defined.

Assumption 2 holds under the usual RCs that guarantee normality and consistency of MLEs. When the imputations $X_{\text{mis}}^1, \dots, X_{\text{mis}}^m$ are drawn independently from (correctly specified) posterior predictive distribution $f(X_{\text{mis}} | X_{\text{obs}})$, Assumption 3 is typically satisfied. Clearly, we can replace ψ by its sub-vector θ in Assumptions 2 and 3. These θ -version assumptions are sufficient to guarantee the validity of the following Theorem 2.4 and Corollary 2.3. For simplicity, Assumption 1, the θ -version of Assumptions 2 and 3, and conditions that are strong enough to guarantee Property 1.1 are collectively written as RC_θ , which are commonly assumed for MI inference.

THEOREM 2.1. *Assume RC_θ . Under H_0 , we have (i) $\hat{d}_L \geq 0$ for all m, n ; (ii) $\hat{d}_L$ is invariant to parametrization of ψ for all m, n ; and (iii) $\hat{d}_L \simeq \tilde{d}_L$ as $n \rightarrow \infty$ for each m .*

Consequently, an improved combining rule is defined as

$$(2.7) \quad \hat{D}_L(\boldsymbol{r}_m) = \frac{\hat{d}_L}{k(1 + \boldsymbol{r}_m)},$$

for a given value of $\boldsymbol{r}_m$. It follows the forms of (1.4) and of (1.9). The issue is then how to estimate $\boldsymbol{r}_m$ that avoids the defects of $\tilde{\boldsymbol{r}}_L$ of (1.8).

2.2. *An Improved Estimator of $\mathfrak{r}_m$.* Using $\hat{d}_L$ in (2.3), we can modify $\tilde{r}_L$ in (1.8) to a potentially better estimator:

$$(2.8) \quad \hat{r}_L = \frac{m+1}{k(m-1)}(\bar{d}_L - \hat{d}_L).$$

Although $\hat{d}_L \geq 0$ is guaranteed by our construction, $\hat{r}_L \geq 0$ does not hold in general for a finite m . However, it is guaranteed in the following situation.

PROPOSITION 2.2. *Write $\psi = (\theta^\top, \eta^\top)^\top$, where η represents a nuisance parameter that is distinct from θ . If there exist functions $L_\dagger$ and $L_\ddagger$ such that, for all X , the log-likelihood function $L(\psi | X) = \log f(X | \psi)$ is of the form $L(\psi | X) = L_\dagger(\theta | X) + L_\ddagger(\eta | X)$, then $\hat{r}_L \geq 0$ for all m, n .*

The condition in Proposition 2.2 means that the likelihood function of ψ is separable. Then, the profile likelihood estimator of η given θ , i.e., $\hat{\eta}_\theta = \arg \max_\eta L(\theta, \eta | X)$, does not depend on θ . Trivially, if there is no nuisance parameter η , the separation condition is satisfied. More generally, we have

COROLLARY 2.3. *Assume RC_θ . We have (i) under H_0 , $\hat{r}_L \xrightarrow{\text{pr}} \mathfrak{r}$ as $m, n \rightarrow \infty$; and (ii) under H_1 , $\hat{r}_L \xrightarrow{\text{pr}} \mathfrak{r}_0$ as $m, n \rightarrow \infty$, where $\mathfrak{r}_0 \geq 0$ is some finite value depending on θ_0 and θ^* .*

Corollary 2.3 ensures that, under H_0 , $\hat{r}_L$ is non-negative asymptotically and converges in probability to the true $\mathfrak{r}$. But it also reveals another fundamental defect of $\hat{r}_L$: under H_1 , the limit $\mathfrak{r}_0$ may not equal to $\mathfrak{r}$, a problem we will address in § 2.2. Fortunately, since $\hat{d}_L \xrightarrow{\text{pr}} \infty$ under H_1 , the LRT statistic $\hat{D}_L(\hat{r}_L)$ is still powerful, albeit the power may be somewhat reduced. Similarly, $\tilde{r}_L$ of (1.8) has the same asymptotic properties and defects, but $\hat{r}_L$ behaves more nicely than $\tilde{r}_L$ for finite m . This hinges closely on the high sensitivity of $\tilde{r}_L$ to the parametrization of ψ for small m , e.g., in some cases, $\tilde{r}_L$ becomes more negative as H_1 moves away from H_1 ; see § 4.1.

Whereas we can fix the occasional negativeness of $\hat{r}_L$ by using $\hat{r}_L^+ = \max(0, \hat{r}_L)$, such an ad hoc fix misses the opportunity to improve upon $\hat{r}_L$, and indeed it cannot fix the inconsistency of $\hat{r}_L$ under H_1 .

2.3. *A Complication Caused by Nuisance Parameter.* To better understand the source of the negativity of $\hat{r}_L$, we extend $\bar{L}(\psi)$ in (2.1) to allow it take m different arguments:

$$(2.9) \quad \bar{L}(\psi^1, \dots, \psi^m) = \frac{1}{m} \sum_{\ell=1}^m L^\ell(\psi^\ell).$$

TABLE 1
The definitions of hypotheses $H_0^0, H_0^1, H_1^0, H_1^1$.

	$\mathcal{E}^0 : \psi^1 = \dots = \psi^m \in \Psi$ (i.e., $\mathcal{r} = 0$)	$\mathcal{E}^1 : \psi^1, \dots, \psi^m \in \Psi$ (i.e., $\mathcal{r} \geq 0$)
$\mathcal{E}_0 : \theta^1 = \dots = \theta^m = \theta_0 \in \Theta$ (i.e., H_0 -constrained)	$H_0^0 = \mathcal{E}_0 \cap \mathcal{E}^0$	$H_0^1 = \mathcal{E}_0 \cap \mathcal{E}^1$
$\mathcal{E}_1 : \theta^1, \dots, \theta^m \in \Theta$ (i.e., not H_0 -constrained)	$H_1^0 = \mathcal{E}_1 \cap \mathcal{E}^0$	$H_1^1 = \mathcal{E}_1 \cap \mathcal{E}^1$

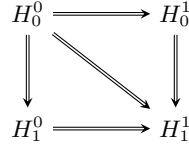


FIG 1. The relationships between the four hypotheses $H_0^0, H_0^1, H_1^0, H_1^1$. Each arrow denotes an implication, e.g., $H_0^0 \Rightarrow H_1^0$ means that H_0^0 implies H_1^0 .

Using the “log-likelihood” $\bar{L}(\psi^1, \dots, \psi^m)$, we can construct, at least conceptually, four hypotheses $H_0^0, H_0^1, H_1^0, H_1^1$ defined in Table 1. Each of them consists of zero, one or two of the constraints

$$\mathcal{E}_0 : \theta^1 = \dots = \theta^m = \theta_0 \quad \text{and} \quad \mathcal{E}^0 : \psi^1 = \dots = \psi^m.$$

The constraint $\mathcal{E}_0$ is equivalent to H_0 , and the constraint $\mathcal{E}^0$ means that all ψ^ℓ s are equal, and hence it is effectively equivalent to $\mathcal{r} = 0$, i.e., no missing information. The relationships among $H_0^0, H_0^1, H_1^0, H_1^1$ can be visualized in Figure 1. Define the maximized value of $\bar{L}(\psi^1, \dots, \psi^m)$ under hypothesis $H \in \{H_0^0, H_0^1, H_1^0, H_1^1\}$ by $\mathbb{L}(H)$. Then we can re-express $(\bar{d}_L - \hat{d}_L)/2$ as

$$(2.10) \quad (\bar{d}_L - \hat{d}_L)/2 = \{\mathbb{L}(H_1^1) - \mathbb{L}(H_1^0)\} - \{\mathbb{L}(H_0^1) - \mathbb{L}(H_0^0)\}.$$

Whereas the two bracketed terms in (2.10) are non-negative because they correspond to two LRT statistics, the difference between these two terms is not guaranteed to be non-negative. A simple example illustrates this well. For the regression model $[Y \mid X_1, X_2] \sim \mathcal{N}(\beta_0 + \beta_1 X_1 + \beta_2 X_2, \sigma^2)$, the LRT statistic for testing $H_1^0 : \beta_1 = 0, \beta_2 \in \mathbb{R}$ against $H_1^1 : \beta_1, \beta_2 \in \mathbb{R}$ is not necessarily larger (or smaller) than that for testing $H_0^0 : \beta_1 = \beta_2 = 0$ against $H_0^1 : \beta_1 \in \mathbb{R}, \beta_2 = 0$. A schematic illustration is provided in Figure 2.

The decomposition (2.10) provides another interpretation of $\hat{r}_L$. The test statistic $\mathbb{L}(H_1^1) - \mathbb{L}(H_1^0)$ seeks evidence for detecting the falsity of $\mathcal{r} = 0$ in

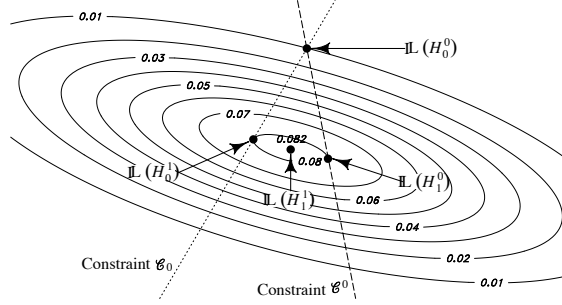


FIG 2. A schematic illustration of the sign of (2.10). The contour lines of $\bar{L}(\psi^1, \dots, \psi^m)$ are plotted. The two straight lines refer to constraints $\mathcal{E}_0$ and $\mathcal{E}^0$. Since $\mathbb{L}(H_1^1) = 0.082$, $\mathbb{L}(H_0^1) = \mathbb{L}(H_1^0) = 0.08$, and $\mathbb{L}(H_0^0) = 0.01$, we have $\{\mathbb{L}(H_1^1) - \mathbb{L}(H_1^0)\} - \{\mathbb{L}(H_0^1) - \mathbb{L}(H_0^0)\} = 0.002 - 0.007 < 0$. Note that the function $\bar{L}(\psi^1, \dots, \psi^m)$ in (2.9) is at least 4-dimensional (i.e., $\theta^1, \theta^2, \eta^1, \eta^2$) generally, so the above illustration in a 2-dimension space is just conceptual.

both θ and η , whereas $\mathbb{L}(H_0^1) - \mathbb{L}(H_0^0)$ seeks evidence only in η . For cases where θ and η are orthogonal (at least locally), the left-hand side of (2.10) can be viewed as a measure of evidence against $\varkappa = 0$ solely from θ ; Proposition 2.2 already hinted this possibility. However, the “test statistic” (2.10) has a very serious problem apart from being possibly negative. Because $\mathcal{E}_0$ requires all θ^ℓ s to coincide with a specific θ_0 , $\mathcal{E}_0$ is not nested within $\mathcal{E}^0$, i.e., $\mathcal{E}^0 \not\supset \mathcal{E}_0$. Hence $\hat{r}_L$ is guaranteed to consistently estimate $\varkappa_m$ only under H_0 . This explains Corollary 2.3, and leads to an improvement below.

2.4. *A Consistent and Non-negative Estimator of $\varkappa_m$.* Our new estimator simply drops the second term in (2.10), that is, we estimate $\varkappa_m$ by

$$(2.11) \quad \hat{r}_L^\diamond = \frac{m+1}{h(m-1)}(\bar{\delta}_L - \hat{\delta}_L), \quad \text{where}$$

$$(2.12) \quad \bar{\delta}_L = 2\bar{L}(\hat{\psi}^1, \dots, \hat{\psi}^m), \quad \hat{\delta}_L = 2\bar{L}(\hat{\psi}^*, \dots, \hat{\psi}^*),$$

where h is the dimension of ψ , and the rhombus “ $\diamond$ ” symbolizes a robust estimator. It is robust, because it is consistent under either H_0 or H_1 , as long as we are willing to impose the EFMI assumption on the entire parameter ψ , a stronger requirement than Assumption 1. This expansion from θ to ψ is inevitable because the LRT must handle the entire ψ , not just θ . The collection of Assumptions 2–4 will be referred to as RC_ψ .

ASSUMPTION 4 (EFMI of ψ). *There is $\varkappa \geq 0$ such that $\mathcal{T}_\psi = (1 + \varkappa)\mathcal{U}_\psi$.*

THEOREM 2.4. Assume RC_ψ . Then for any value of ψ , we have (i) $\hat{r}_L^\diamond \geq 0$ for all m, n ; (ii) $\hat{r}_L^\diamond$ is invariant to parametrization of ψ for all m, n ; and (iii) $\hat{r}_L^\diamond \xrightarrow{\text{pr}} \boldsymbol{\nu}$ as $m, n \rightarrow \infty$, where $\boldsymbol{\nu}$ is given in Assumption 4.

With the improved combining rule $\hat{D}_L(\boldsymbol{\nu}_m)$ of (2.7) and improved estimators for $\boldsymbol{\nu}_m$, we are ready to propose two MI LRT statistics:

$$(2.13) \quad \hat{D}_L^+ = \hat{D}_L(\hat{r}_L^+) \quad \text{and} \quad \hat{D}_L^\diamond = \hat{D}_L(\hat{r}_L^\diamond).$$

For comparison, we also study the test statistic $\hat{D}_L = \hat{D}_L(\hat{r}_L)$.

2.5. *Reference Null Distributions.* The estimators $\hat{r}_L^+$ and $\tilde{r}_L$ have the same functional form asymptotically ($n \rightarrow \infty$) and rely on the same set of assumptions, hence they have the same asymptotic distribution.

LEMMA 2.5. Suppose RC_θ and $m > 1$. Under H_0 , we have, jointly,

$$(2.14) \quad \frac{\hat{r}_L^+}{\boldsymbol{\nu}_m} \Rightarrow M_2 \quad \text{and} \quad \hat{D}_L^+ \Rightarrow \frac{(1 + \boldsymbol{\nu}_m) M_1}{1 + \boldsymbol{\nu}_m M_2}$$

as $n \rightarrow \infty$, where $M_1 \sim \chi_k^2/k$ and $M_2 \sim \chi_{k(m-1)}^2/\{k(m-1)\}$ are independent.

Consequently, the null distribution of $\hat{D}_L^+ = \hat{D}_L(\hat{r}_L^+)$ can be approximated by $F_{k, \tilde{\text{df}}(\hat{r}_L^+, k)}$, but a better approximation will be provided shortly.

For the other proposal, although $\hat{r}_L^+ - \hat{r}_L^\diamond \xrightarrow{\text{pr}} 0$ as $n \rightarrow \infty$ under H_0 , their non-degenerated distributions (after proper scaling) are different because $\hat{r}_L^\diamond$ relies on an average FMI in ψ , but $\hat{r}_L^+$ only on an average FMI in θ .

THEOREM 2.6. Suppose RC_ψ and $m > 1$. Then for any value of ψ ,

$$(2.15) \quad \frac{\hat{r}_L^\diamond}{\boldsymbol{\nu}_m} \Rightarrow M_3 \sim \frac{\chi_{h(m-1)}^2}{h(m-1)}$$

as $n \rightarrow \infty$, where M_3 is independent of the M_1 defined in (2.14).

Theorem 2.6 implies that, if n can be regarded as infinity and $\hat{r}_L^\diamond$ is uniformly integrable in $\mathcal{L}^2$, then

$$\text{Bias}(\hat{r}_L^\diamond) = \mathbb{E}(\hat{r}_L^\diamond) - \boldsymbol{\nu}_m = 0 \quad \text{and} \quad \text{Var}(\hat{r}_L^\diamond) = \frac{2\boldsymbol{\nu}_m^2}{h(m-1)} = O(m^{-1})$$

as $m \rightarrow \infty$. Therefore $\hat{r}_L^\diamond$ is a $\sqrt{m}$ -consistent estimator of ν in $\mathcal{L}^2$. Moreover, for each $m > 1$ and as $n \rightarrow \infty$, we have

$$\frac{\text{Bias}(\hat{r}_L^+)}{\text{Bias}(\hat{r}_L^\diamond)} \rightarrow 1 \quad \text{and} \quad \frac{\text{Var}(\hat{r}_L^+)}{\text{Var}(\hat{r}_L^\diamond)} \rightarrow \frac{h}{k} \geq 1,$$

which implies that $\hat{r}_L^\diamond$ is no less efficient than $\hat{r}_L^+$ when RC_ψ holds. This is of no surprise because of the extra information in $\hat{r}_L^\diamond$ brought in by the stronger Assumption 4. Result (2.15) also gives us the exact (i.e., for any $m > 1$, but assuming $n \rightarrow \infty$) reference null distribution of $\hat{D}_L^\diamond$, as given below.

THEOREM 2.7. *Assume RC_ψ and $m > 1$. Under H_0 , we have*

$$(2.16) \quad \hat{D}_L^\diamond \Rightarrow \frac{(1 + \nu_m) M_1}{1 + \nu_m M_3} \equiv D$$

as $n \rightarrow \infty$, where $M_1 \sim \chi_k^2/k$ and $M_3 \sim \chi_{h(m-1)}^2/\{h(m-1)\}$ are independent.

The impact of the nuisance parameter ν_m on the null distribution diminishes with m , because $\hat{D}_L^\diamond$ and $\hat{D}_L^+$ converge in distribution to $M_1 = \chi_k^2/k$ as $m, n \rightarrow \infty$. Since $M_3 \xrightarrow{\text{pr}} 1$ faster than $M_2 \xrightarrow{\text{pr}} 1$, $\hat{D}_L^\diamond$ is expected to be more robust to ν_m . Nevertheless, because m typically is small in practice (e.g., $m \leq 10$), we cannot ignore the impact of ν_m . This issue has been largely dealt with in the literature by seeking an $F_{k,\text{df}}$ distribution as an approximate null distribution, as in Li *et al.* (1991b). However, directly adopting their $\tilde{\text{df}}$ of (1.6) leads to poorer approximation for our purposes; see below. A better approximation is to match the first two moments of the denominator of (2.16), $1 + \nu_m M_3$, with that of a scaled χ^2 : $a\chi_b^2/b$. This yields $a = 1 + \nu_m$ and $b = (1 + \nu_m^{-1})^2 h(m-1)$, and the approximated $F_{k,\hat{\text{df}}(\nu_m,h)}$, where

$$(2.17) \quad \hat{\text{df}}(\nu_m, h) = \left\{ \frac{1 + \nu_m}{\nu_m} \right\}^2 h(m-1) = \frac{h(m-1)}{\ell_m^2}.$$

This degrees of freedom is appealing because it simply inflates the denominator degrees of freedom M_3 by dividing it by ℓ_m^2 , where $\ell_m = \nu_m/(1 + \nu_m)$ the finite imputation corrected FMI. Intuitively, the less missing information, the closer $F_{k,\hat{\text{df}}(\nu_m,h)}$ should be to χ_k^2/k , the usual large- n asymptotic χ^2 test; as mentioned earlier, for small n , see Barnard and Rubin (1999).

To compare $F_{k,\hat{\text{df}}(\nu_m,h)}$ with $F_{k,\tilde{\text{df}}(\nu_m,h)}$ as approximations to the distribution of D given in (2.16), we compute via simulations

$$\tilde{\alpha} = \text{P} \left\{ D > F_{k,\tilde{\text{df}}(\nu_m,h)}^{-1}(1 - \alpha) \right\} \quad \text{and} \quad \hat{\alpha} = \text{P} \left\{ D > F_{k,\hat{\text{df}}(\nu_m,h)}^{-1}(1 - \alpha) \right\},$$

where $F_{k,\text{df}}^{-1}(q)$ denotes the q -quantile of $F_{k,\text{df}}$. We draw $N = 2^{18}$ independent copies D for each of the following possible combinations: $m \in \{3, 5, 7\}$, $k \in \{1, 2, 4, 8\}$, $\tau = h/k \in \{1, 2, 3\}$, $\ell_m \in \{0, 0.1, \dots, 0.9\}$, and following Benjamin *et al.* (2018)'s recommendation, we use both $\alpha \in \{0.5\%, 5\%\}$. The results for $\alpha = 0.5\%$ and for $\alpha = 5\%$ are shown respectively in Figure 3 and in the Appendix. In general, $\hat{\alpha}$ approximates α much better than $\tilde{\alpha}$, especially when m, k, h are small. When m, h are larger, their performances are similar because both $F_{k,\text{df}(\hat{r}_m, h)}$ and $F_{k,\text{df}(\hat{r}_m, h)}$ get close to χ_k^2/k . But the performances of $\tilde{\alpha}$ and $\hat{\alpha}$ are not monotonic in ℓ_m . Generally speaking, the performance of $F_{k,\text{df}(\hat{r}_m, h)}$ is particularly good for $0\% \lesssim \ell_m \lesssim 30\%$. Consequently, we recommend using $F_{k,\text{df}(\hat{r}_L^\diamond, h)}$ as an approximate null distribution for $\hat{D}_L^\diamond$, and $F_{k,\text{df}(\hat{r}_L^+, k)}$ for $\hat{D}_L^+$, as employed in the rest of this paper. However, these approximations obviously suffer from the usual “plug-in problem” by ignoring the uncertainties in estimating $\hat{r}_m$. Since the $F_{k,\text{df}}$ is not too sensitive to the value of df once it is reasonably large ($\text{df} \geq 20$), the “plug-in problem” is less an issue here than in many other context, leading to acceptable approximations as empirically demonstrated in Section 4. Nevertheless, further improvements are likely and should be sought.

3. Computational Considerations and Comparisons. The statistic $\bar{d}_L$ of (1.7) is easy to compute because only the standard complete-data procedure $\mathcal{D}_L : X \mapsto d_L(\hat{\psi}_0(X), \hat{\psi}(X) \mid X)$ is needed. However, $\hat{d}_L$ of (2.3) and $\hat{r}_L^\diamond$ of (2.8) in general cannot be computed solely by $\mathcal{D}_L$, e.g., $\hat{d}_L$ requires

$$\bar{\mathcal{D}}_L : \mathbb{X} \mapsto \frac{1}{m} \sum_{\ell=1}^m d_L(\hat{\psi}_0^*(\mathbb{X}), \hat{\psi}^*(\mathbb{X}) \mid X^\ell).$$

Creating a subroutine for this computation requires additional effort and information that may beyond a user's capacity. Here we show how to compute or approximate $\hat{d}_L$ and $\hat{r}_L^\diamond$ solely by $\mathcal{D}_L$ or a trivial modification of $\mathcal{D}_L$.

3.1. Computationally Feasible Combining Rule. We begin with precise notation for our complete data X and its sampling model $f(X|\psi)$. For the vast majority of real-world datasets, X is of the form of an $n \times p$ matrix, with rows indicating subjects and columns variables/attributes. We then write $X = (X_1, \dots, X_n)^\top$, and its sampling model by $f_n(X \mid \psi)$. Correspondingly, the ℓ th completed-dataset by MI is $X^\ell = (X_1^\ell, \dots, X_n^\ell)^\top$. Define the stacked dataset by $X^S = [(X^1)^\top, \dots, (X^m)^\top]^\top$, a matrix having mn rows, which is conceptually different from the collection of datasets $\mathbb{X} = \{X^1, \dots, X^m\}$.

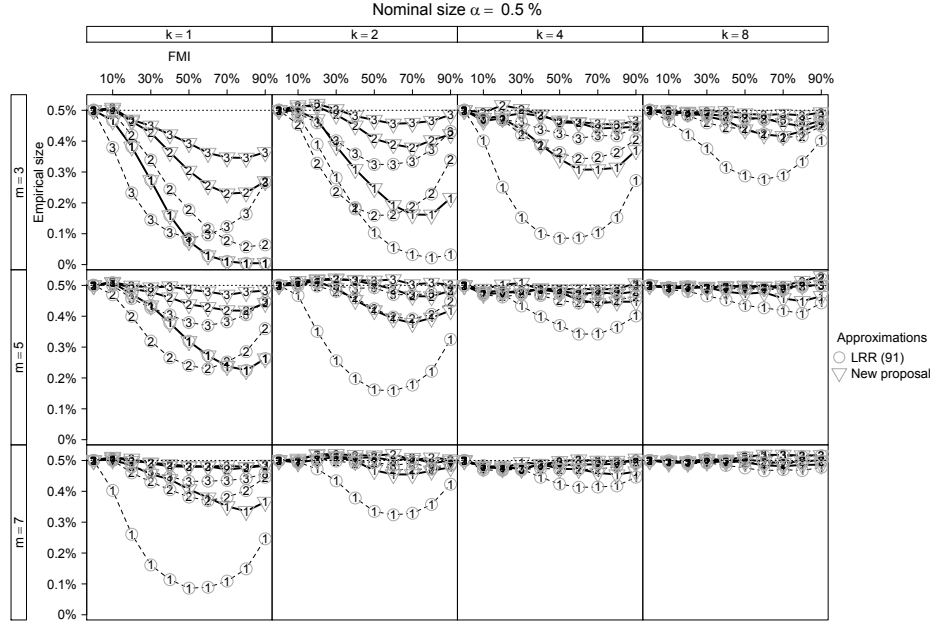


FIG 3. The performances of two different approximated null distributions when the nominal size is $\alpha = 0.5\%$. The vertical axis denotes $\hat{\alpha}$ or $\tilde{\alpha}$, and the horizontal axis denotes the value of f_m . The number attached to each line denotes the value of $\tau = h/k$. The proposed approximation $\hat{\alpha}$ is denoted by thick solid lines with triangles, and the existing approximation $\tilde{\alpha}$ is denoted by thin dashed lines with circles.

Treating X^S as a dataset with size mn , we can define

$$(3.1) \quad \bar{L}^S(\psi) = \frac{1}{m} \log f_{mn}(X^S | \psi),$$

which, other than the scaling factor $1/m$, is just the ordinary log-likelihood function of ψ based on the dataset X^S (for computation purposes). Consequently, as long as the user's complete-data procedure can handle size mn instead of just n , the user can apply it to X^S to obtain

$$(3.2) \quad \hat{\psi}_0^S = \arg \max_{\psi \in \Psi : \theta(\psi) = \theta_0} \bar{L}^S(\psi) \quad \text{and} \quad \hat{\psi}^S = \arg \max_{\psi \in \Psi} \bar{L}^S(\psi).$$

Consequently, the quantities

$$(3.3) \quad \hat{\delta}_{0,S} = 2\bar{L}^S(\hat{\psi}_0^S) \quad \text{and} \quad \hat{\delta}_S = 2\bar{L}^S(\hat{\psi}^S)$$

are readily available from the user's complete-data procedure. It is then desirable if we can replace $\bar{L}(\psi)$ by $\bar{L}^S(\psi)$ in the proposed test statistics.

Precisely, in parallel to (2.7), (2.8) and (2.11), we define

$$(3.4) \quad \hat{D}_S(\boldsymbol{\nu}_m) = \frac{\hat{d}_S}{k(1 + \boldsymbol{\nu}_m)}, \quad \text{with } \hat{d}_S = \hat{\delta}_S - \hat{\delta}_{0,S} \text{ of (3.3);}$$

$$(3.5) \quad \hat{r}_S = \frac{m+1}{k(m-1)}(\bar{d}_S - \hat{d}_S), \quad \text{with } \bar{d}_S = \bar{d}_L \text{ of (1.7);}$$

$$(3.6) \quad \hat{r}_S^\diamond = \frac{m+1}{h(m-1)}(\bar{\delta}_S - \hat{\delta}_S), \quad \text{with } \bar{\delta}_S = \bar{\delta}_L \text{ of (2.12);}$$

and $\hat{r}_S^+ = \max(0, \hat{r}_S)$. The “stacked” counterparts of $\hat{D}_L^\diamond$ and its existing counterparts $\hat{D}_L$ and $\hat{D}_L^+$ (see (2.13)) then are given by

$$(3.7) \quad \hat{D}_S^\diamond = \hat{D}_S(\hat{r}_S^\diamond), \quad \hat{D}_S = \hat{D}_S(\hat{r}_S), \quad \hat{D}_S^+ = \hat{D}_S(\hat{r}_S^+).$$

PROPOSITION 3.1. *If $X = (X_1, \dots, X_n)^\top$ is row-independent for arbitrary n , i.e., $f(X \mid \psi) = \prod_{i=1}^n f(X_i \mid \psi)$, then (2.1) and (3.1) are the same: $\bar{L}(\psi) \equiv \bar{L}^S(\psi)$. Consequently, $\hat{D}_S \equiv \hat{D}_L$ and $\hat{D}_S^\diamond \equiv \hat{D}_L^\diamond$.*

Since for many applications, the rows correspond to individual subjects, the row-independence assumption typically holds for *arbitrary* n . Hence we can extend from n to mn , assuming the user’s complete-data procedure is not size-limited. Even if it does not hold, we can still have $\hat{d}_L \simeq \hat{d}_S$ under some RCs that guarantee $\bar{L}(\psi)$ and $\bar{L}^S(\psi)$ are close; see Appendix A, where we also reveal a subtle but important difference between $\hat{r}_L^\diamond$ and $\hat{r}_S^\diamond$.

Similar to $\mathcal{D}_L$ in (1.11), we define complete-data functions (with data X being the only input)

$$(3.8) \quad \mathcal{D}_{L,0}(X) = 2 \log f(X \mid \hat{\psi}_0(X)), \quad \mathcal{D}_{L,1}(X) = 2 \log f(X \mid \hat{\psi}(X)).$$

The subroutine for evaluating the complete-data LRT function $X \mapsto \mathcal{D}_L(X)$ is usually available, as is the subroutine for $X \mapsto \mathcal{D}_{L,1}(X)$; for example, the function `logLik` in R extracts the maximum of complete data log-likelihood for objects belonging to classes “glm”, “lm”, “nls” and “Arima”.

The Algorithms 1 and 2 listed below compute tests given $\hat{D}_S^\diamond$ and $\hat{D}_S^+$, respectively. Whenever possible, we recommend the use of the robust MI LRT given by Algorithm 1, since it has the best theoretical guarantee. The second test can be useful when $\mathcal{D}_L$ but not $\mathcal{D}_{L,1}$ is available.

3.2. Computational Comparison with Existing Tests. First, we list some existing estimators of $\boldsymbol{\nu}_m$ and their computation. Let $s_{W,1}^2$ and $s_{W,1/2}^2$ be the sample variances of $\{d_W^\ell\}_{\ell=1}^m$ and $\{\sqrt{d_W^\ell}\}_{\ell=1}^m$, respectively. By the method of

Algorithm 1: (Robust) MI LRT statistic $\hat{D}_S^\diamond$

Input: Datasets $X^1, \dots, X^m$; dimensions h, k ; functions $\mathcal{D}_{L,1}, \mathcal{D}_L$ in (3.8), (1.11).

begin

- Compute $\bar{d}_S = m^{-1}\{\mathcal{D}_{L,1}(X^1) + \dots + \mathcal{D}_{L,1}(X^m)\}$.
 - (i) Stack the datasets to form $X^S = [(X^1)^\top, \dots, (X^m)^\top]^\top$.
 - (ii) Compute $\hat{d}_S = m^{-1}\mathcal{D}_L(X^S)$ and $\hat{\delta}_S = m^{-1}\mathcal{D}_{L,1}(X^S)$.
 - Calculate $\hat{r}_S^\diamond$ according to (3.6), and $\hat{D}_S^\diamond$ according to (3.4) and (3.7).
 - Calculate $\widehat{\text{df}}(\hat{r}_S^\diamond, h)$ according to (2.17).
 - Compute the p -value as $1 - F_{k, \widehat{\text{df}}(\hat{r}_S^\diamond, h)}(\hat{D}_S^\diamond)$.
-

Algorithm 2: MI LRT statistic $\hat{D}_S^+$

Input: Datasets $X^1, \dots, X^m$; dimension k ; function $\mathcal{D}_L$ in (1.11).

begin

- Compute $\bar{d}_L = m^{-1}\{\mathcal{D}_L(X^1) + \dots + \mathcal{D}_L(X^m)\}$.
 - (i) Stack the datasets to form $X^S = [(X^1)^\top, \dots, (X^m)^\top]^\top$.
 - (ii) Compute $\hat{d}_S = m^{-1}\mathcal{D}_L(X^S)$.
 - Calculate $\hat{r}_S^+$ according to (3.5), and $\hat{D}_S^+$ according to (3.4) and (3.7).
 - Calculate $\widehat{\text{df}}(\hat{r}_S^+, k)$ according to (2.17).
 - Compute the p -value as $1 - F_{k, \widehat{\text{df}}(\hat{r}_S^+, k)}(\hat{D}_S^+)$.
-

moments concerning $s_{W,1}^2$ and $s_{W,1/2}^2$, Rubin (2004) and Li *et al.* (1991a) respectively proposed estimating $\boldsymbol{\nu}_m$ by

$$(3.9) \quad \tilde{r}_{W,1} = \frac{(1 + 1/m)s_{W,1}^2}{2\bar{d}_W + \sqrt{\{4\bar{d}_W^2 - 2ks_{W,1}^2\}^+}}, \quad \tilde{r}_{W,1/2} = (1 + 1/m)s_{W,1/2}^2,$$

where $\{a\}^+ = \max(0, a)$. Note that when k is large and m is small, using (3.9) may lead to power loss, although the users have no choice when they are given only $\{d_W^\ell\}_{\ell=1}^m$. A trivial modification of $\tilde{r}_L$ of (1.8), i.e., $\tilde{r}_L^+ = \max(0, \tilde{r}_L)$, is a better alternative if the user is able to compute $\tilde{r}_L$.

Second, we list some alternative MI combining rules. Having the above estimators of $\boldsymbol{\nu}_m$, we can insert them into the following combining rules:

$$\tilde{D}'_W(\boldsymbol{\nu}_m) = \frac{\tilde{d}'_W}{k(1 + \boldsymbol{\nu}_m)}, \quad \tilde{D}_L(\boldsymbol{\nu}_m) = \frac{\tilde{d}_L}{k(1 + \boldsymbol{\nu}_m)}, \quad \tilde{D}_L^+(\boldsymbol{\nu}_m) = \left\{ \tilde{D}_L(\boldsymbol{\nu}_m) \right\}^+.$$

Using (1.3) and (1.8), we can also define the following combining rules:

$$(3.10) \quad \bar{D}'_W(\boldsymbol{\nu}_m) = \frac{\bar{d}'_W - \frac{k(m-1)}{m+1}\boldsymbol{\nu}_m}{k(1 + \boldsymbol{\nu}_m)}, \quad \bar{D}_L(\boldsymbol{\nu}_m) = \frac{\bar{d}_L - \frac{k(m-1)}{m+1}\boldsymbol{\nu}_m}{k(1 + \boldsymbol{\nu}_m)}.$$

The combining rule $\overline{D}'_W(\boldsymbol{r}_m)$ is useful when computing $\tilde{d}'_W$ is difficult but computing $\tilde{d}_W$ and estimating $\boldsymbol{r}_m$ are simple. However, the resulting power may deteriorate if the estimator of $\boldsymbol{r}_m$ is inefficient or inaccurate. This type of test statistics is also mentioned in Li *et al.* (1991a). Indeed, there are infinitely many asymptotically equivalent test statistics, e.g., any convex combination of $\tilde{D}_W(\boldsymbol{r}_m)$ and $\overline{D}_W(\boldsymbol{r}_m)$, i.e., $\phi\tilde{D}_W(\boldsymbol{r}_m) + (1-\phi)\overline{D}_W(\boldsymbol{r}_m)$, for $\phi \in [0, 1]$. When $\tilde{r}_{W,1}$ or $\tilde{r}_{W,1/2}$ is used for estimating $\boldsymbol{r}_m$, the null distributions of the resulting MI test statistics can be approximated by $F_{k, \tilde{\text{df}}'(\boldsymbol{r}_m, k)}$, where $\tilde{\text{df}}'(\boldsymbol{r}_m, k) = (m-1)(1 + \boldsymbol{r}_m^{-1})^2 k^{-3/m}$; see Li *et al.* (1991a).

Next, we introduce and recall some notation to facilitate the comparison: (a) standard complete-data moments estimation ($\mathcal{M}_W$ and $\mathcal{M}_L$) and testing ($\mathcal{D}_W$ and $\mathcal{D}_L$) procedures, and (b) non-standard complete-data procedures ($\tilde{\mathcal{D}}_L$, $\overline{\mathcal{D}}_L$, $\mathcal{D}_{L,1}$ and $\overline{\mathcal{D}}_{L,1}$), where

$$\begin{aligned} \mathcal{M}_W(X) &= \{\hat{\theta}(X), U(X)\}, \quad \mathcal{M}_L(X) = \{\hat{\psi}(X), \hat{\psi}_0(X)\}, \\ \mathcal{D}_W(X) &= d_W(\hat{\theta}(X), U(X)), \quad \overline{\mathcal{D}}_{L,1}(\mathbb{X}) = \frac{2}{m} \sum_{\ell=1}^m \log f(X^\ell \mid \hat{\psi}^*(\mathbb{X})). \end{aligned}$$

Clearly, $\mathcal{M}_W$ produces $\mathcal{D}_W$, and $\{\mathcal{M}_L, \tilde{\mathcal{D}}_L\}$ produces $\mathcal{D}_L$. If users can perform optimization, $\tilde{\mathcal{D}}_L$ produces $\{\mathcal{M}_L, \mathcal{D}_L, \mathcal{D}_{L,1}\}$. Note that an un-normalized density can be used in $\mathcal{D}_{L,1}$ and $\overline{\mathcal{D}}_{L,1}$.

Table 2 summarizes whether a particular pair of $D(\cdot)$ and r , resulting the statistic $D(r)$, has the following statistical or computational properties.

- (Inv) $D(r)$ and r are invariant to re-parametrization of ψ ;
- (Rob) r is robust against θ_0 , i.e., consistent under both H_0 and H_1 ;
- (≥ 0) $D(r)$ and r are non-negative for all m and n ;
- (Pow) the test has high power to reject H_0 under H_1 ;
- (Def) $D(r)$ and r are always well-defined and numerically well-conditioned;
- (Sca) the MI procedure requires users only to deal with scalars;
- (Dep) $X_1, \dots, X_n$ can be dependent; and
- (EFMI) whether EFMI is assumed for θ or for ψ .

In summary, $\hat{D}_S(\hat{r}_S^+)$ is the most computationally attractive test statistic. If the user is willing to make stronger assumptions, $\hat{D}_S(\hat{r}_S^\diamond)$ has better statistical properties, and is still computationally feasible. Nevertheless, $\hat{D}_L(\hat{r}_L^\diamond)$ is the most general test statistic and has the best statistical properties.

TABLE 2

Computational requirements and statistical properties of MI test statistics, their associated combining rules and estimators of FMI r_m .

The symbols “+” and “−” mean that the test statistic (or estimator) is equipped and not equipped with the indicated property, respectively; see the end of § 3.2 for heading descriptions. The reference papers/book are abbreviated as follows: [Rubin \(2004\)](#) (R04), [Li et al. \(1991a\)](#) (LMRR91) and [Meng and Rubin \(1992\)](#) (MR92).

Test	No.	Combining Rule		Estimator of r_m		Approx. null distribution ^a		Reference	Properties							
		Formula	Routine	Formula	Routine	Original	Proposed		Inv	Rob	≥ 0	Pow	Def	Sca	Dep	EFMI
WT	W-1	$\tilde{D}'_W(r_m)$	$\mathcal{M}_W$	$\tilde{r}'_W$	$\mathcal{M}_W$	$F_{k,\tilde{\text{df}}(r_m,k)}$ ^b	$F_{k,\widehat{\text{df}}(r_m,k)}$	R04	−	+ ^c	+	−	−	−	+	θ
	W-2	$\tilde{D}'_W(r_m)$ ^d	$\mathcal{M}_W$	$\tilde{r}'_{W,1}$	$\mathcal{D}_W$	$F_{k,\tilde{\text{df}}'(r_m,k)}$	NA	R04	−	−	+	−	−	−	+	θ
	W-3	$\tilde{D}'_W(r_m)$	$\mathcal{M}_W$	$\tilde{r}'_{W,1/2}$	$\mathcal{D}_W$	$F_{k,\tilde{\text{df}}'(r_m,k)}$	NA	LMRR91	−	−	+	−	−	−	+	θ
	W-4	$D_W(T)$ ^e	$\mathcal{M}_W$	$\tilde{r}'_W$	$\mathcal{D}_W$	$F_{k,\tilde{\text{df}}(r_m,k)}$	$F_{k,\widehat{\text{df}}(r_m,k)}$	R04	−	+	+	−	−	−	+	θ ^f
	W-5	$\overline{D}'_W(r_m)$	$\mathcal{D}_W$	$\tilde{r}'_{W,1}$	$\mathcal{D}_W$	$F_{k,\tilde{\text{df}}'(r_m,k)}$	NA	R04	−	−	−	−	−	+	+	θ
	W-6	$\overline{D}'_W(r_m)$	$\mathcal{D}_W$	$\tilde{r}'_{W,1/2}$	$\mathcal{D}_W$	$F_{k,\tilde{\text{df}}'(r_m,k)}$	NA	LMRR91	−	−	−	−	−	+	+	θ
LRT	L-1	$\tilde{D}_L(r_m)$	$\mathcal{M}_L, \tilde{\mathcal{D}}_L$	$\tilde{r}_L$	$\mathcal{M}_L, \tilde{\mathcal{D}}_L$	$F_{k,\tilde{\text{df}}(r_m,k)}$	$F_{k,\widehat{\text{df}}(r_m,k)}$	MR92	−	−	−	−	+	− ^g	+	θ
	L-2	$\tilde{D}_L^+(r_m)$	$\mathcal{M}_L, \tilde{\mathcal{D}}_L$	$\tilde{r}_L^+$	$\mathcal{M}_L, \tilde{\mathcal{D}}_L$	$F_{k,\tilde{\text{df}}(r_m,k)}$	$F_{k,\widehat{\text{df}}(r_m,k)}$	MR92 ^h	−	−	+	−	+	−	+	θ
	L-3	$\hat{D}_S(r_m)$	$\mathcal{D}_L$	$\hat{r}_S$	$\mathcal{D}_L$	$F_{k,\tilde{\text{df}}(r_m,k)}$	$F_{k,\widehat{\text{df}}(r_m,k)}$	Proposal	+	−	−	−	+	+	−	θ
	L-4	$\hat{D}_S(r_m)$	$\mathcal{D}_L$	$\hat{r}_S^+$	$\mathcal{D}_L$	$F_{k,\tilde{\text{df}}(r_m,k)}$	$F_{k,\widehat{\text{df}}(r_m,k)}$	Proposal	+	−	+	−	+	+	−	θ
	L-5	$\hat{D}_S(r_m)$	$\mathcal{D}_L$	$\hat{r}_S^\diamond$	$\mathcal{D}_{L,1}$	$F_{k,\tilde{\text{df}}(r_m,h)}$	$F_{k,\widehat{\text{df}}(r_m,h)}$	Proposal	+	+	+	+	+	+	−	ψ
	L-6 ⁱ	$\hat{D}_L(r_m)$	$\overline{\mathcal{D}}_L$	$\hat{r}_L^+$	$\overline{\mathcal{D}}_L$	$F_{k,\tilde{\text{df}}(r_m,k)}$	$F_{k,\widehat{\text{df}}(r_m,k)}$	Proposal	+	−	+	−	+	+	+	θ
	L-7 ^j	$\hat{D}_L(r_m)$	$\overline{\mathcal{D}}_L$	$\hat{r}_L^\diamond$	$\overline{\mathcal{D}}_{L,1}$	$F_{k,\tilde{\text{df}}(r_m,h)}$	$F_{k,\widehat{\text{df}}(r_m,h)}$	Proposal	+	+	+	+	+	+	+	ψ

^aIn actual computation, the r_m in the denominator degree of freedom of F is replaced by its corresponding estimator given in the previous column.

^bThe original approximate null distribution documented in [Rubin \(2004\)](#) was modified by [Li et al. \(1991a\)](#). This footnote also applies to W-2,4,5.

^cThe estimator $\tilde{r}'_W$ does not depend on θ_0 , but its MSE may be inflated under H_1 if a bad parametrization of θ is used.

^dThe originally proposed combining rule is $\overline{D}'_W(r_m)$; see (3.10). Although $\overline{D}'_W(r_m)$ is more computational feasible, the power loss is more significant than $\tilde{D}'_W(r_m)$ after inserting an inefficient estimator $\tilde{r}'_{W,1}$ for r_m . This footnote also applies to W-3.

^eComputing the test statistic $D_W(T) = d_W(\bar{\theta}, T)/k$ does not require estimating r_m .

^fEFMI is not required for the test statistic $D_W(T)$, but it is required for its approximate null distribution.

^gAveraging and processing vector estimators of ψ , but not their covariance matrixes, is needed. This footnote also applies to L-2.

^hIt is a trivial modification of the original proposal in MR92 by replacing $\tilde{r}_L$ with $\tilde{r}_L^+ = \max\{0, \tilde{r}_L\}$.

ⁱL-6 is equivalent to L-4 when the rows of X are independent.

^jL-7 is equivalent to L-5 when the rows of X are independent.

TABLE 3
The values of parameters used in the simulation experiment in § 4.1.

Experiment		Parameters							
		Fixed			Variable				
No.	Varying	ρ	p	ℓ	Case 1	Case 2	Case 3	Case 4	Case 5
I	Correlation ρ	–	2	0.5	–0.8	–0.4	0	0.4	0.8
II	Dimension p	0.4	–	0.5	2	3	4	5	6
III	FMI ℓ	0.4	2	–	0.1	0.3	0.5	0.7	0.9

4. Empirical Investigation and Findings.

4.1. *Simulation Studies.* Suppose that $X_1, \dots, X_n \sim \mathcal{N}_p(\mu, \Sigma)$ independently, where $\Sigma = \sigma^2\{(1 - \rho)I_p + \rho\mathbf{1}_p\mathbf{1}_p^\top\}$, and $\mathbf{1}_p$ is the p -vector of ones. The values of p , σ^2 , ρ and μ are specified below. Further assume that only $n_{\text{obs}} = \lfloor (1 - \ell)n \rfloor$ data points are observed, where $\ell \in (0, 1)$ is the FMI. Let $X_{\text{obs}} = \{X_i : i = 1, \dots, n_{\text{obs}}\}$ and $X_{\text{mis}} = \{X_i : i = n_{\text{obs}} + 1, \dots, n\}$. Suppose that we want to test whether the means of all components are equal, i.e., $H_0 : \mu = \mu_0\mathbf{1}_p$, where $\mu_0 \in \mathbb{R}$ is an unknown constant.

Obviously, one may directly use the observed dataset to construct the LRT statistic D_L without MI. The test D_L (denoted by L-0) is regarded as a benchmark for comparison. Throughout this subsection, W-1,2,3,4 and L-1,2,3,4,5 listed in Table 2 are compared. In the imputation step, a Bayesian model is employed for imputation. Assume a multivariate Jeffreys prior on (μ, Σ) , i.e., $f(\mu, \Sigma) \propto |\Sigma|^{-(p+1)/2}$. Let $\bar{X}_{\text{obs}}$ and S_{obs} be the sample mean and sample covariance matrix based on X_{obs} . Then, the ℓ th imputed missing dataset can be produced by the following procedure, for $\ell = 1, \dots, m$.

1. Draw a posterior sample Σ^ℓ from the inverse-Wishart distribution with $(n_{\text{obs}} - 1)$ degrees of freedom and scale matrix S_{obs}^{-1} .
2. Draw one posterior sample μ^ℓ from $\mathcal{N}_p(\bar{X}_{\text{obs}}, \Sigma^\ell/n_{\text{obs}})$.
3. Draw $(n - n_{\text{obs}})$ imputed missing values $\{X_i^\ell : i = n_{\text{obs}} + 1, \dots, n\}$ from $\mathcal{N}_p(\mu^\ell, \Sigma^\ell)$ independently. Also, denote $X_i^\ell = X_i$ for $i = 1, \dots, n_{\text{obs}}$.

With the ℓ th completed dataset, the unconstrained MLEs for μ and Σ are

$$\hat{\mu}^\ell = \frac{1}{n} \sum_{i=1}^n X_i^\ell, \quad \hat{\Sigma}^\ell = \frac{1}{n} \sum_{i=1}^n \left(X_i^\ell - \hat{\mu}^\ell \right) \left(X_i^\ell - \hat{\mu}^\ell \right)^\top.$$

Whereas we generate data using a covariance matrix with common variance and correlation, our model does not assume any structure for Σ . The only restriction we can impose is the common-mean assumption under the null,

for which the constrained MLEs are

$$\hat{\mu}_0^\ell = \left\{ \frac{\mathbf{1}_p^\top (\hat{\Sigma}^\ell)^{-1} \hat{\mu}^\ell}{\mathbf{1}_p^\top (\hat{\Sigma}^\ell)^{-1} \mathbf{1}_p} \right\} \mathbf{1}_p, \quad \hat{\Sigma}_0^\ell = \hat{\Sigma}^\ell + \left(\hat{\mu}^\ell - \hat{\mu}_0^\ell \right) \left(\hat{\mu}^\ell - \hat{\mu}_0^\ell \right)^\top.$$

In the experiment, we study the impact of parametrization on different test statistics. For the Wald tests, three parametrizations of θ are examined:

- (i) $\theta = (\mu_2 - \mu_1, \dots, \mu_p - \mu_{p-1})^\top$ — differences of means,
- (ii) $\theta = (\mu_2/\mu_1, \dots, \mu_p/\mu_{p-1})^\top - \mathbf{1}_{p-1}$ — relative differences of means, and
- (iii) $\theta = (\mu_2^3 - \mu_1^3, \dots, \mu_p^3 - \mu_{p-1}^3)^\top$ — differences of cubic means.

For any case above, H_0 can be re-expressed as $\theta_0 = \mathbf{0}_{p-1}$, an $(p-1)$ -vector of zeros. For LRTs, the following parametrizations of ψ are used:

- (i) $\psi = \{\mu; \Sigma\}$ — means and covariances,
- (ii) $\psi = \{\sqrt{\sigma_{ii}}/\mu_i, 1 \leq i \leq p; \Sigma\}$ — noise-to-signal and covariances, and
- (iii) $\psi = \{\mu^\top \Sigma^{-1/2}; \Sigma^{-1}\}$ — standardized means and precisions,

where $\Sigma = (\sigma_{ij})$ and $\Sigma^{1/2}$ is the symmetric square root of Σ . The dimension of ψ is $h = (p^2 + 3p)/2$.

In the first part of the experiment, we study the distribution of p -values derived from each test under H_0 . In particular, we use $n = 100$, $m = 3$, $\sigma^2 = 5$ and $\mu = \mathbf{1}_p$, with various values of ρ , p and ℓ specified in Table 3. All simulations are repeated 2^{12} times. The comparison under parametrization (ii) is shown in Figure 4; whereas those under parametrizations (i) and (iii) are deferred to Appendix C. Note that, for Wald tests under parametrization (ii), the matrix U^ℓ is singular in less than 0.25% of the replications, and those cases are removed from the analysis (which should favor the Wald tests).

The empirical sizes (i.e., type-I errors) of the MI Wald tests generally deviate from the nominal size α under parametrization (ii). In contrast, the sizes of all LRTs are closer to α . However, the original L-1 and its trivial modification L-2 do not have accurate sizes when $|\rho|$ or ℓ is large. They can be over-sized or under-sized depending on which parametrization is used. Moreover, the trivial modification L-2 does not help to correct the size, and it may even worsen the test. For our test statistics L-3 and L-4, they are invariant to parametrizations and have quite accurate sizes, although they are under-sized in challenging cases where both p and ℓ are large. Moreover, they are identical throughout our simulation experiments, i.e., we never observed $\hat{r}_L < 0$. For our recommended statistic L-5, it gives the most satisfactory overall results. It generally has very accurate size, except that it is slightly over-sized for large p , a problem that should diminish when we use m beyond the smallest recommended $m = 3$.

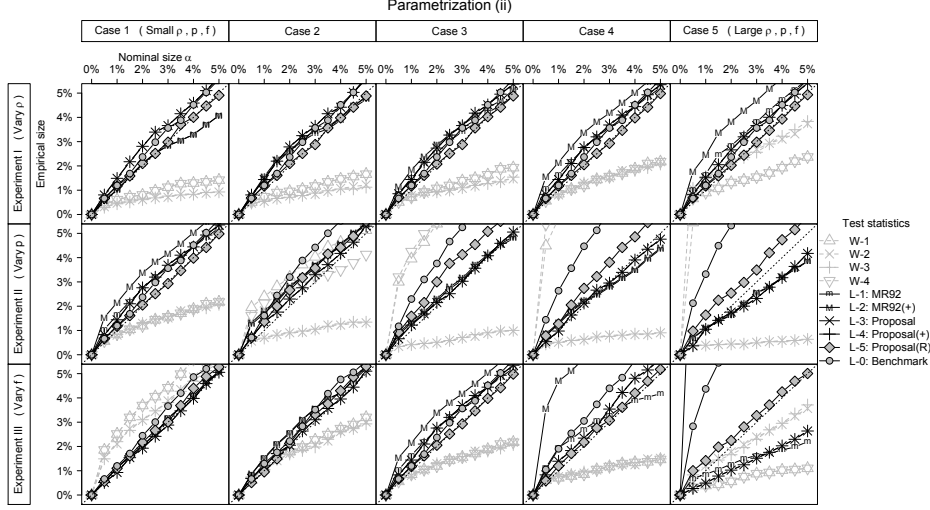


FIG 4. The comparison between empirical size and nominal size α under parametrization (ii) for $\alpha \in (0, 5\%]$. The Wald tests (W-1,2,3,4) and LRTs (L-0,1,2,3,4,5) are represented by grey dashed and black solid lines, respectively. The LRT statistic $\tilde{D}_L$ (L-1: MR92) and its modification $\tilde{D}_L^+$ (L-2: MR92(+)) are the tests that greatly improved upon by our proposals $\hat{D}_S$ (L-3: Proposal), $\hat{D}_S^+$ (L-4: Proposal(+)) and $\hat{D}_S^\diamond$ (L-5: Proposal(R)).

Interestingly, as seen clearly in Figure 4, the benchmark L-0 performs very badly for large p and ℓ . The sample size per parameter, n/h , is small; for $p \geq 4$, $n/h \leq 100/14 < 8$. The asymptotic null distribution χ_k^2/k then can fail badly under arbitrary or even all parametrizations; (ii) apparently falls into this category. An F approximation would be more appropriate. But this is exactly what is being used for MI tests, albeit with different choices of the denominator degrees of freedom. Table 4 documents how often $\tilde{r}_L$, $\tilde{D}_L$ and $\hat{r}_S$ are negative. In some cases, nearly half of the simulated values of $\tilde{r}_L$ and $\tilde{D}_L$ are negative. In contrast, $\hat{r}_S$ is always non-negative in our simulation, despite the fact that it can be negative in theory.

To study the power of each test, we set $\ell = 0.5$, $p = 2$, $\rho = 0.8$, $\sigma^2 = 5$ and $\mu = (-2 + \delta, -2 + 2\delta)^\top$ for different values of $m \in \{3, 10, 30\}$, $n \in \{100, 400, 1600\}$ and $\delta = \mu_2 - \mu_1 \in [0, 4]$. The empirical power functions for size 0.5% tests under parametrizations (i), (ii) and (iii) are plotted in Figure 5. The results for size 5% tests are deferred to Table 12 of the Appendix. Generally, none of the Wald tests exhibits monotonically increasing power as δ increases, and their performance is affected significantly by parametrization. In particular, the powers can be as low as zero when $1 \lesssim \delta \lesssim 2$ under

TABLE 4
The empirical proportions of negative $\tilde{r}_L$ and $\tilde{D}_L$. The results under parametrizations (ii) and (iii) are shown. For parametrization (i), $\tilde{r}_L \geq 0$ and $\tilde{D}_L \geq 0$ in the experiments.

Experiment	Parametrization	Case									
		1	2	3	4	5	1	2	3	4	5
		% of $\tilde{r}_L < 0$					% of $\tilde{D}_L < 0$				
I	(ii)	1	2	3	4	5	26	16	13	12	12
	(iii)	6	6	7	7	7	1	1	1	1	2
II	(ii)	4	1	0	0	0	12	5	3	4	3
	(iii)	7	3	1	1	1	1	0	0	0	0
III	(ii)	13	6	4	4	3	55	25	12	5	2
	(iii)	18	9	7	5	4	20	5	1	1	0

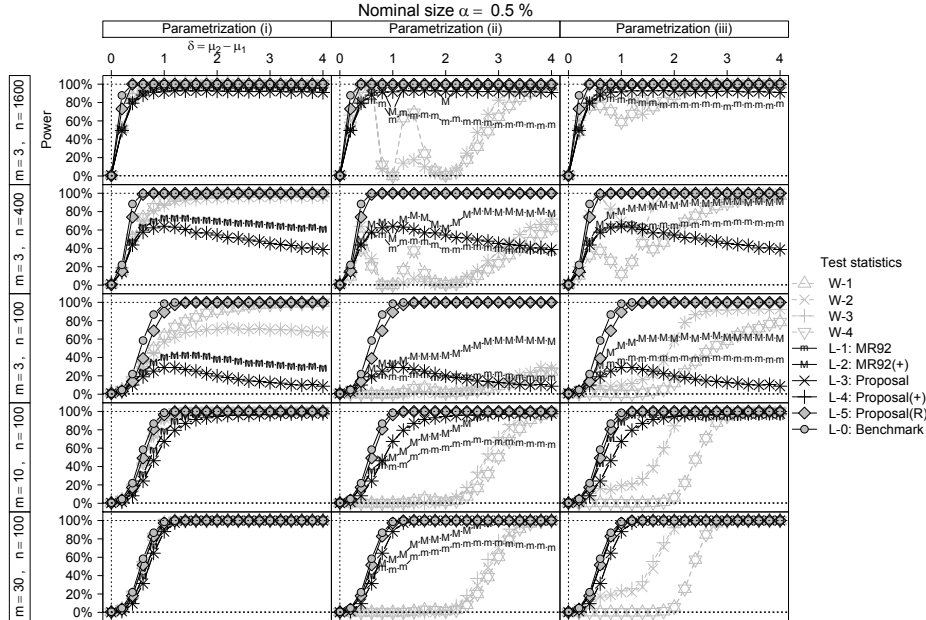


FIG 5. The power curves under different parametrizations. The nominal size is $\alpha = 0.5\%$. In each plot, the vertical axis denotes the power, whereas the horizontal axis denotes the value of $\delta = \mu_2 - \mu_1$. The legend in Figure 4 also applies here.

parametrizations (ii) and (iii). Under parametrization (ii), L-1 is not powerful even for large δ . Moreover, its trivial modifications L-2 cannot retrieve all the power it should have. On the other hand, our first proposed test

TABLE 5

The range of empirical size $[\min \hat{\alpha}, \max \hat{\alpha}]$ in percentage, where max and min are taken over the three parametrizations. Only one value is recorded for those tests that are invariant to parametrization. The nominal size is $\alpha = 0.5\%$.

(n, m)	Range of empirical size: $[\min \hat{\alpha}, \max \hat{\alpha}]/\%$				
	(1600, 3)	(400, 3)	(100, 3)	(100, 10)	(100, 30)
W-1	[0.90, 1.05]	[0.76, 1.05]	[0.20, 1.22]	[0.07, 0.56]	[0.02, 0.49]
W-2	[0.90, 1.05]	[0.98, 1.22]	[0.93, 1.25]	[0.32, 0.73]	[0.20, 0.85]
W-3	[0.98, 1.05]	[0.98, 1.25]	[0.90, 1.29]	[0.34, 0.71]	[0.22, 0.73]
W-4	[0.90, 1.05]	[0.76, 1.05]	[0.20, 1.22]	[0.07, 0.56]	[0.02, 0.49]
L-1	[0.90, 1.03]	[1.10, 1.64]	[1.15, 1.49]	[0.37, 1.05]	[0.10, 0.46]
L-2	[0.90, 1.05]	[1.10, 1.76]	[1.15, 2.37]	[0.37, 0.98]	[0.10, 0.49]
L-3	0.90	1.10	0.83	0.24	0.07
L-4	0.90	1.10	0.83	0.24	0.07
L-5	0.46	0.44	0.68	0.46	0.42
L-0	0.39	0.66	0.66	0.66	0.66

statistics L-3 and L-4 perform better than L-1 and L-2 at least for large m , however, they also lose a significant amount of power when m is small.

Compared with all these, our recommended test statistic L-5 performs extremely well for all m and n , with power very close to the benchmark L-0 even for small m . To ensure the comparisons of power are fair, we also investigate the empirical (actual) size, $\hat{\alpha}$, in comparison to the nominal type-I error α . Table 5 shows the minimum and maximum of the empirical sizes over the three parametrizations considered in each test — and only one value is needed for those tests that are invariant to parametrization — when the nominal size $\alpha = 0.5\%$. We see the deviations from the nominal α can be noticeable, especially when $m = 3$. To take that into account, we report the empirical size adjusted power, that is, $O = \text{power}/\hat{\alpha}$, which also has the interpretation as (an approximated) posterior odds of H_1 to H_0 (Bayarri *et al.*, 2016). Figure 6 plots the result for size 0.5% tests. Compared with the benchmark L-0, the odds O of the proposed robust MI test (L-5) is closer to the nominal value $1/\alpha$ as $\delta \rightarrow \infty$. Nevertheless, the finite sample performances of all size 0.5% tests are less satisfactory than those for size 5% tests (given in Appendix) because a very large sample size n is required in order to approximate the tail behavior of test statistics satisfactorily.

We also compare the performance of estimators of ν_m for different δ and parametrizations. In our experiment, we have $\nu_m = 1 + 1/m$ because we have set $\nu = 1$. The MSEs of estimators $\hat{f} = \hat{r}/(1 + \hat{r})$ of $\ell_m = \nu_m/(1 + \nu_m)$ are

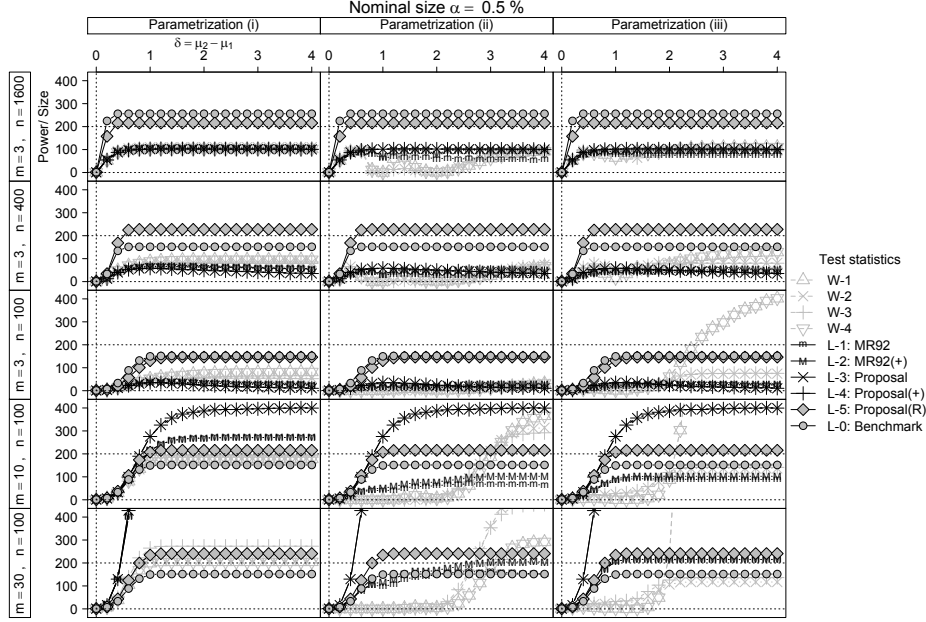


FIG 6. The ratios of empirical power to empirical size under different parametrizations. The nominal size is $\alpha = 0.5\%$. In each plot, the vertical axis denotes the ratio, whereas the horizontal axis denotes $\delta = \mu_2 - \mu_1$. The legend in Figure 4 also applies here.

shown in Figure 7, in log scale. Clearly, the only estimator that is consistent, invariant to parametrization and robust against δ is our proposal $\hat{f}_L^\diamond = \hat{r}_L^\diamond / (1 + \hat{r}_L^\diamond)$. It concentrates at the true value ℓ_m quite closely even for small m and n . Since $\hat{f}_L^\diamond$ is the only reliable estimator of ℓ_m , it verifies why L-5 has the greatest power. On the other hand, the estimator $\tilde{f}_L = \tilde{r}_L / (1 + \tilde{r}_L)$ has very large MSE when $\delta \neq 0$. It also explains why L-1 is not powerful.

4.2. Monte Carlo Experiments Without EFMI. To check how robust various tests are to the assumption of EFMI, we simulate $X_i = (X_{i1}, \dots, X_{ip})^\top \stackrel{\text{iid}}{\sim} \mathcal{N}_p(\mu, \Sigma)$ for $i = 1, \dots, n$. Let R_{ij} be defined by $R_{ij} = 1$ if X_{ij} is observed, otherwise $R_{ij} = 0$. Suppose that the first variable $X_{.1}$ is always observed, and the rest form a monotone missing pattern as defined by a logistic model on the missing propensity:

$$\mathbf{P}(R_{ij} = 0 \mid R_{i,j-1} = a) = \begin{cases} [1 + \exp(\alpha_0 + \alpha_1 X_{i,j-1})]^{-1} & \text{if } a = 1; \\ 1 & \text{if } a = 0, \end{cases}$$

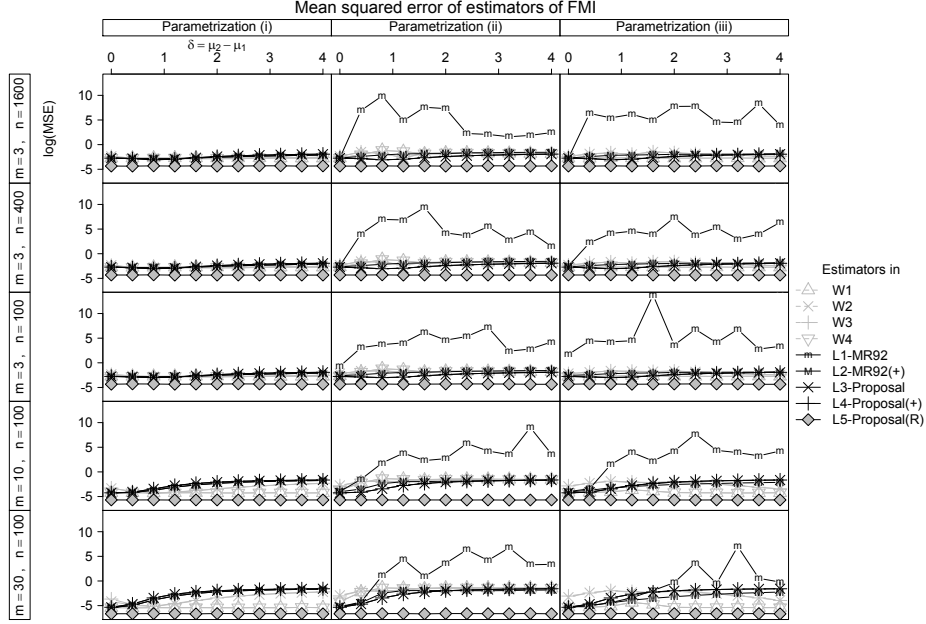


FIG 7. The MSEs of estimators of f_m used in the test statistics. In each plot, the vertical axis denotes the log of MSE, whereas the horizontal axis denotes the value of $\delta = \mu_2 - \mu_1$. The legend in Figure 4 also applies here.

for $j = 2, \dots, p$, where $\alpha_0, \alpha_1 \in \mathbb{R}$. If $\alpha_1 = 0$, then the data are missing completely at random (MCAR); otherwise they are missing at random (MAR), as defined in Rubin (1976). Let $n_j = \sum_{i=1}^n R_{ij}$ be the number of observed j th component. Without loss of generality, assume X_{obs} is arranged in such a way that $R_{ij} \geq R_{i'j}$ for all $i < i'$ and j .

To impute the missing data, it is useful to represent X_i by

$$\begin{aligned} [X_{i1} \mid \beta_1, \tau_1^2] &\sim \mathcal{N}(\beta_1, \tau_1^2), \\ [X_{ij} \mid X_{i,1:(j-1)}, \beta_j, \tau_j^2] &\sim \mathcal{N}(\beta_j^\top Z_{ij}, \tau_j^2), \quad j = 2, \dots, p, \end{aligned}$$

where $\tau_1^2, \dots, \tau_p^2 \in \mathbb{R}^+$, $\beta_j \in \mathbb{R}^j$, $X_{i,1:(j-1)} = (X_{i1}, \dots, X_{i,j-1})^\top$ and $Z_{ij} = (1, X_{i,1:(j-1)}^\top)^\top$ for $j \geq 2$. Denote the (complete-case) least squares estimators of β_j and τ_j^2 by $\hat{\beta}_j = (Z_j^\top Z_j)^{-1} Z_j^\top W_j$ and $\hat{\tau}_j^2 = \frac{1}{n_j - j} (W_j - Z_j \hat{\beta}_j)^\top (W_j - Z_j \hat{\beta}_j)$, where $Z_j = (Z_{1j}, \dots, Z_{n_j j})^\top$ and $W_j = (X_{1j}, \dots, X_{n_j j})^\top$.

To perform MI, we assume a Bayesian model with the non-informative prior $f(\beta_1, \dots, \beta_p, \tau_1^2, \dots, \tau_p^2) \propto 1/(\tau_1^2 \cdots \tau_p^2)$. For each $\ell = 1, \dots, m$, the ℓ th imputed dataset X^ℓ , whose (i, j) th element is X_{ij}^ℓ , is produced as follows.

1. Let $X_{ij}^\ell = X_{ij}$ for all $1 \leq j \leq p$ and $i \leq n_j$.
2. For each $j = 2, \dots, p$, repeat Step 3 to Step 5.
3. Draw a sample $(\tau_j^\ell)^2$ from $\hat{\tau}_j^2(n_j - j)/\chi_{n_j-j}^2$.
4. Draw a sample β_j^ℓ from $\mathcal{N}_j(\hat{\beta}_j, (\tau_j^\ell)^2(Z_j^\top Z_j)^{-1})$.
5. Draw a sample X_{ij}^ℓ from $\mathcal{N}((\beta_j^\ell)^\top Z_{ij}^\ell, (\tau_j^\ell)^2)$ for $i = n_j + 1, \dots, n$, where $Z_{ij}^\ell = (1, (X_{i,1:(j-1)}^\ell)^\top)^\top$.

We test $H_0 : \mu = \mathbf{0}_p$ against $H_1 : \mu \neq \mathbf{0}_p$. In the experiments, we set $\mu = \delta \mathbf{1}_p$, where $\delta \in [0, 0.6]$; the (i, j) th element of Σ to be $\Sigma_{ij} = 0.5^{|i-j|}$ for $i, j = 1, \dots, p$; $n = 500$; $m \in \{3, 5\}$; $p = 5$; $(\alpha_0, \alpha_1) \in \{(2, -1), (1, 0)\}$. Our model treats Σ unknown, and hence $k = p$ and $h = (3p + p^2)/2$. With the ℓ th imputed dataset, the H_0 -constrained MLEs of μ and Σ are $\hat{\mu}_0^\ell = \mathbf{0}_p$ and $\hat{\Sigma}_0^\ell = (X^\ell)^\top (X^\ell)/n$; whereas the unconstrained counterparts are $\hat{\mu}^\ell = \mathbf{1}_n^\top X^\ell/n$ and $\hat{\Sigma}^\ell = (X^\ell - \hat{\mu}^\ell)^\top (X^\ell - \hat{\mu}^\ell)/n$. Under H_0 and MAR, the fractions of missing *observations* of the five variables are (0, 16%, 28%, 38%, 47%), whereas the average fractions of missing *information*, i.e., the eigenvalues of $\mathcal{B}_\theta \mathcal{T}_\theta^{-1}$, are (0, 19%, 34%, 45%, 55%). So, the assumption of EFMI does not hold.

We compare (L4) $\hat{D}_L^+ \approx F_{k, \text{df}(\hat{\tau}_L^+, k)}$, (L5) $\hat{D}_L^\diamond \approx F_{k, \text{df}(\hat{\tau}_L^\diamond, h)}$, (C1) complete-data (asymptotic) LRT using $\{X_i : i = 1, \dots, n\}$, and (C2) complete-case (asymptotic) LRT using $\{X_i : i = 1, \dots, n_p\}$. The results are shown in Figure 8. The size of $\hat{D}_L^\diamond$ is quite accurate when the nominal size is small. If the data are MCAR, complete-case test C2 is valid, however, with slightly less power. (Test C2 is typically invalid without MCAR.) In terms of power-to-size ratio, the performance of $\hat{D}_L^\diamond$ is the best among the three implementable tests L4, L5 and C2. Its performance is comparable to the (unavailable) complete-data test C1. Note also that the power-to-size ratio of $\hat{D}_L^+$ and $\hat{D}_L^\diamond$ become closer to the nominal value 1/0.5% when m increases. All these indicate that the performance of our proposed tests are acceptable despite of the serious violation of the EFMI assumption.

4.3. Applications to a Care-Survival Data. Meng and Rubin (1992) applied their test to the data given in Table 6, where i, j and k index, respectively, clinic (A or B), amount of parental care (more or less) and survival status (died or survived). However, the clinic label k is missing for some of the observations (and the missing-data mechanism was assumed to be ignorable). Two hypotheses were tested in Meng and Rubin (1992). The first is whether the clinic and parental care are conditionally independent given the survival status, and the second is whether all three variables are mutually independent. The MI datasets are generated from a Bayesian model in § 4.2 of Meng and Rubin (1992). Our aim here is to investigate the impact on the

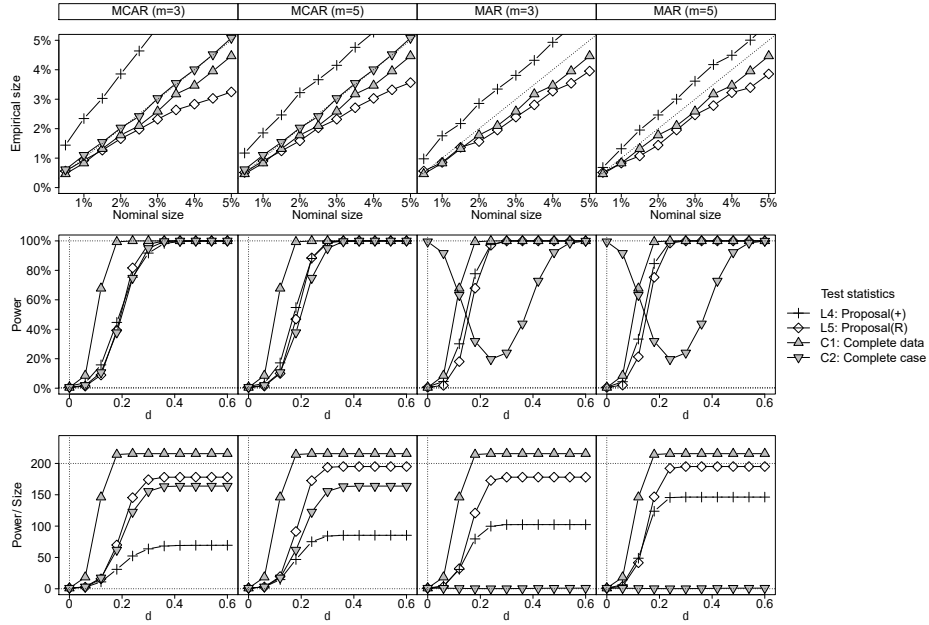


FIG 8. The empirical size, empirical power, and their ratio. The first row of plots show the empirical sizes. The size of the complete-case test (C2) under MAR is off the chart (always equals to one) in the experiment because it is invalid. The second and third rows of plots show the powers and the power-to-size ratios, respectively, where the nominal size is 0.5%.

test statistics $\tilde{D}_S$, $\hat{D}_S^+$ and $\hat{D}_S^\diamond$ by different parametrizations of $\{\pi_{ijk}\}$; and the impact on the estimators $\tilde{r}_L$, $\hat{r}_S^+$ and $\hat{r}_S^\diamond$ under different null hypotheses.

TABLE 6
Data from [Meng and Rubin \(1992\)](#). The notation “?” indicates missing label.

Clinic (k)	Parental care (i)	Survival Status (j)	
		Died	Survived
A	Less	3	176
	More	4	293
B	Less	17	197
	More	2	23
?	Less	10	150
	More	5	90

Specifically, the ℓ th imputed log-likelihood function is $\log f(X^\ell \mid \pi) = \sum_c n_c^\ell \log \pi_c$, where X^ℓ are the cell counts n_c^ℓ in the ℓ th imputed dataset.

Hence the unconstrained MLE of π_c is $\hat{\pi}_c^\ell = n_c^\ell/n_+^\ell$, where $n_+^\ell = \sum_c n_c^\ell$. Consequently, the joint log-likelihood based on the stacked data is

$$(4.1) \quad \log f(X^S | \pi) = \sum_{\ell=1}^m \sum_c n_c^\ell \log \pi_c = \sum_c n_c^+ \log \pi_c,$$

where $n_c^+ = \sum_{\ell=1}^m n_c^\ell$. Thus the unconstrained MLE with respect to (4.1) is $\hat{\pi}_c^S = n_c^+/n_+^+$, where $n_+^+ = \sum_c n_c^+$. Similarly, we can find the constrained MLEs under a given null. We consider the following parametrizations:

- (i) $\psi_{ijk} = \pi_{ijk}$ — the identity map,
- (ii) $\psi_{ijk} = \log\{\pi_{ijk}/(1 - \pi_{ijk})\}$ — the logit transformation, and
- (iii) $\psi_{ij1} = \pi_{ij1}$ and $\psi_{ij2} = \pi_{ij2}/\pi_{ij1}$ — ratios of probabilities.

The p -values $\tilde{p}_L$, $\hat{p}_S^+$ and $\hat{p}_S^\diamond$ of the tests $\tilde{D}_L$, $\hat{D}_S^+$ and $\hat{D}_S^\diamond$ and the associated estimates of ν_m , i.e., $\tilde{r}_L$, $\hat{r}_S^+$ and $\hat{r}_S^\diamond$, are shown in Table 7. A more detailed comparison is deferred to Table 9 of the Appendix.

TABLE 7
The LRTs using $\tilde{D}_L$, $\hat{D}_S^+$ and $\hat{D}_S^\diamond$ under different parametrizations in § 4.3.

Parametrization (i)							
H_0 : Conditional independence				H_0 : Full independence			
m	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$	
3	0.54, 0.54, 0.38	0.08, 0.08, 0.09	0.93, 0.93, 0.92	0.31, 0.31, 0.38	54.2, 54.2, 51.4	0, 0, 0	
10	0.50, 0.50, 0.70	0.14, 0.14, 0.12	0.87, 0.87, 0.88	0.56, 0.56, 0.70	45.4, 45.4, 41.7	0, 0, 0	
50	0.31, 0.31, 0.45	0.11, 0.11, 0.10	0.90, 0.90, 0.91	0.33, 0.33, 0.45	51.5, 51.5, 47.3	0, 0, 0	
Parametrization (ii)							
H_0 : Conditional independence				H_0 : Full independence			
m	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$	
3	1.08, 0.54, 0.38	-0.07, 0.08, 0.09	1.00, 0.93, 0.92	0.61, 0.31, 0.38	43.9, 54.2, 51.4	0, 0, 0	
10	0.99, 0.50, 0.70	-0.10, 0.14, 0.12	1.00, 0.87, 0.88	1.09, 0.56, 0.70	33.7, 45.4, 41.7	0, 0, 0	
50	0.63, 0.31, 0.45	-0.10, 0.11, 0.10	1.00, 0.90, 0.91	0.65, 0.33, 0.45	41.3, 51.5, 47.3	0, 0, 0	
Parametrization (iii)							
H_0 : Conditional independence				H_0 : Full independence			
m	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$	
3	-2.35, 0.54, 0.38	-1.16, 0.08, 0.09	1.00, 0.93, 0.92	-1.22, 0.31, 0.38	-321, 54.2, 51.4	1, 0, 0	
10	-2.04, 0.50, 0.70	-2.20, 0.14, 0.12	1.00, 0.87, 0.88	-1.85, 0.56, 0.70	-86, 45.4, 41.7	1, 0, 0	
50	-1.22, 0.31, 0.45	-7.39, 0.11, 0.10	1.00, 0.90, 0.91	-1.06, 0.33, 0.45	-1136, 51.5, 47.3	1, 0, 0	

The simulation outputs demonstrate that $\hat{D}_S^+$ and $\hat{D}_S^\diamond$ are invariant to parametrizations, whereas $\tilde{D}_L$ is not. Moreover, the impact on $\tilde{D}_L$ is large

under the parametrization (iii). In particular, the value of $\tilde{r}_L$ is inflated; and some of the values of $\tilde{r}_L$ and $\tilde{D}_L$ are negative, leading to the meaningless $\tilde{p}_L = 1$, especially under parametrization (iii). In contrast, $\hat{r}_S \geq 0$ for all cases in this example (and hence $\hat{r}_S^+ = \hat{r}_S$). In addition, $\hat{D}_S^+ \approx \hat{D}_S^\diamond$ for testing the conditional independence, a hypothesis that is not rejected by either test. In contrast, for testing the full independence, $\hat{D}_S^+$ and $\hat{D}_S^\diamond$ are not very close to each other, but they both lead to essentially zero p -value, and hence both reject the null hypothesis. These results reconfirm the conclusions in [Meng and Rubin \(1992\)](#). Last but not least, the estimator $\hat{r}_S^\diamond$ does not change under different null hypotheses, however it is not true for $\tilde{r}_L$ and $\tilde{r}_S^+$.

5. Conclusions, Limitations and Future Work. In addition to conducting a general comparative study of MI tests, we proposed two particularly promising MI LRT based on $\hat{D}_S^\diamond = \hat{D}_S(\hat{r}_S^\diamond)$ and $\hat{D}_S^+ = \hat{D}_S(\hat{r}_S^+)$. Both test statistics are non-negative, invariant to re-parametrizations, and powerful to reject a false null hypothesis (at least for large enough m). Test $\hat{D}_S^\diamond$ is most principled, and the resulting test has the desirable monotonically increasing power as H_1 departs from H_0 . However, it is derived under the stronger assumption of EFMI for ψ , not just for θ ; and row independence of X_{com} is needed for the ease of computation. (The computationally more demanding test based on $\hat{D}_L(\hat{r}_L^\diamond)$ relaxes the independence assumption.) The main advantage of $\hat{D}_S^+$ is that it is easier to compute, as it requires only standard complete-data computer subroutines for likelihood ratio tests. One drawback is that the ad hoc fix $\hat{r}_S^+ = \max(0, \hat{r}_S)$ is inconsistent in general. However, the inconsistency does not appear to significantly affect the asymptotic power, at least in our experiments. Whereas $\hat{D}_S^+$ and $\hat{D}_S^\diamond$ significantly improve over existing counterparts, more studies are needed, for reasons listed below.

- When the missing data mechanism is not ignorable but the imputers fail to fully take that into account, the issue of uncongeniality becomes critical ([Meng, 1994a](#)). [Xie and Meng \(2017\)](#) provides theoretical tools for addressing such an issue in the context of estimation, and research is needed to extend their findings to the setting of hypothesis testing.
- Although the violation of the EFMI assumption may not (seriously) invalidate a test, it will affect its power. It is therefore desirable to explore MI tests without this assumption.
- The robust $\hat{D}_S^\diamond$ relies on a stronger assumption of EFMI on ψ . We can modify it so only EFMI on θ is required, but the modification may be very difficult to compute and may require users to have access to non-trivial complete-data procedures. Hence a computational feasible robust test that only assumes EFMI on θ needs to be developed.

- Because the FMI is a fundamental nuisance parameter here and there is no (known) pivotal quantity, all MI tests are approximate in nature. In particular, they all have the potential of doing poorly when FMI is large and/or m is small. It is therefore of both theoretical and practical interest to seek powerful MI tests that are least affected by FMI.

References.

- Barnard, J. and Rubin, D. B. (1999) Small-sample degrees of freedom with multiple imputation. *Biometrika*, **86**, 948–955.
- Bayarri, M. J., Benjamin, D. J., Berger, J. O. and Sellke, T. M. (2016) Rejection odds and rejection ratios: A proposal for statistical practice in testing hypotheses. *Journal of Mathematical Psychology*, **72**, 90–103.
- Benjamin, D. J., Berger, J. O., Johannesson, M., Nosek, B. A., Wagenmakers, E.-J., Berk, R., Bollen, K. A., Brembs, B., Brown, L., Camerer, C. *et al.* (2018) Redefine statistical significance. *Nature Human Behaviour*, **2**, 6–10.
- Berglund, P. and Heeringa, S. G. (2014) *Multiple imputation of missing data using SAS*. SAS Institute.
- Blocker, A. W. and Meng, X.-L. (2013) The potential and perils of preprocessing: Building new foundations. *Bernoulli*, **19**, 1176–1211.
- Carlin, J. B., Galati, J. C. and Royston, P. (2008) A new framework for managing and analyzing multiply imputed data in stata. *The Stata Journal*, **8**, 49–67.
- Cox, D. R. and Reid, N. (1987) Parameter orthogonality and approximate conditional inference. *Journal of the Royal Statistical Society B*, **49**, 1–39.
- Grund, S., Robitzsch, A. and Luedtke, O. (2017) *Tools for Multiple Imputation in Multi-level Modeling*.
- Harel, O. and Zhou, X.-H. (2007) Multiple imputation - review of theory, implementation and software. *Statistics in medicine*, **26**, 3057–3077.
- Holan, S. H., Toth, D., Ferreira, M. A. R. and F., K. A. (2010) Bayesian multiscale multiple imputation with implications for data confidentiality. *Journal of the American Statistical Association*, **105**, 564–577.
- Horton, N. and Kleinman, K. P. (2007) Much ado about nothing: A comparison of missing data methods and software to fit incomplete data regression models. *The American Statistician*, **61**, 79–90.
- Kenward, M. G. and Carpenter, J. R. (2007) Multiple imputation: current perspectives. *Statistical Methods in Medical Research*, **16**, 199–218.
- Kim, J. K. and Shao, J. (2013) *Statistical Methods for Handling Incomplete Data*. Chapman and Hall/CRC.
- Kim, J. K. and Yang, S. (2017) A note on multiple imputation under complex sampling. *Biometrika*, **104**, 221–228.
- King, G., Honaker, J., Joseph, A. and Scheve, K. (2001) Analyzing incomplete political science data: An alternative algorithm for multiple imputation. *American Political Science Review*, **95**, 49–69.
- Li, K. H., Meng, X.-L., Raghunathan, T. E. and Rubin, D. B. (1991a) Significance levels from repeated p -values with multiply-imputed data. *Statistica Sinica*, **1**, 65–92.
- Li, K. H., Raghunathan, T. E. and Rubin, D. B. (1991b) Large-sample significance levels from multiply imputed data using moment-based statistics and an F reference distribution. *Journal of the American Statistical Association*, **86**, 1065–1073.
- Medeiros, R. (2008) Likelihood ratio tests for multiply imputed datasets: Introducing milrtest.

- Meng, X.-L. (1994a) Multiple-imputation inferences with uncongenial sources of input. *Statistical Science*, **9**, 538–573.
- Meng, X.-L. (1994b) Posterior predictive p -values. *The Annals of Statistics*, **22**, 1142–1160.
- Meng, X.-L. (2002) Discussion of “Bayesian measures of model complexity and fit” by Spiegelhalter, D. J., Best, N. G., Carlin, B. P. and Van Der Linde, A. *Journal of the Royal Statistical Society B*, **64**, 633.
- Meng, X.-L. and Rubin, D. B. (1992) Performing likelihood ratio tests with multiply-imputed data sets. *Biometrika*, **79**, 103–111.
- Peugh, J. L. and Enders, C. K. (2004) Missing data in educational research: A review of reporting practices and suggestions for improvement. *Review of Educational Research*, **74**, 525–556.
- Rose, R. A. and Fraser, M. W. (2008) A simplified framework for using multiple imputation in social work research. *Social Work Research*, **32**, 171–178.
- Royston, P. and White, I. R. (2011) Multiple imputation by chained equations (mice): Implementation in stata. *Journal of Statistical Software*, **45**, 1–20.
- Rubin, D. B. (1976) Inference and missing data. *Biometrika*, **63**, 581–592.
- Rubin, D. B. (1978) Multiple imputations in sample surveys - a phenomenological Bayesian approach to nonresponse. *Proceedings of the Survey Research Methods Section of the American Statistical Association*, 20–34.
- Rubin, D. B. (1996) Multiple imputation after 18+ years. *Journal of the American statistical Association*, **91**, 473–489.
- Rubin, D. B. (2004) *Multiple Imputation for Nonresponse in Surveys*. John Wiley & Sons.
- Rubin, D. B. and Schenker, N. (1986) Multiple imputation for interval estimation from simple random samples with ignorable nonresponse. *Journal of the American Statistical Association*, **81**, 366–374.
- Schafer, J. L. (1999) Multiple imputation: A primer. *Statistical Methods in Medical Research*, **8**, 3–15.
- Serfling, R. J. (2001) *Approximation Theorems of Mathematical Statistics*. Wiley-Interscience.
- Su, Y.-S., Gelman, A., Hill, J. and Yajima, M. (2011) Multiple imputation with diagnostics (mi) in R: Opening windows into the black box. *Journal of Statistical Software*, **45**, 1–31.
- Tu, X. M., Meng, X.-L. and Pagano, M. (1993) The AIDS epidemic: Estimating survival after AIDS diagnosis from surveillance data. *Journal of the American Statistical Association*, **88**, 26–36.
- van Buuren, S. and Groothuis-Oudshoorn, K. (2011) Mice: Multivariate imputation by chained equations in R. *Journal of Statistical Software*, **45**, 1–67.
- van Buuren S (2012) *Flexible Imputation of Missing Data*. Chapman and Hall/CRC.
- van der Vaart, A. W. (2000) *Asymptotic Statistics*. Cambridge University Press.
- Wallace, D. L. (1980) The Behrens-Fisher and Fieller-Creasy problems. In *R. A. Fisher: An Appreciation* (eds. S. E. Fienberg and D. V. Hinkley), 119–147. Springer New York.
- Wang, N. and Robins, J. M. (1998) Large-sample theory for parametric multiple imputation procedures. *Biometrika*, **85**, 935–948.
- Xie, X. and Meng, X.-L. (2017) Dissecting multiple imputation from a multi-phase inference perspective: What happens when God’s, imputer’s and analyst’s models are uncongenial? (with discussion). *Statistica Sinica*, **27**, 1485–1594.

APPENDIX A: SUPPLEMENTARY RESULTS

A.1. Another Motivation for $\hat{r}_L^\diamond$. The definition of $\hat{r}_L^\diamond$ can also be motivated by the following observation. First, observe that one simple method to construct an always non-negative estimator of ν_m is to perturb $\hat{\psi}_0^*$ and $\hat{\psi}_0^\ell$ by a suitable amount, say Δ , so that the perturbed version of $\hat{r}_L$ is always non-negative, and is still asymptotically equivalent to the original $\hat{r}_L$. We show, in Theorem A.1 below, that the right amount of Δ is $\Delta = \hat{\psi}^* - \hat{\psi}_0^*$. Using the perturbed version of $\hat{r}_L$, we obtain

$$\hat{r}_L^\Delta = \frac{m+1}{k(m-1)} \hat{\delta}_L^\Delta,$$

where

$$\hat{\delta}_L^\Delta = \frac{2}{m} \sum_{\ell=1}^m \log \left\{ \frac{f(X^\ell | \hat{\psi}^\ell) f(X^\ell | \hat{\psi}_0^* + \Delta)}{f(X^\ell | \hat{\psi}^*) f(X^\ell | \hat{\psi}_0^\ell + \Delta)} \right\} = \frac{1}{m} \sum_{\ell=1}^m d_L(\hat{\psi}_0^\ell + \Delta, \hat{\psi}^\ell | X^\ell).$$

Then we have the following result.

THEOREM A.1. *Suppose RC_θ . Under H_0 , we have (i) $\hat{r}_L^\Delta \geq 0$ for all m, n ; and (ii) $\hat{r}_L^\Delta \simeq \hat{r}_L$ as $n \rightarrow \infty$ for each m .*

Although $\hat{r}_L^\Delta \geq 0$, it is only invariant to affine transformations, and not robust against θ_0 , and less computational feasible than $\hat{r}_L$; see § 3. However, it gives us some insights on how to construct a potentially better estimator. Note that, in (A.1), the constrained MLE is not used in $d_L(\cdot, \cdot | X^\ell)$, but it is still always non-negative. We call this a “pseudo” LRT statistics. Then, $\hat{\delta}_L^\Delta$ is just a multiple of an average of many “pseudo” LRT statistics. In order to find a good estimator of ν_m , we may seek for an estimator which admits this form. Indeed, our proposed estimator $\hat{r}_L^\diamond$ also takes the same form:

$$\hat{r}_L^\diamond = \frac{m+1}{h(m-1)} \frac{1}{m} \sum_{\ell=1}^m d_L(\hat{\psi}^*, \hat{\psi}^\ell | X^\ell).$$

A.2. Results for Dependent Data. This is a supplement for § 3.1. If the data are not independent, $\hat{d}_L \simeq \hat{d}_S$ is still true under the following conditions.

ASSUMPTION 5. (a) Define $R(\psi) = \bar{L}^S(\psi) - \bar{L}(\psi)$, where $\bar{L}(\psi) = (mn)^{-1} \sum_{\ell=1}^m \log f(X^\ell | \psi)$ and $\bar{L}^S(\psi) = (mn)^{-1} \log f(X^S | \psi)$. For each m , as $n \rightarrow \infty$,

$$\sup_{\psi \in \Psi} |R(\psi)| = O_p(1/n), \quad \sup_{\psi \in \Psi} \left| \frac{\partial}{\partial \psi} R(\psi) \right| = O_p(1/n).$$

(b) For each m , there exists a continuous function $\psi \mapsto \overline{\mathcal{L}}(\psi)$, which is free of n but may depend on m , such that, as $n \rightarrow \infty$,

$$\sup_{\psi \in \Psi} |\overline{\mathcal{L}}(\psi) - \underline{\mathcal{L}}(\psi)| = o_p(1).$$

(c) Let $\psi_0^* = \arg \max_{\psi \in \Psi : \psi(\theta) = \theta_0} \overline{\mathcal{L}}(\psi)$ and $\psi^* = \arg \max_{\psi \in \Psi} \overline{\mathcal{L}}(\psi)$. For any fixed m , and for all $\varepsilon > 0$, there exists $\delta > 0$ such that

$$\sup_{\substack{\psi \in \Psi : |\psi_0^* - \psi| > \varepsilon \\ \theta(\psi) = \theta_0}} \{\overline{\mathcal{L}}(\psi_0^*) - \overline{\mathcal{L}}(\psi)\} \geq \delta, \quad \sup_{\psi \in \Psi : |\psi^* - \psi| > \varepsilon} \{\overline{\mathcal{L}}(\psi^*) - \overline{\mathcal{L}}(\psi)\} \geq \delta.$$

Conditions (b) and (c) in Assumption 5 are standard RCs that are usually assumed for M-estimators (see § 5 of [van der Vaart \(2000\)](#)); whereas condition (a) is satisfied by many models (see Example A.1 below).

THEOREM A.2. *Suppose RC_θ and Assumption 5. Under both H_0 and H_1 , we have (i) $\hat{d}_S, \hat{r}_S \geq 0$ for all m, n ; (ii) $\hat{d}_S, \hat{r}_S$ are invariant to the parametrization of ψ for all m, n ; and (iii) $\hat{d}_L \simeq \hat{d}_S$ and $\hat{r}_L \simeq \hat{r}_S$ as $n \rightarrow \infty$ for each m .*

Theorem A.2 implies that the handy test statistics $\hat{D}_S$ and $\hat{D}_S^+$ approximate $\hat{D}_L$ and $\hat{D}_L^+$ for dependent data, provided that Assumption 5 holds.

EXAMPLE A.1. Consider a stationary autoregressive model of order one. Suppose the complete data $X = (X_1, \dots, X_n)^\top$ is generated as following: $X_1 \sim \mathcal{N}(0, v^2)$ and $[X_i | X_{i-1}] \sim \mathcal{N}(\phi X_{i-1}, \sigma^2)$ for $i \geq 2$, where $v^2 = \sigma^2(1 + \phi)/(1 - \phi)$. Then $\psi = (\phi, \sigma^2)^\top$, and

$$\begin{aligned} \overline{\mathcal{L}}(\psi) &= -\frac{1}{2} \log(2\pi) - \frac{1}{2n} \log v^2 - \frac{1}{mn} \sum_{\ell=1}^m \frac{X_1^\ell}{2v^2} - \frac{n-1}{2n} \log \sigma^2 \\ &\quad - \frac{1}{mn} \sum_{\ell=1}^m \sum_{i=2}^n \frac{(X_i^\ell - \phi X_{i-1}^\ell)^2}{2\sigma^2}, \\ \overline{\mathcal{L}}^S(\psi) &= -\frac{1}{2} \log(2\pi) - \frac{1}{2mn} \log v^2 - \frac{(X_1^1)^2}{2mnv^2} - \frac{mn-1}{2mn} \log \sigma^2 \\ &\quad - \frac{1}{mn} \sum_{\ell=1}^m \sum_{i=2}^n \frac{(X_i^\ell - \phi X_{i-1}^\ell)^2}{2\sigma^2} - \frac{1}{mn} \sum_{\ell=2}^m \frac{(X_1^\ell - \phi X_n^{\ell-1})^2}{2\sigma^2}. \end{aligned}$$

Then, it is easy to see that condition (a) of Assumption 5 is satisfied.

APPENDIX B: PROOFS

PROOF OF THEOREM 2.1. (i, ii) From (2.3), we know $\hat{d}_L \geq 0$ is invariant to parametrization ψ . (iii) Since $\hat{d}_L$ is invariant to transformation of ψ , we assume, without loss of generality, that ψ admits a parameterization such that $\text{Cov}(\hat{\theta}^\ell, \hat{\eta}^\ell) \simeq \mathbf{0}$ by taking suitable linear transformation of ψ . Also write U_η^ℓ as an efficient estimator of $\text{Var}(\hat{\eta})$ based on X^ℓ ; and recall that $U_\theta^\ell = U^\ell$ is an efficient estimator of $\text{Var}(\hat{\theta})$ based on X^ℓ .

Using Taylor's expansion on $\psi \mapsto \bar{L}(\psi) = m^{-1} \sum_{\ell=1}^m \log f(X^\ell | \psi)$ around $\hat{\psi}^* = ((\hat{\theta}^*)^\top, (\hat{\eta}^*)^\top)^\top$, we know that for $\psi \simeq \hat{\psi}^*$,

$$(B.1) \quad \bar{L}(\psi) \simeq \bar{L}(\hat{\psi}^*) - \frac{1}{2} (\psi - \hat{\psi}^*)^\top \bar{I}(\hat{\psi}^*) (\psi - \hat{\psi}^*),$$

where $\bar{I}(\psi) = -\partial^2 \bar{L}(\psi) / \partial \psi \partial \psi^\top$, which satisfies

$$(B.2) \quad \bar{I}(\hat{\psi}^*) \simeq \begin{pmatrix} \bar{U}_\theta^{-1} & \mathbf{0} \\ \mathbf{0} & \bar{U}_\eta^{-1} \end{pmatrix}$$

with $\bar{U}_\eta = m^{-1} \sum_{i=1}^m U_\eta^\ell$. Under the null, $\hat{\psi}^* \simeq \hat{\psi}_0^*$. So, using (B.1), we have

$$(B.3) \quad \begin{aligned} \hat{d}_L &\simeq (\hat{\psi}_0^* - \hat{\psi}^*)^\top \bar{I}(\hat{\psi}^*) (\hat{\psi}_0^* - \hat{\psi}^*), \\ &\simeq \begin{pmatrix} \theta_0 - \hat{\theta}^* \\ \hat{\eta}(\theta_0) - \hat{\eta}(\hat{\theta}^*) \end{pmatrix}^\top \begin{pmatrix} \bar{U}_\theta^{-1} & \mathbf{0} \\ \mathbf{0} & \bar{U}_\eta^{-1} \end{pmatrix} \begin{pmatrix} \theta_0 - \hat{\theta}^* \\ \hat{\eta}(\theta_0) - \hat{\eta}(\hat{\theta}^*) \end{pmatrix} \\ &\simeq (\bar{\theta}^\top - \theta_0) \bar{U}_\theta^{-1} (\bar{\theta}^\top - \theta_0) = \tilde{d}_W, \end{aligned}$$

where we have used (a) $\hat{\theta}^* \simeq \bar{\theta}$; see, e.g., Lemma 1 of Wang and Robins (1998), and (b) $\hat{\eta}(\theta_0) - \hat{\eta}(\hat{\theta}^*) = O_p(1/n)$ if $\theta_0 - \hat{\theta}^* = O_p(1/\sqrt{n})$; see Cox and Reid (1987). Since $\tilde{d}_W \simeq \tilde{d}_L$ (Meng and Rubin, 1992), we have $\hat{d}_L \simeq \tilde{d}_L$. $\square$

PROOF OF PROPOSITION 2.2. The given condition implies that

$$\begin{aligned} \hat{\psi}^\ell &= ((\hat{\theta}^\ell)^\top, (\hat{\eta}^\ell)^\top)^\top, & \hat{\psi}_0^\ell &= (\theta_0^\top, (\hat{\eta}^\ell)^\top)^\top, \\ \hat{\psi}^* &= ((\hat{\theta}^*)^\top, (\hat{\eta}^*)^\top)^\top, & \hat{\psi}_0^* &= (\theta_0^\top, (\hat{\eta}^*)^\top)^\top. \end{aligned}$$

Clearly, we also have the decomposition: $L^\ell(\psi) = L_\dagger^\ell(\theta) + L_\ddagger^\ell(\eta)$ for all ℓ , where $L_\dagger^\ell(\theta) = L_\dagger(\theta | X^\ell)$ and $L_\ddagger^\ell(\eta) = L_\ddagger(\eta | X^\ell)$. Then,

$$\begin{aligned} \bar{d}_L - \hat{d}_L &= \frac{2}{m} \sum_{\ell=1}^m \left\{ L^\ell(\hat{\psi}^\ell) - L^\ell(\hat{\psi}_0^\ell) - L^\ell(\hat{\psi}^*) + L^\ell(\hat{\psi}_0^*) \right\} \\ &= \frac{2}{m} \sum_{\ell=1}^m \left\{ L_\dagger^\ell(\hat{\theta}^\ell) - L_\dagger^\ell(\hat{\theta}^*) \right\} \geq 0 \end{aligned}$$

since $L_{\dagger}^{\ell}(\hat{\theta}^{\ell}) \geq L_{\dagger}^{\ell}(\hat{\theta}^*)$ for all ℓ . $\square$

PROOF OF COROLLARY 2.3. Applying Taylor's expansion on $\psi \mapsto L^{\ell}(\psi)$, we can find $\check{\psi}^{\ell}$ lying on the line segment joining $\hat{\psi}^{\ell}$ and $\hat{\psi}_0^{\ell}$ such that

$$L^{\ell}(\hat{\psi}_0^{\ell}) = L^{\ell}(\hat{\psi}^{\ell}) - \frac{1}{2} \left(\hat{\psi}_0^{\ell} - \hat{\psi}^{\ell} \right)^{\top} I^{\ell}(\check{\psi}^{\ell}) \left(\hat{\psi}_0^{\ell} - \hat{\psi}^{\ell} \right),$$

where $I^{\ell}(\psi) = -\partial^2 L^{\ell}(\psi) / \partial \psi \partial \psi^{\top}$. By the lower order variability of $I^{\ell}(\check{\psi}^{\ell})$, we can find $\check{\psi}^*$ such that $I^{\ell}(\check{\psi}^{\ell}) \simeq I^{\ell}(\check{\psi}^*)$ for all ℓ . Then, using similar techniques as in (B.2) and (B.3), we have

$$\begin{aligned} L^{\ell}(\hat{\psi}^{\ell}) - L^{\ell}(\hat{\psi}_0^{\ell}) &\simeq \frac{1}{2} \left(\hat{\psi}_0^{\ell} - \hat{\psi}^{\ell} \right)^{\top} I^{\ell}(\check{\psi}^*) \left(\hat{\psi}_0^{\ell} - \hat{\psi}^{\ell} \right) \\ (B.4) \quad &\simeq \frac{1}{2} \left(\theta_0 - \hat{\theta}^{\ell} \right)^{\top} \check{U}^{-1} \left(\theta_0 - \hat{\theta}^{\ell} \right) \end{aligned}$$

for some matrix $\check{U}$. Similarly, we have

$$(B.5) \quad L^{\ell}(\hat{\psi}^*) - L^{\ell}(\hat{\psi}_0^*) \simeq \frac{1}{2} \left(\theta_0 - \hat{\theta}^* \right)^{\top} \check{U}^{-1} \left(\theta_0 - \hat{\theta}^* \right).$$

Write $A^{\otimes 2} = AA^{\top}$ for any appropriate matrix A . Using (B.4), (B.5) and the cyclic property of trace, we have

$$\begin{aligned} \bar{d}_L - \hat{d}_L &\simeq \frac{1}{m} \sum_{\ell=1}^m \left\{ \left(\theta_0 - \hat{\theta}^{\ell} \right)^{\top} \check{U}^{-1} \left(\theta_0 - \hat{\theta}^{\ell} \right) - \left(\theta_0 - \hat{\theta}^* \right)^{\top} \check{U}^{-1} \left(\theta_0 - \hat{\theta}^* \right) \right\} \\ &= \text{tr} \left[\check{U}^{-1} \left\{ \frac{1}{m} \sum_{\ell=1}^m \left(\theta_0 - \hat{\theta}^{\ell} \right)^{\otimes 2} - \left(\theta_0 - \hat{\theta}^* \right)^{\otimes 2} \right\} \right] \\ &\simeq \text{tr} \left[\check{U}^{-1} \frac{1}{m} \sum_{\ell=1}^m \left\{ \left(\hat{\theta}^{\ell} \right)^{\otimes 2} - \bar{\theta}^{\otimes 2} \right\} \right] \simeq \text{tr} \left(\check{U}^{-1} B \right) \simeq \text{tr} \left(\mathcal{U}_{\theta,0}^{-1} \mathcal{B}_{\theta} \right) \end{aligned}$$

as $m, n \rightarrow \infty$, where $\mathcal{U}_{\theta,0}$ is a deterministic matrix that depends on both θ_0 and θ^* , and satisfies $n(\check{U} - \mathcal{U}_{\theta,0}) \xrightarrow{\text{pr}} 0$. Note that $\text{tr}(\mathcal{U}_{\theta,0}^{-1} \mathcal{B}_{\theta}) = k\kappa_0$, for some finite κ_0 by Assumption 2. Then $\hat{r}_L \xrightarrow{\text{pr}} \kappa_0 = \text{tr}(\mathcal{U}_{\theta,0}^{-1} \mathcal{B}_{\theta})/k$, proving (ii). (But $\mathcal{U}_{\theta,0}$ may not equal to $\mathcal{U}_{\theta}$, and hence $\hat{r}_L$ may not be consistent for κ_m .)

If H_0 is true, then $\bar{\theta} \xrightarrow{\text{pr}} \theta_0 = \theta^*$ and $\check{U} \simeq \bar{U} \simeq \mathcal{U}_{\theta} = \mathcal{U}_{\theta,0}$. Then, $\hat{r}_L \xrightarrow{\text{pr}} \kappa$ as $m, n \rightarrow \infty$. So, (i) follows. $\square$

PROOF OF THEOREM 2.4. (i, ii) It is trivial by the definition of $\hat{r}_L^{\diamond}$. (iii) Applying Taylor's expansion to $\psi \mapsto L^{\ell}(\psi)$ again, we know there is $\check{\psi}^{\ell}$ lying

on the line segment joining $\hat{\psi}^\ell$ and $\hat{\psi}^*$ such that

$$(B.6) \quad L^\ell(\hat{\psi}^*) = L^\ell(\hat{\psi}^\ell) - \frac{1}{2} \left(\hat{\psi}^* - \hat{\psi}^\ell \right)^\top I^\ell(\check{\psi}^\ell) \left(\hat{\psi}^* - \hat{\psi}^\ell \right).$$

By the lower order variability of $I^\ell(\check{\psi}^\ell)$, we know that $I^\ell(\check{\psi}^\ell) \simeq \bar{I}(\hat{\psi}^*)$ for all ℓ , where $\bar{I}(\psi) = m^{-1} \sum_{\ell=1}^m I^\ell(\psi)$. We also know that $\hat{\psi}^* \simeq \bar{\psi}$. Thus

$$(B.7) \quad \begin{aligned} \bar{\delta}_L - \hat{\delta}_L &\simeq \frac{1}{m} \sum_{\ell=1}^m \left(\hat{\psi}^* - \hat{\psi}^\ell \right)^\top \bar{I}(\hat{\psi}^*) \left(\hat{\psi}^* - \hat{\psi}^\ell \right) \\ &= \text{tr} \left\{ \bar{I}(\hat{\psi}^*) \frac{1}{m} \sum_{\ell=1}^m \left(\hat{\psi}^* - \hat{\psi}^\ell \right)^{\otimes 2} \right\} \\ &\simeq \text{tr} \left\{ \bar{I}(\hat{\psi}^*) \frac{1}{m} \sum_{\ell=1}^m \left(\hat{\psi}^\ell - \bar{\psi} \right)^{\otimes 2} \right\} \simeq \text{tr} \left(\mathcal{U}_\psi^{-1} \mathcal{B}_\psi \right) \end{aligned}$$

as $m, n \rightarrow \infty$. By the assumption of EFMI of ψ , we have $\hat{r}_L^\diamond \xrightarrow{\text{pr}} \boldsymbol{r}$. $\square$

PROOF OF LEMMA 2.5. First, recall that, as $n \rightarrow \infty$, the observed data MLE $\hat{\theta}_{\text{obs}}$ of θ satisfies (2.4), which can be written as $[\hat{\theta}_{\text{obs}} \mid \theta^*] \stackrel{\mathcal{D}}{\approx} \mathcal{N}_k(\theta^*, \mathcal{T}_\theta)$, where $A_{1,n} \stackrel{\mathcal{D}}{\approx} A_{2,n}$ means that $A_{1,n}$ and $A_{2,n}$ have the same asymptotic distribution, i.e., there exist deterministic sequences μ_n and Σ_n such that $(A_{1,n} - \mu_n) \Sigma_n^{-1/2} \Rightarrow A$ and $(A_{2,n} - \mu_n) \Sigma_n^{-1/2} \Rightarrow A$ for some non-degenerate random variable A . From Assumption 3, a proper imputation model is used. So, we have (2.5), which is equivalent to say that, as $n \rightarrow \infty$,

$$(B.8) \quad \left[\hat{\theta}^\ell \mid X_{\text{obs}} \right] \stackrel{\mathcal{D}}{\approx} \mathcal{N}_k(\hat{\theta}_{\text{obs}}, \mathcal{B}_\theta),$$

independently for $\ell = 1, \dots, m$. Therefore we can represent

$$(B.9) \quad \hat{\theta}_{\text{obs}} \stackrel{\mathcal{D}}{\approx} \theta^* + \mathcal{T}_\theta^{1/2} W,$$

$$(B.10) \quad \hat{\theta}^\ell \stackrel{\mathcal{D}}{\approx} \hat{\theta}_{\text{obs}} + \mathcal{B}_\theta^{1/2} Z_\ell, \quad \ell = 1, \dots, m$$

where $Z_1, \dots, Z_m, W \stackrel{\text{iid}}{\sim} \mathcal{N}_k(0, I_k)$. Also write $Z_\ell = (Z_{1\ell}, \dots, Z_{k\ell})^\top$, for $\ell = 1, 2, \dots, m$, and $W = (W_1, \dots, W_k)^\top$. Averaging (B.10) over ℓ , we have $\bar{\theta} \stackrel{\mathcal{D}}{\approx} \hat{\theta}_{\text{obs}} + \mathcal{B}_\theta^{1/2} \bar{Z}_\bullet$, where $\bar{Z}_\bullet = m^{-1} \sum_{\ell=1}^m Z_\ell$. Since $\mathcal{B}_\theta = \boldsymbol{r} \mathcal{U}_\theta$, we have

$$\begin{aligned} \mathcal{U}_\theta^{-1/2}(\hat{\theta}^\ell - \theta^*) &\stackrel{\mathcal{D}}{\approx} (1 + \boldsymbol{r})^{1/2} W + \boldsymbol{r}^{1/2} Z_\ell, \\ \mathcal{U}_\theta^{-1/2}(\bar{\theta} - \theta^*) &\stackrel{\mathcal{D}}{\approx} (1 + \boldsymbol{r})^{1/2} W + \boldsymbol{r}^{1/2} \bar{Z}_\bullet. \end{aligned}$$

Note that (2.6) implies $\mathcal{U}_\theta \simeq \bar{U}$. Under H_0 , we have $\theta^\star = \theta_0$ and

$$\begin{aligned}\bar{d}_L &\simeq \bar{d}'_W \stackrel{\mathcal{D}}{\approx} \sum_{i=1}^k \left\{ (1 + \varkappa)^{1/2} W_i + \varkappa^{1/2} Z_{i\ell} \right\}^2, \\ \hat{d}_L &\simeq \tilde{d}_L \simeq \tilde{d}'_W \stackrel{\mathcal{D}}{\approx} \sum_{i=1}^k \left\{ (1 + \varkappa)^{1/2} W_i + \varkappa^{1/2} \bar{Z}_i \right\}^2.\end{aligned}$$

After some simple algebra, we obtain

$$\hat{r}_L^+ \stackrel{\mathcal{D}}{\approx} \frac{(m+1)\varkappa}{mk} \sum_{i=1}^k s_{Z_i}^2 \quad \text{and} \quad \hat{D}_L^+ \stackrel{\mathcal{D}}{\approx} \frac{m \sum_{i=1}^k \left\{ (1 + \varkappa)^{1/2} W_i + \varkappa^{1/2} \bar{Z}_{i\bullet} \right\}^2}{mk + (m+1)\varkappa \sum_{i=1}^k s_{Z_i}^2},$$

where $s_{Z_i}^2 = (m-1)^{-1} \sum_{\ell=1}^m (Z_{i\ell} - \bar{Z}_{i\bullet})^2$ is the sample variance of $\{Z_{i\ell}\}_{\ell=1}^m$. Since W_i , $\bar{Z}_{i\bullet}$ and $s_{Z_i}^2$ are mutually independent for each fixed i , we can simplify the representation of $\hat{D}_L^+$ to

$$\hat{r}_L^+ \stackrel{\mathcal{D}}{\approx} \frac{(m+1)\varkappa}{m(m-1)k} \sum_{i=1}^k H_i^2 \quad \text{and} \quad \hat{D}_L^+ \stackrel{\mathcal{D}}{\approx} \frac{(m-1)\{m + (m+1)\varkappa\} \sum_{i=1}^k G_i^2}{m(m-1)k + (m+1)\varkappa \sum_{i=1}^k H_i^2},$$

where $G_i^2 \stackrel{\text{iid}}{\sim} \chi_1^2$ and $H_i^2 \stackrel{\text{iid}}{\sim} \chi_{m-1}^2$, for $i = 1, \dots, k$, are all mutually independent. Clearly, they can be further simplified to (2.14). $\square$

PROOF OF THEOREM 2.6. Similar to (B.9) and (B.10), we can have a more general representation:

$$\hat{\psi}_{\text{obs}} \stackrel{\mathcal{D}}{\approx} \psi^\star + \mathcal{T}_\psi^{1/2} W; \quad \hat{\psi}^\ell \stackrel{\mathcal{D}}{\approx} \hat{\psi}_{\text{obs}} + \mathcal{B}_\psi^{1/2} Z_\ell, \quad \ell = 1, \dots, m,$$

where $Z_1, \dots, Z_h, W \stackrel{\text{iid}}{\sim} \mathcal{N}_h(0, I_h)$. Also write $Z_\ell = (Z_{1\ell}, \dots, Z_{h\ell})^\top$, for $\ell = 1, 2, \dots, m$, and $W = (W_1, \dots, W_h)^\top$. Using (B.7), we have

$$\begin{aligned}\bar{\delta}_L - \hat{\delta}_L &\simeq \text{tr} \left\{ \bar{I}(\hat{\psi}^\star) \frac{1}{m} \sum_{\ell=1}^m \left(\hat{\psi}^\ell - \bar{\psi} \right) \left(\hat{\psi}^\ell - \bar{\psi} \right)^\top \right\} \\ &\stackrel{\mathcal{D}}{\approx} \text{tr} \left[\mathcal{U}_\psi^{-1} \frac{1}{m} \sum_{\ell=1}^m \left\{ (\mathcal{T}_\psi - \mathcal{U}_\psi)^{1/2} (Z_\ell - \bar{Z}_\bullet) \right\}^{\otimes 2} \right] \\ &= \frac{1}{m} \sum_{\ell=1}^m \text{tr} \left\{ \varkappa I_h (Z_\ell - \bar{Z}_\bullet)^{\otimes 2} \right\} = \frac{\varkappa}{m} \sum_{\ell=1}^m \sum_{i=1}^h (Z_{i\ell} - \bar{Z}_{i\bullet})^2.\end{aligned}$$

Equivalently, we can say $\bar{\delta}_L - \hat{\delta}_L \Rightarrow \kappa \chi_{h(m-1)}^2 / m$ as $n \rightarrow \infty$. Hence

$$\hat{r}_L^\diamond \Rightarrow \kappa \cdot \frac{m+1}{hm(m-1)} \cdot \chi_{h(m-1)}^2,$$

which is equivalent to (2.15). Note that it is true under both H_0 and H_1 . $\square$

PROOF OF THEOREM 2.7. From the representations of $\hat{d}_L^\diamond$ and $\hat{r}_L^\diamond$ in Lemma 2.5 and Theorem 2.6, we know that they are asymptotically ($n \rightarrow \infty$) independent. The proof then follows the derivation for Lemma 2.5. $\square$

PROOF OF PROPOSITION 3.1. It is trivial. $\square$

PROOF OF THEOREM A.1. (i) Using the representation (A.1), we can easily see that $\hat{r}_L^\Delta \geq 0$. (ii) It suffices to show

$$m^{-1} \sum_{\ell=1}^m d_L(\hat{\psi}_0^\ell + \Delta_m, \hat{\psi}^\ell \mid X^\ell) \simeq \bar{d}_L - \tilde{d}_L,$$

where $\Delta_m = \hat{\psi}^* - \hat{\psi}_0^*$. Under H_0 , $\Delta_m \simeq 0$ and $\hat{\psi}_0^\ell \simeq \hat{\psi}^\ell$, so $\hat{\psi}_0^\ell + \Delta_m \simeq \hat{\psi}^\ell$. Using Taylor's expansion on $\psi \mapsto L^\ell(\psi)$ around its maximizer $\hat{\psi}^\ell$, we have for $\psi \simeq \hat{\psi}^\ell$ that

$$L^\ell(\psi) \simeq L^\ell(\hat{\psi}^\ell) - \frac{1}{2} (\psi - \hat{\psi}^\ell)^\top I^\ell(\hat{\psi}^\ell) (\psi - \hat{\psi}^\ell).$$

Under the parametrization of ψ in the proof of Theorem 2.1, we know that the upper $k \times k$ sub-matrix of $I^\ell(\hat{\psi}^\ell)$ is $(U^\ell)^{-1}$. Using the lower order variability of U^ℓ , we have $(U^\ell)^{-1} \simeq \bar{U}^{-1}$ and

$$\begin{aligned} & \frac{1}{m} \sum_{\ell=1}^m d_L(\hat{\psi}_0^\ell + \Delta_m, \hat{\psi}^\ell \mid X^\ell) \\ & \simeq \frac{1}{m} \sum_{\ell=1}^m (\hat{\psi}_0^\ell + \Delta_m - \hat{\psi}^\ell)^\top I^\ell(\hat{\psi}^\ell) (\hat{\psi}_0^\ell + \Delta_m - \hat{\psi}^\ell) \\ & \simeq \frac{1}{m} \sum_{\ell=1}^m (\hat{\theta}^\ell - \bar{\theta})^\top \bar{U}^{-1} (\hat{\theta}^\ell - \bar{\theta}) = \bar{d}'_W - \tilde{d}'_W \simeq \bar{d}_L - \hat{d}_L. \end{aligned}$$

Therefore, the desired result follows. $\square$

PROOF OF THEOREM A.2. Throughout this proof, conditions (a), (b) and (c) refer to the list given in Assumption 5. (i, ii) It trivially follows

from the definitions of $\hat{d}_S$ and $\hat{r}_S$. (iii) First, by the definition of maximizer and condition (a), we have

$$\begin{aligned}\underline{\bar{L}}(\hat{\psi}^*) - \underline{\bar{L}}(\hat{\psi}^S) &= \underline{\bar{L}}(\hat{\psi}^*) - \underline{\bar{L}}^S(\hat{\psi}^S) + \underline{\bar{L}}^S(\hat{\psi}^S) - \underline{\bar{L}}(\hat{\psi}^S) \\ &\leq \underline{\bar{L}}(\hat{\psi}^*) - \underline{\bar{L}}^S(\hat{\psi}^*) + \underline{\bar{L}}^S(\hat{\psi}^S) - \underline{\bar{L}}(\hat{\psi}^S) \\ &\leq 2 \sup_{\psi \in \Psi} |\underline{\bar{L}}(\psi) - \underline{\bar{L}}^S(\psi)| = O_p(1/n),\end{aligned}$$

which, together with condition (b), imply that

$$\begin{aligned}&\underline{\mathcal{L}}(\psi^*) - \underline{\mathcal{L}}(\hat{\psi}^S) \\ &= \{\underline{\mathcal{L}}(\psi^*) - \underline{\bar{L}}(\psi^*)\} + \{\underline{\bar{L}}(\psi^*) - \underline{\bar{L}}(\hat{\psi}^S)\} + \{\underline{\bar{L}}(\hat{\psi}^S) - \underline{\mathcal{L}}(\hat{\psi}^S)\} \\ (B.11) \quad &\leq 2 \sup_{\psi \in \Psi} |\underline{\bar{L}}(\psi) - \underline{\mathcal{L}}(\psi)| + \{\underline{\bar{L}}(\hat{\psi}^*) - \underline{\bar{L}}(\hat{\psi}^S)\} = o_p(1).\end{aligned}$$

Using (B.11) and (c), we have $\hat{\psi}^S \xrightarrow{\text{pr}} \psi^*$. By (b) and (c), we also have $\hat{\psi}^* \xrightarrow{\text{pr}} \psi^*$. So, $|\hat{\psi}^S - \hat{\psi}^*| \xrightarrow{\text{pr}} \mathbf{0}$ as $n \rightarrow \infty$. By the definition of maximizer,

$$(B.12) \quad \mathbf{0} = \nabla \underline{\bar{L}}^S(\hat{\psi}^S) = \nabla \underline{\bar{L}}(\hat{\psi}^S) + \nabla R(\hat{\psi}^S),$$

where $\nabla g(\psi) = \partial g(\psi)/\partial \psi$ is the gradient of $\psi \mapsto g(\psi)$. By condition (a), we know $\nabla R(\hat{\psi}^S) = O_p(1/n)$. Thus, together with (B.12), we have $\nabla \underline{\bar{L}}(\hat{\psi}^S) = O_p(1/n)$. Also, by the definition of MLE, we have $\nabla \underline{\bar{L}}(\hat{\psi}^*) = \mathbf{0}$.

By Taylor's expansion, there exists $\check{\psi}$ such that

$$(B.13) \quad \underline{\bar{L}}(\hat{\psi}^*) - \underline{\bar{L}}(\hat{\psi}^S) = \left\{ \nabla \underline{\bar{L}}(\check{\psi}) \right\}^\top (\hat{\psi}^* - \hat{\psi}^S) = o_p(1/n),$$

where we have used the continuity of $\psi \mapsto \nabla \underline{\bar{L}}(\psi)$ to yield $\nabla \underline{\bar{L}}(\check{\psi}) = O_p(1/n)$. Rewriting (B.13), we have

$$(B.14) \quad \underline{\bar{L}}(\hat{\psi}^*) - \underline{\bar{L}}^S(\hat{\psi}^S) = R(\hat{\psi}^S) + o_p(1/n).$$

Similar to (B.14), we have

$$(B.15) \quad \underline{\bar{L}}(\hat{\psi}_0^*) - \underline{\bar{L}}^S(\hat{\psi}_0^S) = R(\hat{\psi}_0^S) + o_p(1/n).$$

Then, using (B.14) and (B.15), we have

$$\begin{aligned}|\hat{d}_L - \hat{d}_S| &= 2n \left| \left\{ \underline{\bar{L}}(\hat{\psi}^*) - \underline{\bar{L}}^S(\hat{\psi}^S) \right\} - \left\{ \underline{\bar{L}}(\hat{\psi}_0^*) - \underline{\bar{L}}^S(\hat{\psi}_0^S) \right\} \right| \\ &= 2n \left| R(\hat{\psi}^S) - R(\hat{\psi}_0^S) + o_p(1/n) \right|.\end{aligned}$$

Now consider two cases.

- (i) Under H_0 , we have $\hat{d}_L = O_p(1)$ and $\hat{\psi}_0^S \simeq \hat{\psi}^S$. Thus condition (a) implies $R(\hat{\psi}^S) - R(\hat{\psi}_0^S) = o_p(1/n)$. Then, we have $|\hat{d}_L - \hat{d}_S| = o_p(\hat{d}_L)$.
- (ii) Under H_1 , we have $\hat{d}_L \xrightarrow{\text{pr}} \infty$. Condition (a) and (B.11) imply that $\bar{L}(\hat{\psi}^*) - \bar{L}^S(\hat{\psi}^S) = O_p(1/n)$. Similarly, we also have $\bar{L}(\hat{\psi}_0^*) - \bar{L}^S(\hat{\psi}_0^S) = O_p(1/n)$. Hence $|\hat{d}_L - \hat{d}_S| = O_p(1)$. Thus we have $|\hat{d}_L - \hat{d}_S| = o_p(\hat{d}_L)$.

Therefore, under either H_0 or H_1 , we also have $|\hat{d}_L - \hat{d}_S| = o_p(\hat{d}_L)$. Since $\hat{d}_L \simeq \hat{d}_S$ and $\bar{d}_L = \bar{d}_S$, we know $\hat{r}_L \simeq \hat{r}_S$. $\square$

Note that, even under the assumption of this theorem, $\hat{r}_S$ and $\hat{r}_S^\diamond$ are not equivalent. From (3.5) and (3.6), $\hat{r}_S$ and $\hat{r}_S^\diamond$ are a “difference of difference” estimator and a “difference” estimator, respectively. So, the “bias” of using $\bar{L}^S(\psi)$ cannot be canceled out in $\hat{r}_S^\diamond$.

APPENDIX C: ADDITIONAL FIGURES AND TABLES

This section presents additional figures and tables in § 2.5 and § 4

- Figure 9: the performance of different approximations to the reference null distribution when $\alpha = 5\%$; see § 2.5.
- Figures 10 and 11: the empirical distributions of the p -values under H_0 and parametrizations (i) and (iii), respectively; see § 4.1.
- Figures 12 and 13: the empirical power functions and the empirical ratio of power-to-size for size 5% tests, respectively; see § 4.1.
- Table 8: the ranges of empirical sizes over different parametrizations for size 5% tests; see § 4.1.
- Table 9: detailed results of the care-survival example in § 4.3.

DEPARTMENT OF STATISTICS
THE CHINESE UNIVERSITY OF HONG KONG
SHATIN, N.T., HONG KONG
E-MAIL: kinwaichan@sta.cuhk.edu.hk

DEPARTMENT OF STATISTICS
HARVARD UNIVERSITY
CAMBRIDGE, MASSACHUSETTS 02138, U.S.A.
E-MAIL: meng@stat.harvard.edu

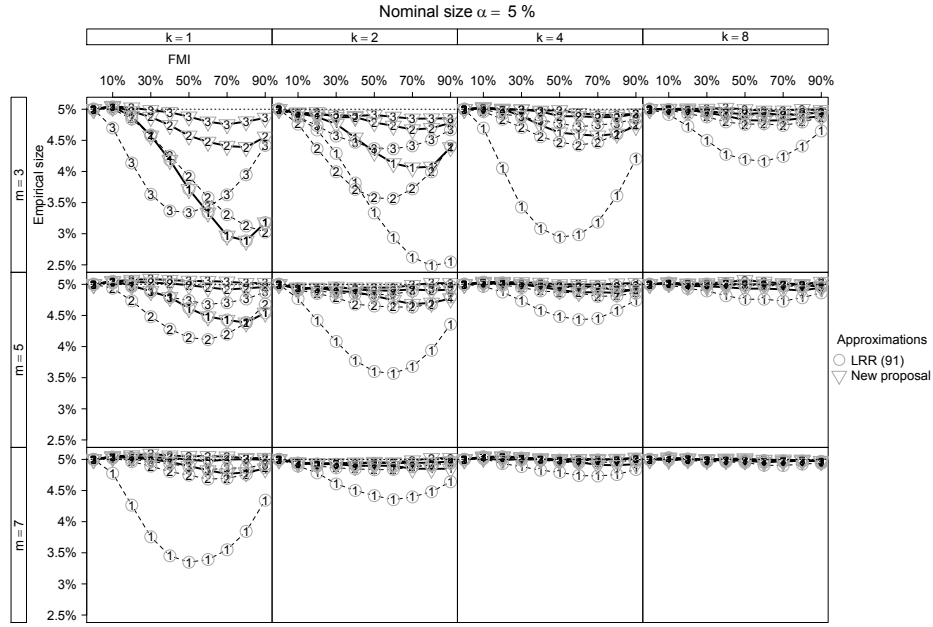


FIG 9. The performances of two different approximate null distributions when the nominal size is $\alpha = 5\%$. The vertical axis denotes $\hat{\alpha}$ or $\tilde{\alpha}$, and the horizontal axis denotes the value of ℓ_m . The number attached to each line denotes the value of $\tau = h/k$. The proposed approximation $\hat{\alpha}$ is denoted by thick solid lines with triangles, and the existing approximation $\tilde{\alpha}$ is denoted by thin dashed lines with circles.

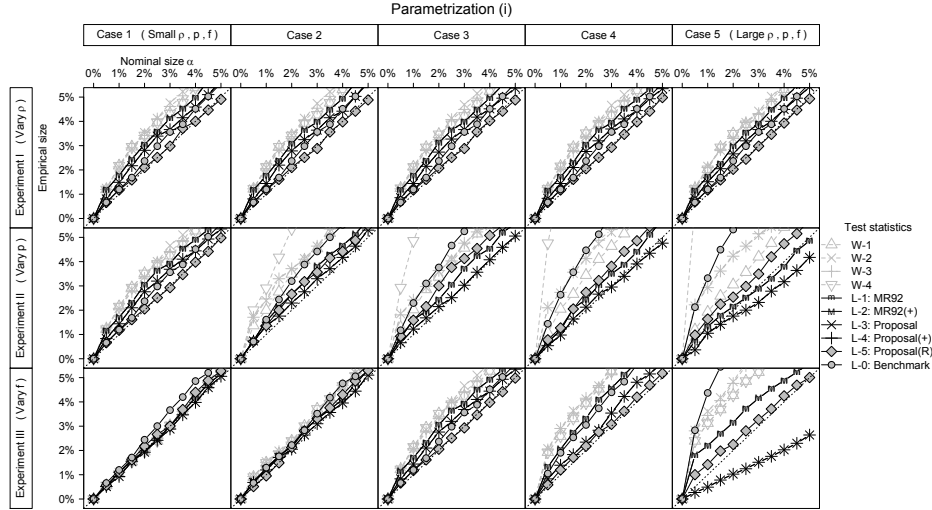


FIG 10. The comparison between empirical size and nominal size α under parametrization (i) for $\alpha \in (0, 5\%]$. The Wald tests ($W-1, 2, 3, 4$) and LRTs ($L-0, 1, 2, 3, 4, 5$) are represented by grey dashed and black solid lines, respectively. The LRT statistic $\hat{D}_L$ ($L-1$: MR92) (Meng and Rubin, 1992) and its modification $\hat{D}_L^+$ ($L-2$: MR92(+)) are the existing counterparts of our proposals $\hat{D}_S$ ($L-3$: Proposal), $\hat{D}_S^+$ ($L-4$: Proposal(+)) and $\hat{D}_S^\diamond$ ($L-5$: Proposal(R)).

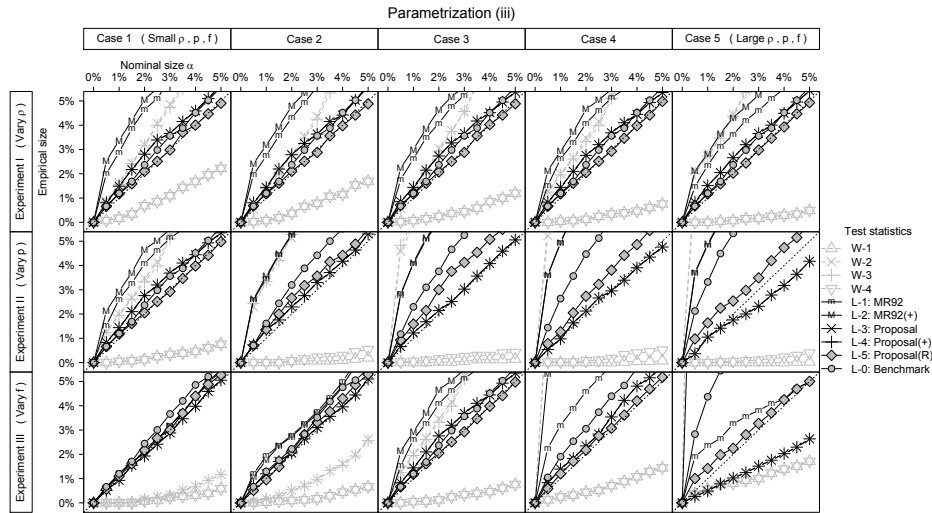


FIG 11. The comparison between empirical size and nominal size α under parametrization (iii) for $\alpha \in (0, 5\%]$. The legend in Figure 10 also applies here.

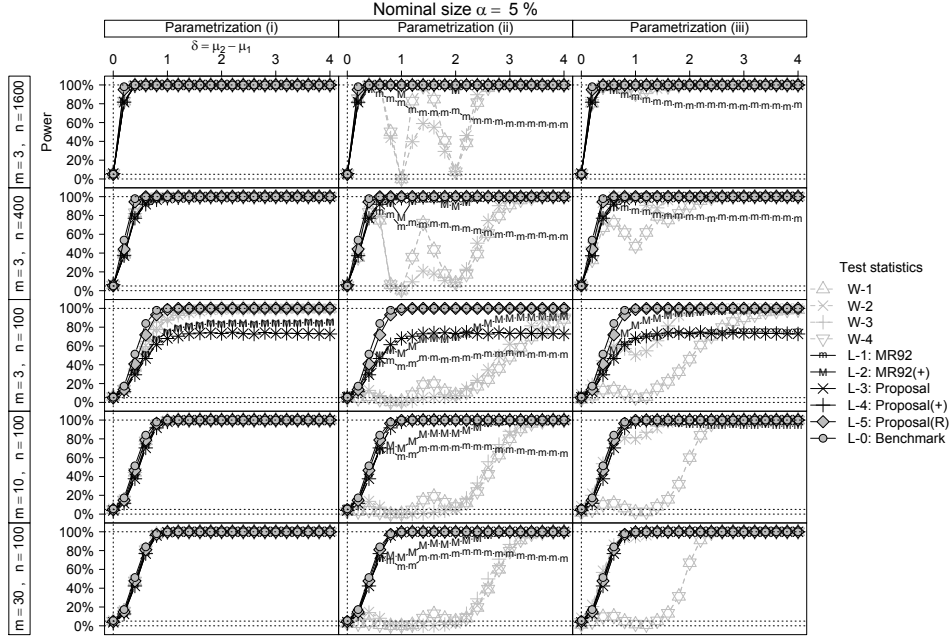


FIG 12. The power curves under different parametrizations. The nominal size is $\alpha = 5\%$. In each plot, the vertical axis denotes the power, whereas the horizontal axis denotes the value of $\delta = \mu_2 - \mu_1$. The legend in Figure 10 also applies here.

TABLE 8

The range of empirical size $[\min \hat{\alpha}, \max \hat{\alpha}]$ in percentage, where max and min are taken over the three parametrizations. Only one value is recorded for those tests that are invariant to parametrization. The nominal size is $\alpha = 5\%$.

(n, m)	Range of empirical size: $[\min \hat{\alpha}, \max \hat{\alpha}]/\%$				
	(1600, 3)	(400, 3)	(100, 3)	(100, 10)	(100, 30)
W-1	[5.62, 5.71]	[5.30, 6.03]	[3.22, 6.20]	[1.64, 4.81]	[1.37, 5.00]
W-2	[5.93, 6.05]	[6.08, 7.18]	[5.52, 8.69]	[4.42, 8.47]	[4.20, 8.50]
W-3	[5.81, 6.03]	[6.01, 6.98]	[5.37, 8.28]	[4.20, 7.67]	[4.10, 7.50]
W-4	[5.62, 5.71]	[5.30, 6.03]	[3.22, 6.20]	[1.64, 4.81]	[1.37, 5.00]
L-1	[5.57, 6.15]	[6.37, 6.57]	[5.88, 6.47]	[4.39, 5.66]	[4.22, 5.32]
L-2	[5.52, 6.10]	[6.37, 6.52]	[5.88, 7.47]	[4.39, 5.66]	[4.22, 5.32]
L-3	5.76	6.37	5.42	3.78	3.71
L-4	5.76	6.37	5.42	3.78	3.71
L-5	4.96	5.32	4.93	4.79	4.54
L-0	5.03	5.03	5.57	5.57	5.57

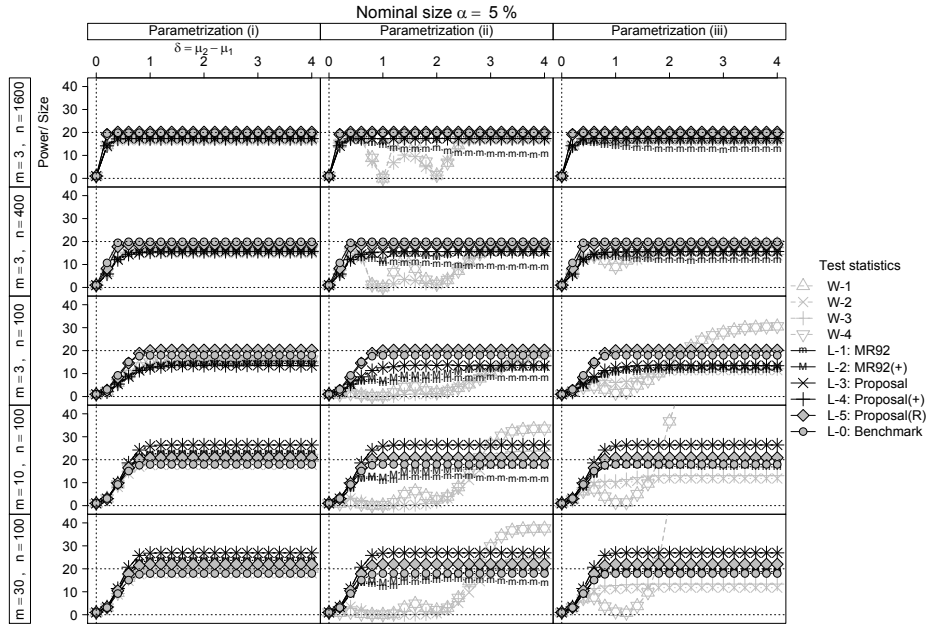


FIG 13. The ratios of empirical power to empirical size under different parametrizations. The nominal size is $\alpha = 5\%$. In each plot, the vertical axis denotes the ratio, whereas the horizontal axis denotes $\delta = \mu_2 - \mu_1$. The legend in Figure 10 also applies here.

TABLE 9
The LRTs using $\tilde{D}_L$, $\hat{D}_S^+$ and $\hat{D}_S^\diamond$ under different parametrizations in § 4.3.

Parametrization (i)						
m	H_0 : Conditional independence			H_0 : Full independence		
	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$
2	0.63, 0.64, 0.83	0.14, 0.14, 0.12	0.87, 0.87, 0.89	0.53, 0.53, 0.83	44.4, 44.4, 37.1	0, 0, 0
3	0.54, 0.54, 0.38	0.08, 0.08, 0.09	0.93, 0.93, 0.92	0.31, 0.31, 0.38	54.2, 54.2, 51.4	0, 0, 0
5	0.49, 0.48, 0.89	0.12, 0.12, 0.10	0.89, 0.89, 0.91	0.72, 0.72, 0.89	40.8, 40.8, 37.1	0, 0, 0
7	0.23, 0.23, 0.47	0.06, 0.06, 0.05	0.94, 0.94, 0.95	0.31, 0.31, 0.47	53.2, 53.2, 47.6	0, 0, 0
10	0.50, 0.50, 0.70	0.14, 0.14, 0.12	0.87, 0.87, 0.88	0.56, 0.56, 0.70	45.4, 45.4, 41.7	0, 0, 0
25	0.35, 0.35, 0.47	0.06, 0.06, 0.06	0.94, 0.94, 0.95	0.35, 0.35, 0.47	51.4, 51.4, 47.0	0, 0, 0
50	0.31, 0.31, 0.45	0.11, 0.11, 0.10	0.90, 0.90, 0.91	0.33, 0.33, 0.45	51.5, 51.5, 47.3	0, 0, 0
Parametrization (ii)						
m	H_0 : Conditional independence			H_0 : Full independence		
	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$
2	1.23, 0.64, 0.83	0.01, 0.14, 0.12	0.99, 0.87, 0.89	0.98, 0.53, 0.83	34.2, 44.4, 37.1	0, 0, 0
3	1.08, 0.54, 0.38	-0.07, 0.08, 0.09	1.00, 0.93, 0.92	0.61, 0.31, 0.38	43.9, 54.2, 51.4	0, 0, 0
5	1.02, 0.48, 0.89	-0.09, 0.12, 0.10	1.00, 0.89, 0.91	1.40, 0.72, 0.89	29.0, 40.8, 37.1	0, 0, 0
7	0.45, 0.23, 0.47	-0.07, 0.06, 0.05	1.00, 0.94, 0.95	0.58, 0.31, 0.47	43.9, 53.2, 47.6	0, 0, 0
10	0.99, 0.50, 0.70	-0.10, 0.14, 0.12	1.00, 0.87, 0.88	1.09, 0.56, 0.70	33.7, 45.4, 41.7	0, 0, 0
25	0.71, 0.35, 0.47	-0.14, 0.06, 0.06	1.00, 0.94, 0.95	0.68, 0.35, 0.47	41.0, 51.4, 47.0	0, 0, 0
50	0.63, 0.31, 0.45	-0.10, 0.11, 0.10	1.00, 0.90, 0.91	0.65, 0.33, 0.45	41.3, 51.5, 47.3	0, 0, 0
Parametrization (iii)						
m	H_0 : Conditional independence			H_0 : Full independence		
	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$	$\tilde{r}_L, \hat{r}_S^+, \hat{r}_S^\diamond$	$\tilde{D}_L, \hat{D}_S^+, \hat{D}_S^\diamond$	$\tilde{p}_L, \hat{p}_S^+, \hat{p}_S^\diamond$
2	1.06, 0.64, 0.83	0.04, 0.14, 0.12	0.96, 0.87, 0.88	-0.38, 0.53, 0.83	109, 44.4, 37.1	0, 0, 0
3	-2.35, 0.54, 0.38	-1.16, 0.08, 0.09	1.00, 0.93, 0.92	-1.22, 0.31, 0.38	-321, 54.2, 51.4	1, 0, 0
5	-2.64, 0.48, 0.89	-1.38, 0.12, 0.10	1.00, 0.89, 0.91	-2.24, 0.72, 0.89	-58, 40.8, 37.1	1, 0, 0
7	-0.01, 0.23, 0.47	0.25, 0.06, 0.05	0.78, 0.94, 0.95	-0.34, 0.31, 0.47	107, 53.2, 47.6	0, 0, 0
10	-2.04, 0.50, 0.70	-2.20, 0.14, 0.12	1.00, 0.87, 0.88	-1.85, 0.56, 0.70	-86, 45.4, 41.7	1, 0, 0
25	-1.39, 0.35, 0.47	-4.30, 0.06, 0.06	1.00, 0.94, 0.95	-1.12, 0.35, 0.47	-603, 51.4, 47.0	1, 0, 0
50	-1.22, 0.31, 0.45	-7.39, 0.11, 0.10	1.00, 0.90, 0.91	-1.06, 0.33, 0.45	-1136, 51.5, 47.3	1, 0, 0